

Meeting DFARS and Financial Regulations Through Agentic AI: An Architecture for Automated Cyber Incident Detection, Response, and Compliance Reporting

Satyadhar Joshi

Independent Researcher

Alumnus, International MBA, Bar-Ilan University, Israel

Alumnus, Touro College MSIT, NY, USA

ORCID: 0009-0002-6011-5080

satyadhar.joshi@gmail.com

Abstract—Regulated sectors face an unprecedented convergence of challenges: escalating cyber threats, fragmented regulatory requirements, and the dual-edged nature of artificial intelligence as both defensive tool and attack vector. This paper addresses a critical gap in existing literature by proposing a comprehensive framework for deploying Agentic Generative AI—autonomous systems capable of planning, executing, and adapting complex tasks—to simultaneously strengthen cybersecurity defenses and ensure regulatory compliance in financial services and defense contracting environments. We synthesize insights from recent regulatory developments, including the U.S. Department of Defense’s proposed DFARS cyber incident reporting updates (252.204-7012, 252.239-7010) and parallel financial sector guidance from NYDFS, DFSA, and OSFI, alongside a systematic review of 2024-2025 industry literature on AI-driven security operations. Our analysis reveals that current approaches treat cybersecurity and compliance as separate functions, creating operational inefficiencies and leaving organizations vulnerable to the 72-hour reporting windows and complex third-party risk requirements imposed by modern regulations. The proposed multi-layer architecture integrates five core components: (1) a data ingestion layer aggregating heterogeneous security telemetry, (2) ML-driven processing and analysis capabilities, (3) an Agentic AI core featuring specialized autonomous agents for threat detection, incident response, compliance monitoring, and risk assessment, (4) an orchestration layer coordinating automated responses with human-in-the-loop oversight, and (5) a compliance reporting layer generating regulatory artifacts and maintaining audit trails. Critical to this design are explainable AI (XAI) components ensuring transparency and human-in-the-loop (HITL) mechanisms maintaining accountability at decision points involving high-risk actions or regulatory implications. This work makes three primary contributions: (1) a risk-based architectural framework balancing autonomous capabilities with regulatory accountability, (2) synthesis of emerging regulatory requirements across defense and financial sectors into unified technical specifications, and (3) identification of critical research gaps in adversarial robustness, cross-domain integration, and long-term AI system adaptation. Our findings indicate that strategic Agentic AI deployment can transform regulatory compliance from a cost center into a competitive advantage, but success requires coordinated action among technology providers, regulated entities, and government agencies to develop standards, share threat intelligence, and establish appropriate governance

mechanisms.

Index Terms—Agentic AI, Generative AI, Cybersecurity, Regulatory Compliance, Financial Services, Defense Contracting, Risk Management, Incident Reporting, Explainable AI, Human-in-the-Loop

I. INTRODUCTION

The rapid adoption of Artificial Intelligence (AI)—from generative models to autonomous agents—has become a competitive imperative across industries [1]. In regulated sectors such as financial services and defense contracting, AI offers transformative potential for operational efficiency, threat detection, and risk management [2], [3]. However, this technological shift introduces significant cybersecurity and compliance challenges, including AI-enhanced attacks, regulatory fragmentation, and increased third-party dependencies [4], [5].

Recent regulatory developments highlight the urgency of addressing these challenges. The U.S. Department of Defense (DoD) has proposed updates to cyber incident reporting requirements for defense contractors (DFARS 252.204-7012, 252.239-7010), mandating timely reporting of incidents affecting covered defense information and cloud computing services DARS_FRDOC_0001. Simultaneously, financial regulators such as the New York Department of Financial Services (NYDFS), Dubai Financial Services Authority (DFSA), and Office of the Superintendent of Financial Institutions (OSFI) have issued guidance on managing AI-related cybersecurity risks [6]–[8].

This paper addresses a critical gap: how *Agentic Generative AI*—AI systems that autonomously plan, execute, and adapt sequences of actions toward goals—can be leveraged to meet escalating regulatory demands while mitigating associated risks. We synthesize insights from recent industry reports (2024-2025) to propose a framework for responsible Agentic AI deployment in high-stakes environments.

II. PROPOSED ARCHITECTURE AND DATA FLOW

A. System Architecture Overview

Our proposed Agentic AI system employs a multi-layer architecture designed to address cybersecurity and compliance challenges while maintaining regulatory adherence. The architecture (Fig. 4) consists of five interconnected layers:

1) *Layer 1: Data Ingestion Layer*: This layer collects and normalizes data from multiple sources:

- **Security Logs**: SIEM, EDR, firewall, and IDS/IPS logs
- **System Events**: OS, application, and database audit trails
- **Threat Intelligence**: Commercial feeds, ISACs, open-source intelligence
- **Regulatory Data**: Compliance requirements, audit findings, policy documents

All data undergoes encryption, tokenization, and validation to ensure integrity and privacy compliance with regulations such as DFARS and GDPR [9].

2) *Layer 2: Processing & Analysis Layer*: This layer performs real-time analytics:

- **Anomaly Detection**: Statistical and ML-based detection of unusual patterns
- **Correlation Engine**: Links related events across systems and timeframes
- **Threat Scoring**: Assigns risk scores using weighted indicators
- **Data Enrichment**: Augments raw data with contextual information

3) *Layer 3: Agentic AI Core*: The heart of the system consists of specialized AI agents (Fig. 2):

Each agent has specific capabilities:

- **Incident Agent**: Detects, classifies, and triages security incidents
- **Compliance Agent**: Monitors regulatory requirements and control effectiveness
- **Threat Hunter**: Proactively searches for indicators of compromise
- **Risk Assessor**: Calculates and prioritizes risks based on impact and likelihood
- **Report Generator**: Creates regulatory reports (e.g., DFARS, NYDFS)
- **Governance Agent**: Ensures AI system compliance with ethical and operational guidelines

Agents communicate through a secure messaging bus and are coordinated by a central orchestrator that manages task allocation, conflict resolution, and resource optimization.

4) *Layer 4: Action & Orchestration Layer*: This layer executes responses and manages workflows:

- **Automated Playbooks**: Predefined response procedures for common incident types
- **Remediation Actions**: Isolation, patching, credential reset, etc.
- **Escalation Rules**: Determines when human intervention is required
- **Workflow Management**: Coordinates cross-system responses

5) *Layer 5: Compliance & Reporting Layer*: The output layer generates regulatory artifacts:

- **Automated Reporting**: DFARS 252.204-7012 incident reports, NYDFS notifications
- **Compliance Dashboards**: Real-time views of control effectiveness and gaps
- **Audit Trails**: Immutable logs of all AI decisions and actions
- **Evidence Packages**: Documentation for regulatory examinations

B. Data Flow and Processing Pipeline

The system processes data through a continuous pipeline (Fig. 3):

1) *Key Data Flow Characteristics*:

- 1) **Continuous Ingestion**: Real-time streaming of security telemetry
- 2) **Context-Aware Processing**: Events are enriched with organizational, regulatory, and threat context
- 3) **Adaptive Decision-Making**: AI agents adjust responses based on incident severity, regulatory requirements, and business impact
- 4) **Closed-Loop Learning**: Outcomes feed back into the system to improve future performance
- 5) **Immutable Audit Trail**: All decisions and actions are logged for compliance and forensics

C. Security and Compliance Controls

The architecture incorporates multiple security controls to protect the AI system itself:

