

Cybersecure by Design: Managing Adversarial Risk and Resilience in AI-Enabled Critical Infrastructure

You Wu^{1*}

¹Tsinghua University, CHINA

Abstract

The use of artificial intelligence in critical infrastructure creates a new class of cyber-physical and institutional risks. AI systems depend on extensive data pipelines, complex models, software supply chains, cloud services, interfaces, and continuous updates. These features expand the attack surface beyond traditional information technology and expose models to data poisoning, evasion, model extraction, inference attacks, manipulation of sensors, and unsafe automated responses. This review examines cybersecurity as a central component of AI governance rather than a separate technical speciality. It draws on the thematic synthesis of 95 high-quality studies and frameworks covering public administration, digital services, smart cities, and infrastructure-related applications. The evidence indicates that many governance frameworks acknowledge privacy and ethics but provide limited operational guidance for threat modelling, adversarial testing, incident response, resilience, and audit evidence. The review presents a lifecycle approach to AI security that spans design, data preparation, model development, validation, deployment, monitoring, change management, and retirement. It also explains how risk tiering can connect system impact to proportionate controls and how assurance records can support accountability. Particular attention is given to smart grids, transport, healthcare, buildings, and structural monitoring, where security failures can move quickly from digital disruption to physical or social harm. The article concludes that cybersecure AI requires integrated governance, continuous testing, and recovery planning from the beginning of system design.

Keywords: AI Cybersecurity; Adversarial Machine Learning; Critical Infrastructure; Resilience; Incident Response; Threat Modelling; Sustainable Systems

Introduction

Artificial intelligence can improve the operation of critical infrastructure by detecting faults, forecasting demand, optimising maintenance, and supporting rapid decisions. The same capabilities can create serious vulnerabilities when models are manipulated, data are compromised, or automated outputs are trusted without adequate safeguards. In ordinary information systems, a cyber incident may expose data or interrupt a service. In AI-enabled infrastructure, an incident can also change how the system interprets reality. A corrupted model may fail to detect equipment damage, misallocate energy, misclassify a medical case, or recommend an unsafe operational response.

Traditional cybersecurity remains essential, but it is not sufficient. AI systems introduce dependencies that are unusual in conventional software. Their behaviour is shaped by training data, feature engineering, model architecture, optimisation choices, and ongoing data streams. Small changes to these components can produce large and difficult-to-predict effects. Attackers may target

the model indirectly through poisoned data or manipulated sensors. They may probe outputs to reconstruct sensitive information. They may create inputs designed to trigger a specific error. They may exploit weak governance around model updates, third-party components, or cloud interfaces.

These risks make cybersecurity a governance issue. Technical teams can test models, but institutions must decide which tests are mandatory, who reviews results, what level