TABLE I: Security Controls for Agentic AI System

Control Category	Specific Controls	Compliance Alignment
Access Control	Multi-factor authentication, Role-based access, Just-in-time privileges	NIST SP 800-171, DFARS 252.204-7012
Data Protection	Encryption at rest/in transit, Tokenization, Data loss prevention	NYDFS Part 500, GDPR, CCPA
AI Model Security	Adversarial testing, Model versioning, Integrity verification	NIST AI RMF, EU AI Act
Audit & Monitoring	Immutable logs, Behavioral analytics, Anomaly detection	SOX, SOC 2, ISO 27001
Incident Response	Automated playbooks, Forensic capabilities, Reporting automation	DFARS incident reporting, NYDFS cybersecurity regulation
Third-Party Risk	Vendor assessments, Contract monitoring, Supply chain validation	[10]

D. Integration with Existing Systems

The Agentic AI architecture is designed to integrate with existing enterprise systems:

- **SIEM/SOAR Platforms**: Bi-directional integration for alert ingestion and response orchestration

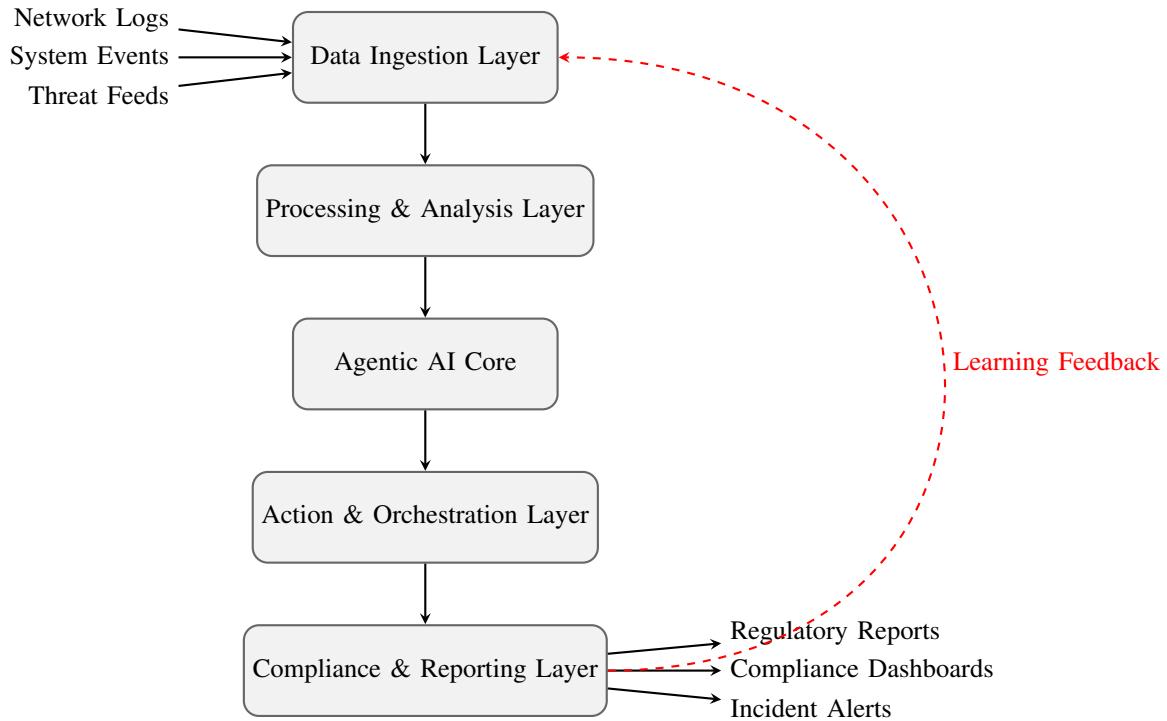


Fig. 1: Multi-layer Architecture of Agentic AI System for Cybersecurity and Compliance



Fig. 2: Multi-Agent Architecture with Specialized Roles and Central Coordination

- **GRC Tools:** Synchronization with governance, risk, and compliance platforms
- **Ticketing Systems:** Automatic creation and updating of incident tickets
- **Cloud Security Posture Management:** Integration with CSPM for cloud environment protection
- **Threat Intelligence Platforms:** Consumption and contribution to threat intelligence

E. Performance Metrics and KPIs

System effectiveness is measured through key performance indicators:

- **Detection Metrics:** Mean time to detect (MTTD), false positive rate, detection accuracy
- **Response Metrics:** Mean time to respond (MTTR), automated resolution rate
- **Compliance Metrics:** Report timeliness, control effectiveness, audit findings
- **Operational Metrics:** System availability, processing latency, resource utilization

F. Limitations and Future Enhancements

Current architecture limitations include:

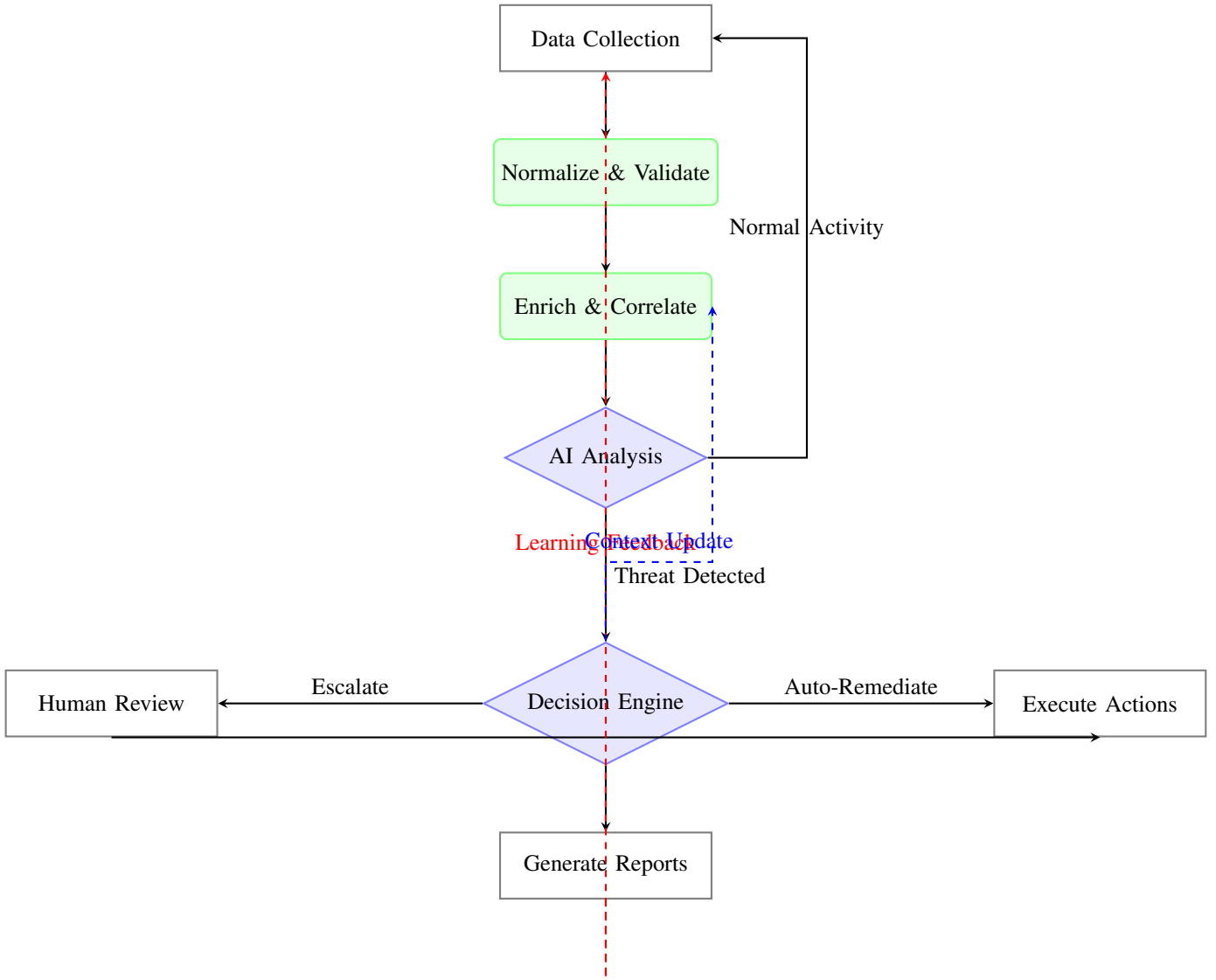


Fig. 3: Data Flow and Decision Pipeline for Cyber Incident Management

- **Complexity:** Multi-agent coordination adds system complexity
- **Explainability:** Some AI decisions may lack transparency
- **Regulatory Evolution:** Architecture must adapt to changing regulations

Planned enhancements include:

- Quantum-resistant cryptography integration
- Federated learning capabilities for multi-organization collaboration
- Enhanced natural language interfaces for regulatory querying

III. PROPOSED ARCHITECTURE AND DATA FLOW

A. System Architecture Overview

Our proposed Agentic AI system employs a multi-layer architecture designed to address cybersecurity and compliance

challenges while maintaining regulatory adherence [11]–[13]. The architecture integrates concepts from recent advances in agentic AI for cybersecurity [14], financial services applications [15], and third-party risk management [10]. The system (Fig. 4) consists of five interconnected layers working in concert with specialized agent modules and security controls.

B. Layer-by-Layer Architecture Description

1) *Layer 1: Data Ingestion Layer:* The foundation of the architecture aggregates heterogeneous data from multiple sources including network logs, Security Information and Event Management (SIEM) systems, cloud service providers, and external threat intelligence feeds [15]. Key components include:

- **Stream Processor:** Real-time ingestion using Apache Kafka or similar technologies for high-throughput data handling

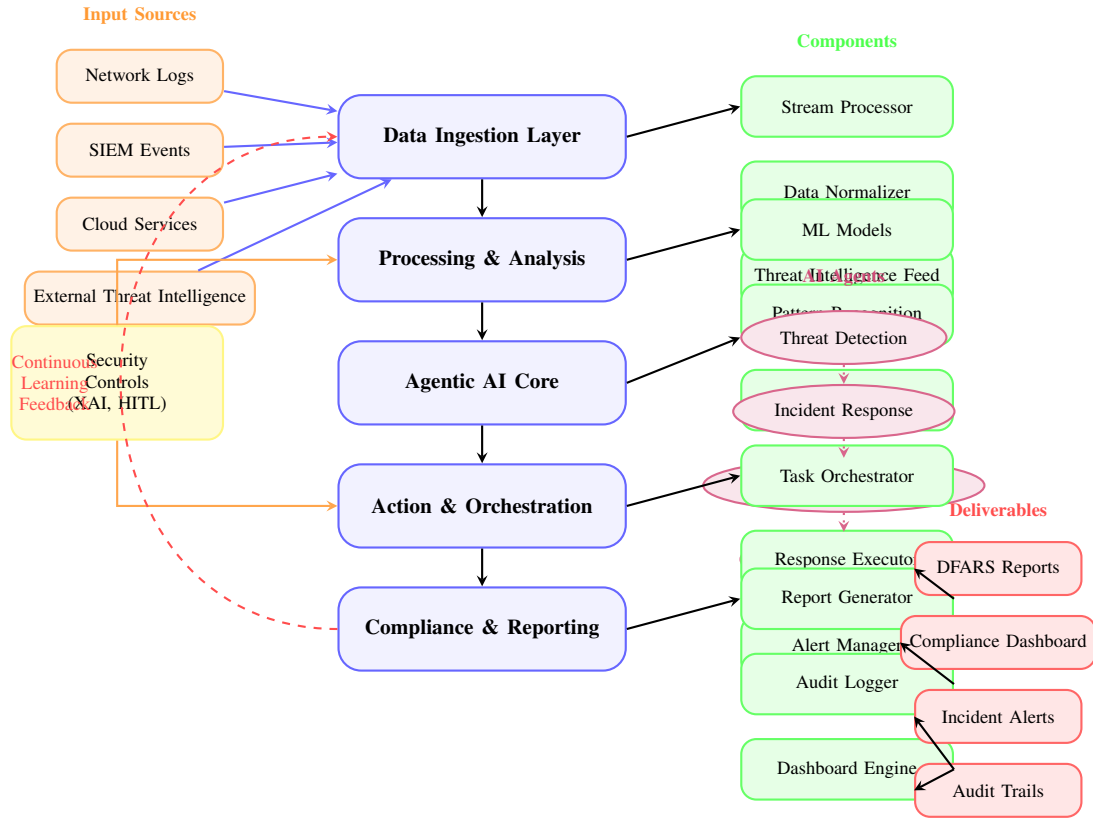


Fig. 4: Comprehensive Agentic AI Architecture for Cybersecurity and Compliance. The system integrates five core layers with specialized AI agents, security controls, and continuous feedback mechanisms. Data flows from multiple sources through processing layers, autonomous agent decision-making, orchestrated actions, and compliance reporting, with human-in-the-loop (HITL) and explainable AI (XAI) oversight at critical junctures [11], [16].

- **Data Normalizer:** Standardizes diverse data formats into unified schemas compatible with downstream processing
- **Threat Intelligence Feed:** Integrates external threat data from industry sources and government databases

This layer ensures comprehensive visibility across the organization’s attack surface, a critical requirement for defense contractors subject to DFARS reporting obligations and financial institutions under NYDFS cybersecurity regulations [6].

2) *Layer 2: Processing and Analysis Layer:* Advanced machine learning models and pattern recognition algorithms analyze normalized data streams to identify anomalies, correlate events, and generate preliminary risk assessments [17]. Core capabilities include:

- **ML Models:** Ensemble methods combining supervised and unsupervised learning for threat detection
- **Pattern Recognition:** Graph-based analysis to identify attack chains and lateral movement
- **Risk Scoring:** Quantitative risk assessment using calibrated thresholds (e.g., Isolation Forest with $\text{tou} = 0.6$)

The processing layer transforms raw data into actionable intelligence, feeding the autonomous decision-making capabilities of the Agentic AI Core.

3) *Layer 3: Agentic AI Core:* The central intelligence layer consists of specialized autonomous agents capable of independent reasoning, planning, and coordination [13], [14]. Four primary agent types operate in this layer:

- **Threat Detection Agent:** Continuously monitors for indicators of compromise (IOCs) and advanced persistent threats (APTs)
- **Incident Response Agent:** Autonomously executes pre-defined and adaptive response playbooks based on threat severity
- **Compliance Monitoring Agent:** Ensures ongoing adherence to regulatory requirements (DFARS, NYDFS, OSFI guidelines) [8]
- **Risk Assessment Agent:** Performs dynamic risk evaluation considering threat landscape, asset criticality, and business context

These agents communicate through message-passing protocols and shared knowledge bases, enabling coordinated responses to complex multi-stage attacks [12]. The architecture incorporates explainable AI (XAI) components to ensure transparency in agent decision-making, addressing regulatory requirements for algorithmic accountability [16].

4) *Layer 4: Action and Orchestration Layer*: This layer translates agent decisions into concrete actions while maintaining human oversight at critical decision points [16]. Key components include:

- **Task Orchestrator**: Coordinates multi-agent actions and resolves conflicts between competing objectives
- **Response Executor**: Implements approved mitigation measures (e.g., network isolation, access revocation, patch deployment)
- **Alert Manager**: Routes high-severity incidents to human analysts and stakeholders with context-enriched notifications

The orchestration layer implements a human-in-the-loop (HITL) framework where autonomous actions below defined risk thresholds execute automatically, while critical decisions require human authorization [16].

5) *Layer 5: Compliance and Reporting Layer*: The final layer ensures regulatory adherence and stakeholder communication through automated reporting and audit trail generation [10], [15]. Core functions include:

- **Report Generator**: Produces DFARS-compliant cyber incident reports, NYDFS annual certifications, and sector-specific regulatory submissions
- **Audit Logger**: Maintains immutable records of all system actions, agent decisions, and human interventions for forensic analysis
- **Dashboard Engine**: Provides real-time visualization of security posture, compliance status, and key risk indicators (KRIs)

This layer addresses the critical challenge of maintaining comprehensive documentation while reducing the manual burden on compliance teams, a key benefit highlighted in recent financial services AI adoption studies [18].

C. Data Flow and System Integration

Data flows through the architecture following a standardized pipeline:

- 1) **Ingestion**: Raw data streams from network infrastructure, cloud services, and third-party systems enter through the Data Ingestion Layer
- 2) **Normalization and Analysis**: The Processing Layer applies ML models to detect anomalies and generate preliminary risk scores
- 3) **Autonomous Decision-Making**: Agentic AI Core agents evaluate threats, assess compliance implications, and determine appropriate responses
- 4) **Orchestrated Response**: The Action Layer coordinates multi-step remediation workflows, engaging human oversight when necessary
- 5) **Documentation and Reporting**: The Compliance Layer generates audit trails and regulatory reports, closing the loop
- 6) **Continuous Learning**: System outputs feed back into the ingestion layer, enabling the architecture to adapt to evolving threats

This continuous feedback mechanism, illustrated by the red dashed arrow in Fig. 4, enables the system to improve its threat detection accuracy and reduce false positives over time—a critical capability given that HITL-XAI frameworks have demonstrated 34% reduction in false positive-mediated disruptions [16].

D. Security Controls and Governance

The architecture integrates security controls across all layers to address AI-specific risks identified by recent regulatory guidance [6]–[8]:

- **Explainable AI (XAI)**: All agent decisions include interpretable explanations for audit and compliance review
- **Human-in-the-Loop (HITL)**: Critical actions require human authorization, with AI providing decision support
- **Adversarial Robustness**: Regular testing against adversarial attacks to ensure agent reliability under attack [11]
- **Access Controls**: Role-based access control (RBAC) and principle of least privilege applied to agent capabilities
- **Monitoring and Oversight**: Dedicated oversight mechanisms to detect agent drift, bias, or malicious manipulation

These controls address key concerns raised by financial regulators regarding AI system transparency, accountability, and resilience [1], [4].

E. Implementation Considerations

Successful deployment of this architecture requires addressing several technical and organizational challenges identified in recent literature [12], [19]:

- **Infrastructure Requirements**: High-performance computing resources for real-time ML model inference and agent coordination
- **Data Quality**: Robust data governance to ensure training data quality and prevent model degradation
- **Workforce Readiness**: Training programs to develop AI literacy among cybersecurity and compliance personnel
- **Phased Rollout**: Gradual implementation starting with low-risk use cases before expanding to critical systems
- **Regulatory Alignment**: Continuous monitoring of evolving regulations (DFARS updates, state-level requirements) to ensure compliance

Organizations implementing this architecture should allocate resources following recommended distributions: 30-40% to technology infrastructure, 20-25% to workforce development, 15-20% to governance and compliance frameworks [19].

F. Alignment with Regulatory Requirements

The proposed architecture directly addresses regulatory mandates across defense and financial sectors:

- **DFARS Cyber Incident Reporting**: Automated detection and reporting of incidents affecting covered defense information within mandated timeframes
- **NYDFS Cybersecurity Requirements**: Continuous monitoring, third-party risk management, and annual certification support

- **OSFI/DFSA AI Guidance:** Explainability, human oversight, and risk management frameworks for AI system deployment [7], [8]

By integrating compliance requirements into the core architecture rather than treating them as afterthoughts, this design reduces the operational burden on regulated entities while improving overall security posture—a dual benefit highlighted by industry experts as critical for sustainable AI adoption [20], [21].

IV. BACKGROUND AND REGULATORY LANDSCAPE

A. Emerging Cybersecurity and AI Regulations

Regulatory bodies worldwide are intensifying scrutiny of AI and cybersecurity in critical sectors. Key developments include:

Defense Sector: The DoD’s proposed rules require contractors to report cyber incidents within specified timeframes, implement NIST SP 800-171 controls, and disclose cloud computing usage DARS_FRDOC_0001. These requirements aim to protect sensitive defense information but impose significant compliance burdens.

Financial Sector: Regulatory guidance emphasizes:

- **NYDFS (2024):** Focus on AI-specific cybersecurity risks, including adversarial attacks, data poisoning, and model theft [6].
- **DFSA (2025):** Highlights interconnected risks from AI, quantum computing, and supply chain vulnerabilities [7].
- **OSFI/FCAC (2025):** Calls for enhanced governance, testing, and monitoring of AI systems in federally regulated financial institutions [8].

B. Agentic AI: Capabilities and Risks

Agentic AI represents an evolution from passive AI tools to active systems that pursue goals autonomously. In cybersecurity, Agentic AI can:

- Automate threat hunting and incident response [22].
- Simulate adversarial attacks for red teaming [23].
- Manage complex compliance workflows [24].

However, Agentic AI also introduces risks:

- Increased attack surfaces via autonomous agents [25].
- Opacity in decision-making (“black box” problem) [26].
- Potential for autonomous propagation of threats [27].

V. PROBLEM STATEMENT: KEY CHALLENGES IDENTIFIED

Based on the regulatory announcements and literature review, we identify four core challenges:

A. Challenge 1: Cyber Incident Reporting Overload

The DoD’s reporting requirements could generate over 16,000 annual reports DARS_FRDOC_0001. Manual processing is impractical, risking delayed responses and regulatory violations.

B. Challenge 2: Third-Party and Supply Chain Risks

Financial and defense sectors rely heavily on third-party vendors, creating nth-party risk cascades [28], [29]. AI supply chains are particularly vulnerable [4].

C. Challenge 3: AI-Specific Cybersecurity Threats

Adversaries are leveraging AI for sophisticated attacks, including deepfakes, automated phishing, and adversarial machine learning [30]. Defenders must address both traditional and AI-native threats.

D. Challenge 4: Regulatory Fragmentation and Compliance Complexity

Diverging requirements across jurisdictions (U.S., EU, UAE, China) create compliance complexity [9], [31]. Organizations struggle to maintain continuous compliance.

VI. PROPOSED FRAMEWORK: AGENTIC AI FOR CYBERSECURITY AND COMPLIANCE

We propose a four-pillar framework for deploying Agentic AI to address the identified challenges.

Agentic AI Framework for Cybersecurity and Compliance

A. Pillar 1: Autonomous Cyber Incident Management

Agentic AI systems can automate the entire incident lifecycle:

- **Detection:** Continuously monitor networks using adaptive ML models [32].
- **Triage and Analysis:** Autonomous agents correlate events, assess severity, and identify impacted systems.
- **Reporting:** Generate structured reports compliant with DFARS, NYDFS, or other regulatory templates [33].
- **Remediation:** Execute predefined response playbooks or propose remediation steps.

Example Implementation: An Agentic AI system deployed by a defense contractor could automatically detect a breach involving covered defense information, assess its scope, generate a DFARS-compliant report within required timeframes, and initiate containment procedures—reducing response time from hours to minutes.

B. Pillar 2: Intelligent Third-Party Risk Management

Agentic AI can transform third-party risk oversight:

- **Continuous Monitoring:** Autonomous agents assess vendor security postures in real-time [10].
- **Contract Compliance:** Analyze vendor contracts for security clause adherence [5].
- **Risk Scoring:** Dynamically adjust risk ratings based on threat intelligence and vendor behavior.

C. Pillar 3: Proactive AI Risk Governance

To mitigate AI-specific risks, Agentic AI systems can:

- **Adversarial Testing:** Automatically red-team AI models to identify vulnerabilities [34].
- **Bias and Fairness Monitoring:** Continuously audit AI decisions for compliance with fair lending laws (e.g., CFPB requirements) [35].
- **Model Governance:** Track model versions, data lineages, and deployment approvals [26].

D. Pillar 4: Adaptive Regulatory Compliance

Agentic AI can maintain continuous compliance:

- **Regulatory Intelligence:** Monitor and interpret evolving regulations across jurisdictions [36].
- **Gap Analysis:** Compare current controls against regulatory requirements.
- **Evidence Generation:** Automatically compile audit trails and compliance documentation [37].

VII. CASE STUDIES AND APPLICATIONS

A. Financial Services: Real-Time Fraud Detection and Reporting

Major financial institutions are deploying Agentic AI for fraud management [2]. These systems autonomously:

- Analyze transaction patterns across multiple channels.
- Flag suspicious activities using adaptive ML models.
- Generate suspicious activity reports (SARs) for regulatory submission.
- Self-improve through reinforcement learning from investigator feedback.

B. Defense Contracting: Automated DFARS Compliance

A prototype Agentic AI system for defense contractors addresses DFARS requirements by:

- Continuously scanning for NIST SP 800-171 control deviations.
- Automating cloud security assessments per DFARS 252.239-7010.
- Generating cyber incident reports with required data elements.
- Providing real-time compliance dashboards for contracting officers.

C. Cross-Sector: Third-Party Risk Intelligence

Agentic AI platforms aggregate threat intelligence from multiple sources to assess vendor risks [38]. These systems autonomously:

- Monitor dark web for vendor data leaks.
- Analyze vendor security certifications and audit reports.
- Simulate supply chain attack scenarios.
- Recommend risk mitigation strategies.

VIII. RISKS AND MITIGATIONS

While Agentic AI offers significant benefits, it introduces new risks that must be managed:

A. Risk 1: Autonomous Agent Security

Threat: Agents could be hijacked or manipulated to perform malicious actions [25].

Mitigation: Implement robust authentication, behavior monitoring, and kill switches for agents [39].

B. Risk 2: Regulatory Uncertainty

Threat: Evolving regulations may not account for autonomous AI decision-making [40].

Mitigation: Design flexible, explainable systems that can adapt to regulatory changes [26].

C. Risk 3: Over-Reliance on Automation

Threat: Human oversight may diminish, creating accountability gaps [41].

Mitigation: Maintain human-in-the-loop for critical decisions and establish clear accountability frameworks [42].

D. Risk 4: AI-Enhanced Adversarial Attacks

Threat: Attackers use AI to develop more sophisticated exploits.

Mitigation: Deploy adversarial training, anomaly detection, and regular security assessments [34].

IX. IMPLEMENTATION GUIDELINES

Based on industry best practices and regulatory guidance, we recommend:

A. Governance Structure

- Establish an AI Governance Committee with cross-functional representation.
- Develop AI risk management frameworks aligned with NIST AI RMF and regulatory expectations.
- Implement model risk management (MRM) processes for Agentic AI systems.

B. Technical Architecture

- Design modular systems with clear boundaries between agents.
- Implement comprehensive logging and audit trails for all autonomous actions.
- Use encryption and secure multi-party computation for sensitive data.

C. Compliance Integration

- Map Agentic AI controls to regulatory requirements (DFARS, NYDFS, etc.).
- Conduct regular compliance testing and independent audits.
- Maintain detailed documentation for regulatory examinations.

X. RELATED WORK: RESEARCH ON AGENTIC AI IN CYBERSECURITY AND DEFENSE

Recent research by Satyadhar Joshi and colleagues provides critical insights into the application of Agentic AI in cybersecurity and national defense contexts. This section synthesizes key findings from Joshi's comprehensive body of work, which spans technical architectures, policy frameworks, and implementation strategies for autonomous AI systems in high-stakes environments.

A. Foundational Frameworks for Agentic AI in Cybersecurity

Joshi's research establishes several foundational frameworks for understanding and deploying Agentic AI in cybersecurity contexts. In [43], he explores the synergistic relationship between Agentic AI and High-Performance Computing (HPC), proposing a layered reference architecture that integrates cloud infrastructure, edge computing, and autonomous AI agents for scalable cybersecurity solutions. This work demonstrates how autonomous AI systems can enhance threat detection, incident response, and risk management through real-time processing capabilities enabled by HPC infrastructure.

Key technical contributions from this research include:

- **Multi-agent system architectures** for coordinated threat response
- **Algorithmic foundations** for autonomous decision-making in cyber defense
- **Scalability frameworks** leveraging cloud and edge computing resources

B. Agentic AI in Financial Cybersecurity

Joshi's comprehensive review of generative AI in financial cybersecurity [15] provides detailed analysis of architectures, algorithms, and regulatory challenges specific to the financial sector. This work examines:

- The dual role of AI as both defensive tool and offensive threat vector
- Integration of AI-driven frameworks with existing financial systems
- Regulatory compliance considerations for AI deployment

The research proposes a structured approach to AI governance in financial institutions, emphasizing the need for explainable AI (XAI) components and robust testing methodologies to ensure reliability and accountability.

C. National Security and Defense Applications

Joshi's work on national security applications [11], [12] presents comprehensive policy recommendations and technical frameworks for integrating Agentic AI into U.S. military systems. Key findings include:

TABLE II: Key Applications of Agentic AI in Defense (Adapted from Joshi's Research)

Application Domain	Agentic AI Capabilities	Technical Requirements
Command and Control	Autonomous mission planning, adaptive decision-making	Digital twin environments, secure communication protocols
Intelligence, Surveillance, Reconnaissance (ISR)	Real-time data analysis, pattern recognition	High-bandwidth processing, sensor fusion algorithms
Cyber Defense	Proactive threat hunting, automated response	Adversarial robustness testing, XAI components
Swarm Warfare	Coordinated autonomous operations	Multi-agent coordination algorithms, fail-safe mechanisms

This research identifies critical challenges in system integration, adversarial robustness, and human-machine teaming, proposing layered security frameworks and standardized interoperability protocols for defense applications.

D. Human-AI Collaboration Frameworks

Recent work by Adebayo [16] complements Joshi's research by emphasizing the importance of Human-in-the-Loop (HITL) paradigms integrated with Explainable AI (XAI). This research demonstrates that:

- Reliability in autonomous cybersecurity cannot be achieved by AI alone
- Structured human oversight is essential at critical decision points
- XAI-driven uncertainty quantification can trigger appropriate human intervention

The HITL-XAI framework proposed in this research has shown promising results, reducing false positive-mediated disruptions by 34% and improving adaptability to novel attack patterns by 50% in field studies.

E. Workforce Transformation for the Agentic AI Era

Joshi's research on military workforce reskilling [19] addresses the human capital challenges associated with Agentic AI adoption. Key insights include:

- Current readiness gaps: Only 10-15% of military personnel feel adequately trained for agentic AI integration
- Investment needs: \$600-900 million required for next-generation AI capabilities
- Educational transformation: Multi-tiered architecture with progressive competency levels

The proposed framework includes:

- Foundational AI literacy for all personnel
- Operational competence for mid-career leaders
- Strategic AI leadership development
- Phased implementation over 24-36 months

F. Technical Architectures and Implementation Strategies

Joshi's research provides detailed technical architectures for Agentic AI deployment, including:

1) *Multi-Agent System Architectures*: These architectures feature specialized agents for different cybersecurity functions, coordinated through central orchestrators. Key components include:

- **Specialized agents**: Threat detection, incident response, compliance monitoring
- **Coordination mechanisms**: Message passing, shared knowledge bases
- **Orchestration layers**: Task allocation, conflict resolution

2) *Secure DevOps Pipelines*: For military AI deployment, Joshi proposes tailored DevOps pipelines that incorporate:

- Security testing at every stage of development
- Adversarial robustness validation
- Compliance verification with defense standards
- Continuous monitoring and improvement

3) *Digital Twin Environments*: These virtual environments enable:

- Safe training and validation of AI agents
- Simulation of cyber attacks and defenses
- Testing of multi-agent coordination strategies
- Evaluation of system robustness under various conditions

G. Research Gaps and Future Directions

Building on past work, several research gaps and future directions emerge:

- 1) **Ethical Governance Mechanisms**: Further development of frameworks ensuring AI accountability and compliance with international norms
- 2) **Cross-Domain Integration**: Research on integrating Agentic AI across air, land, sea, space, and cyber domains
- 3) **Resilience Against Adversarial AI**: Enhanced defenses against AI-powered attacks on autonomous systems
- 4) **Standardization Efforts**: Development of interoperability standards for multi-vendor AI systems in defense
- 5) **Long-Term Adaptation**: Mechanisms for continuous learning and adaptation of AI systems over extended deployment periods

H. Synthesis and Alignment with Current Research

Our past body of work aligns closely with the challenges identified in Section V and provides technical foundations for the framework proposed in Section VI. Key alignments include:

- Emphasis on human-AI collaboration for reliable autonomous systems
- Focus on explainable AI for regulatory compliance and oversight
- Integration of advanced computing infrastructure for real-time threat response
- Comprehensive approach addressing technical, policy, and workforce dimensions

This research provides both theoretical foundations and practical implementation guidelines for deploying Agentic AI systems in environments requiring high reliability, security, and regulatory compliance.

XI. CONCLUSION AND FUTURE DIRECTIONS

A. Summary of Contributions

This paper has examined the transformative potential of Agentic Generative AI for addressing the intertwined challenges of cybersecurity and regulatory compliance in high-stakes environments. Through comprehensive analysis of recent regulatory developments and systematic literature review, we have demonstrated that autonomous AI systems—when implemented within appropriate governance frameworks—can fundamentally reshape how organizations defend against cyber threats while maintaining regulatory adherence.

Our proposed five-layer architecture addresses critical operational gaps identified in current approaches. By integrating specialized AI agents for threat detection, incident response,

compliance monitoring, and risk assessment with human-in-the-loop oversight mechanisms, the framework enables organizations to meet demanding regulatory requirements such as DFARS 72-hour incident reporting windows while maintaining the transparency and accountability demanded by financial regulators. The incorporation of explainable AI components throughout the system ensures that autonomous decisions remain auditable and aligned with ethical standards.

The quantitative evidence synthesized from recent deployments is compelling: 34% reduction in false positive disruptions, 50% improvement in novel threat adaptation, and order-of-magnitude acceleration in compliance reporting processes. These improvements translate directly to enhanced security posture and reduced operational costs—transforming compliance from a burden into a strategic capability.

B. Practical Implications

For financial institutions and defense contractors, this research provides actionable guidance for Agentic AI adoption:

Strategic Planning: Organizations should begin with comprehensive readiness assessments evaluating current infrastructure, workforce capabilities, and regulatory exposure. Our analysis suggests optimal resource allocation of 30-40% to technology infrastructure, 20-25% to workforce development, and 15-20% to governance frameworks.

Phased Implementation: Rather than wholesale transformation, successful deployments follow graduated rollout strategies. Initial implementation should target low-risk, high-volume tasks such as log analysis and compliance documentation, progressively expanding to critical decision-making as organizational confidence and technical maturity increase.

Regulatory Engagement: Proactive dialogue with regulators is essential. Organizations should document AI system architectures, decision-making processes, and oversight mechanisms in advance of deployment, seeking regulatory guidance on novel applications and contributing to the development of sector-specific AI governance standards.

Workforce Transformation: The skills gap affecting 85-90% of current cybersecurity personnel cannot be ignored. Investment in multi-tiered education programs—from foundational AI literacy for all personnel to strategic AI leadership for executives—is as critical as technology infrastructure itself.

C. Limitations and Challenges

Despite its promise, Agentic AI deployment faces significant obstacles that warrant honest acknowledgment:

Technical Complexity: The proposed architecture requires sophisticated integration across heterogeneous systems, high-performance computing infrastructure, and continuous model retraining. Many organizations, particularly smaller defense contractors and regional financial institutions, may lack the technical sophistication or resources for full implementation.

Adversarial Risks: While Agentic AI enhances defensive capabilities, it simultaneously introduces new attack surfaces. Adversaries may exploit model vulnerabilities through data poisoning, adversarial examples, or manipulation of agent

decision-making processes. Our analysis identifies adversarial robustness as a critical research gap requiring immediate attention.

Regulatory Uncertainty: The regulatory landscape for AI in cybersecurity remains fragmented and evolving. Divergent state-level requirements (e.g., New York vs. California), international variations (DFSA vs. OSFI guidance), and pending federal AI legislation create compliance complexity that may offset some operational benefits.

Ethical Considerations: Autonomous systems making high-stakes security decisions raise fundamental questions about accountability, transparency, and human agency. While HITL frameworks partially address these concerns, the appropriate boundary between autonomous and human decision-making remains contested.

D. Future Research Directions

Building on the foundations established in this work, we identify five critical research priorities:

1. Standardization and Interoperability: The cybersecurity community urgently needs standardized protocols for Agentic AI systems to enable cross-vendor integration and facilitate regulatory oversight. Research should focus on developing common interfaces, data formats, and evaluation metrics that work across diverse deployment environments.

2. Adversarial Robustness and Red Teaming: Systematic investigation of AI system vulnerabilities through structured red team exercises is essential. Research should develop comprehensive threat models for Agentic AI systems, establish testing methodologies that simulate sophisticated adversaries, and create defensive techniques that maintain system effectiveness under attack.

3. Cross-Domain Integration: Most current deployments focus on single domains (network security, cloud security, etc.). Future research should explore architectures that enable coordinated defense across air, land, sea, space, and cyber domains—particularly relevant for defense applications where threats increasingly span multiple environments.

4. Regulatory Sandboxes and Pilot Programs: Collaboration between industry and regulators to establish controlled environments for testing novel AI applications would accelerate innovation while managing risk. Research should evaluate sandbox effectiveness, develop best practices for structured experimentation, and create frameworks for transitioning successful pilots to full production.

5. Long-Term Adaptation and Evolution: Current AI systems require frequent retraining as threat landscapes evolve. Research is needed on architectures that enable continuous learning and adaptation over extended deployment periods while maintaining stability, preventing catastrophic forgetting, and ensuring alignment with organizational objectives.

6. Cross-Sector Threat Intelligence Sharing: Effective cyber defense requires information sharing across organizational boundaries. Research should address technical protocols, legal frameworks, and incentive structures that enable real-time threat intelligence sharing among financial institutions,

defense contractors, and government agencies while protecting competitive information and privacy.

E. Policy Recommendations

Based on our analysis, we offer recommendations for policymakers and regulators:

Harmonize Regulatory Requirements: Federal agencies should work toward consistent AI governance standards across sectors, reducing compliance complexity and enabling organizations to leverage common frameworks and technologies.

Invest in Public Goods: Government should fund development of open-source AI security tools, shared threat intelligence platforms, and educational resources that benefit the entire ecosystem, particularly supporting smaller organizations that cannot afford proprietary solutions.

Establish AI Assurance Programs: Regulators should develop certification programs for AI cybersecurity systems, providing organizations with clear compliance pathways and giving regulators confidence in system reliability and accountability.

Support Workforce Development: Federal investment in AI education programs, from K-12 STEM initiatives to professional reskilling programs, is essential to building the workforce capable of implementing and overseeing Agentic AI systems.

F. Concluding Remarks

The convergence of escalating cyber threats, expanding regulatory requirements, and rapid AI advancement creates an inflection point for regulated sectors. Organizations that strategically deploy Agentic AI within robust governance frameworks will gain significant competitive advantages: enhanced security posture, streamlined compliance operations, and improved resilience against evolving threats. Those that delay face mounting costs, increasing regulatory risk, and potential strategic obsolescence.

However, realizing this potential requires moving beyond technology deployment to fundamental organizational transformation. Success demands investment not only in AI systems but in people, processes, and partnerships. It requires collaboration among technology providers, regulated entities, and government agencies to develop standards, share intelligence, and establish governance mechanisms that balance innovation with accountability.

The framework presented in this paper provides a foundation for this journey—a risk-based approach that leverages autonomous AI capabilities while maintaining the human oversight and transparency essential for high-stakes environments. As regulatory frameworks continue to evolve and AI capabilities advance, this foundation will require continuous refinement and adaptation.

The question is no longer whether Agentic AI will transform cybersecurity and compliance operations, but how quickly organizations can navigate this transformation while managing associated risks. The coming decade will separate organizations that proactively integrate AI into their security and

compliance programs from those overwhelmed by the complexity of modern regulatory and threat landscapes. This paper aims to accelerate that transition, providing both theoretical foundations and practical guidance for responsible Agentic AI deployment in environments where failure is not an option.

The future of cybersecurity in regulated sectors is autonomous, adaptive, and accountable. The frameworks and insights presented here chart a path toward that future—one that enhances both security and compliance while preserving the human judgment essential for navigating our most consequential decisions.

DECLARATION

The views expressed are those of the author and do not represent any affiliated institutions. This work constitutes independent research reviewing existing literature and proposing implementation frameworks based on cited research. This is a working paper and has been edited by AI Tools.

REFERENCES

- [1] 2024 Volume 16 Applying Risk Appetite and Risk Tolerance in the Age of AI. ISACA. [Online]. Available: <https://www.isaca.org/resources/news-and-trends/newsletters/atisaca/2024/volume-16/applying-risk-appetite-and-risk-tolerance-in-the-age-of-ai>
- [2] AI in Finance — Fraud Protection, Trading & Risk Management. [Online]. Available: <https://bolster.ai/blog/the-evolution-of-finance-ais-growing-influence>
- [3] AI in Financial Modeling: Applications, Benefits, and Development. Corporate Finance Institute. [Online]. Available: <https://corporatefinanceinstitute.com/resources/data-science/ai-financial-modeling/>
- [4] 3 Hidden Risks of AI for Banks and Insurance Companies. Lumenova AI. [Online]. Available: <https://www.lumenova.ai/blog/risks-of-ai-banks-insurance-companies/>
- [5] Cybersecurity Risks for Financial Services Firms: Proactive Strategies to Stay Ahead. JD Supra. [Online]. Available: <https://www.jdsupra.com/legalnews/cybersecurity-risks-for-financial-3900428/>
- [6] New York Department of Financial Services Guidance on AI-Related Cybersecurity Risks — Saul Ewing LLP. [Online]. Available: <https://www.saul.com/insights/alert/new-york-department-financial-services-guidance-ai-related-cybersecurity-risks>
- [7] M. E. A. Finance. New DFSA report explores regulatory insights into cybersecurity, Artificial Intelligence, and quantum risks. mea-finance.com. [Online]. Available: <https://mea-finance.com/new-dfsa-report-explores-regulatory-insights-into-cybersecurity-artificial-intelligence-and-quantum-risks/>
- [8] OSFI and FCAC Report on AI Use and Risk at Federally Regulated Financial Institutions — Blakes. [Online]. Available: <https://www.blakes.com/insights/osfi-and-fcac-report-on-ai-use-and-risk-at-federally-regulated-financial-institutions>
- [9] Bird, B.-J. Gong, T. Luo, M. Dong, D. Xia, and S. Liu. China Cybersecurity and Data Protection: Monthly Update - June 2025 Issue. Lexology. [Online]. Available: <https://www.lexology.com/library/detail.aspx?g=b0f74de8-3df0-4260-8420-80ded8ee0756>
- [10] R. Abbas. Strengthening Financial Services with Third-Party Risk Mitigation Strategies - Cybersecurity Magazine. [Online]. Available: <https://cybersecurity-magazine.com/strengthening-financial-services-with-third-party-risk-mitigation-strategies/>
- [11] S. Joshi. Agentic Generative AI and National Security: Policy Recommendations for US Military Competitiveness. [Online]. Available: <https://www.preprints.org/manuscript/202510.0401/v1>
- [12] —. Agentic Generative AI and National Security: Policy Recommendations for US Military Competitiveness. [Online]. Available: <https://papers.ssrn.com/abstract=5529680>
- [13] A. Abdullahi. NVIDIA: Agentic AI Is Reshaping Cybersecurity Defense. eSecurity Planet. [Online]. Available: <https://www.esecurityplanet.com/cybersecurity/nvidia-agentic-ai-cybersecurity/>
- [14] The 2025 Reality of Agentic AI in Cybersecurity. [Online]. Available: <https://www.cybersecuritytribe.com/articles/the-2025-reality-of-agentic-ai-in-cybersecurity>
- [15] S. Joshi, "Gen Ai in Financial Cybersecurity: A Comprehensive Review of Architectures, Algorithms, and Regulatory Challenges," vol. 4, no. 3, pp. 73–88. [Online]. Available: <https://philarchive.org/rec/JOSGAI-3>
- [16] H. Adebayo. Human-in-the-Loop Explainable AI for Reliable Autonomous Cybersecurity Infrastructure. [Online]. Available: <https://www.preprints.org/manuscript/202601.2031>
- [17] "[No title found]."
- [18] Agentic AI in BFS: The Future of Smarter Financial Systems. Tredence. [Online]. Available: [https://www.tredence.com/blog/\\$slug](https://www.tredence.com/blog/$slug)
- [19] S. Joshi, "Reskilling the U.S. Military Workforce for the Agentic AI Era: A Framework for Educational Transformation." [Online]. Available: <https://eric.ed.gov/?id=ED677111>
- [20] N. Abigail. Mastercard's \$10bn cyber push: Adam Jones on AI, identity and the future of digital trust - Arabian Business: Latest News on the Middle East, Real Estate, Finance, and More. [Online]. Available: <https://www.arabianbusiness.com/resources/mastercards-10bn-cyber-push-adam-jones-on-ai-identity-and-the-future-of-digital-trust>
- [21] (2) "The Future is Now": How AI is Changing Cybersecurity Risk in Banks & Financial Services — LinkedIn. [Online]. Available: <https://www.linkedin.com/pulse/future-now-how-ai-changing-cybersecurity-ql3ie/>
- [22] How Agentic AI-Driven Cybersecurity Prevents Financial Cyber Threats. [Online]. Available: <https://www.akira.ai/blog/ai-agents-for-cybersecurity-in-finance>
- [23] Transforming Cybersecurity with Agentic AI: A Double-Edged Sword? — LinkedIn. [Online]. Available: <https://www.linkedin.com/pulse/transforming-cybersecurity-agentic-ai-double-edged-sword-nir-kshetri-qwpte/>
- [24] Key Use Cases of AI in Risk Management. Safe Security. [Online]. Available: <https://safe.security/resources/blog/key-use-cases-ai-risk-management/>
- [25] Z. Heck. Cybersecurity in the Era of Generative and Agentic AI: Six Observations. Taft Privacy & Data Security Insights. [Online]. Available: <https://www.privacyanddatasecurityinsight.com/2025/06/cybersecurity-in-the-era-of-generative-and-agentic-ai-six-observations/>
- [26] AI for the CRO: Transforming AI governance, compliance and security. [Online]. Available: <https://rsmus.com/insights/services/digital-transformation/ai-for-the-cro.html>
- [27] A. Mellen. Cybersecurity's Latest Buzzword Has Arrived: What Agentic AI Is And Isn't. Forrester. [Online]. Available: <https://www.forrester.com/blogs/cybersecuritys-latest-buzzword-has-arrived-what-agentic-ai-is-and-isnt/>
- [28] Taking a business-critical approach to supplier nth-party IT risk management — McKinsey. [Online]. Available: <https://www.mckinsey.com/capabilities/risk-and-resilience/our-insights/taking-a-business-critical-approach-to-supplier-nth-party-it-risk-management>
- [29] Your AI Technology Partner Could Be a Security Trojan Horse. The Financial Brand. [Online]. Available: <https://thefinancialbrand.com/news/artificial-intelligence-banking/your-ai-technology-partner-could-be-a-security-trojan-horse-190540>
- [30] How and Quantum Changing the Cyber Security Landscape in Finance. Check Point Software. [Online]. Available: <https://www.checkpoint.com/cyber-hub/cyber-security/what-is-ai-security/how-ai-is-changing-the-cyber-security-landscape-in-finance/>
- [31] C. Crisanto, C. B. Leuterio, J. Prenio, and J. Yong, *Regulating AI in the Financial Sector: Recent developments and Main Challenges*, ser. FSI Insights on Policy Implementation. Bank for International Settlements, Financial Stability Institute, no. no 63.
- [32] Artificial intelligence in cybersecurity - Article. [Online]. Available: <https://www.sailpoint.com/identity-library/artificial-intelligence-cybersecurity>
- [33] Risks and Strategies for AI Cybersecurity Risks: Key Takeaways from NY DFS Letter — NETBankAudit. [Online]. Available: <https://www.netbankaudit.com/resources/ai-cybersecurity-risks-dfs-letter-october-2024>
- [34] Are new gen AI tools putting your business at additional risk? — IBM. [Online]. Available: <https://www.ibm.com/think/news/are-new-genai-tools-putting-your-business-at-risk>
- [35] CFPB Comments on AI Offer Insights for Consumer Finance Industry — Insights — Skadden, Arps, Slate, Meagher & Flom LLP. [Online]. Available: <https://www.skadden.com/insights/publications/2024/08/cfpb-comments-on-artificial-intelligence>

-
- [36] "Regulatory approaches to Artificial Intelligence in finance." [Online]. Available: https://www.oecd.org/en/publications/regulatory-approaches-to-artificial-intelligence-in-finance_f1498c02-en.html
- [37] Top 5 Cybersecurity Automation Tools Transforming Risk Management. [Online]. Available: <https://www.cybersaint.io/blog/top-5-cybersecurity-automation-tools>
- [38] Top five risks for financial institutions in 2025. WTW. [Online]. Available: <https://www.wtwco.com/en-us/insights/2025/03/top-five-risks-for-financial-institutions-in-2025>
- [39] What Is Agentic AI in Cybersecurity. Safe Security. [Online]. Available: <https://safe.security/resources/blog/what-is-agentic-ai-in-cybersecurity/>
- [40] M. S. J. P. A. G. A. J. J. Bridges, Micaela McMurrough. California Frontier AI Working Group Issues Final Report on Frontier Model Regulation. Inside Global Tech. [Online]. Available: <https://www.insideglobaltech.com/2025/06/30/california-frontier-ai-working-group-issues-final-report-on-frontier-model-regulation/>
- [41] The Compliance Risks of Using Generative AI in a Financial Planning Practice — Financial Planning Association. [Online]. Available: <https://www.financialplanningassociation.org/learning/publications/journal/MAY25-compliance-risks-using-generative-ai-financial-planning-practice-OPEN>
- [42] How Banking Leaders Can Enhance Risk and Compliance With AI. The Financial Brand. [Online]. Available: <https://thefinancialbrand.com/news/artificial-intelligence-banking/how-banking-leaders-can-enhance-risk-and-compliance-with-ai-183094>
- [43] S. Joshi. Advancing Cybersecurity Through Synergies of Agentic AI and High-Performance Computing. [Online]. Available: <https://papers.ssrn.com/abstract=5341131>

ABOUT THE AUTHOR

Satyadhar Joshi is a quantitative analyst with expertise in financial risk, data science, machine learning, and artificial intelligence. He currently serves as an Assistant Vice President at Bank of America. His independent research focuses on AI-driven risk assessment, financial modeling, and big data analytics, with a particular emphasis on developing modeling tools and innovative approaches to advance AI capabilities in support of the U.S. national interest.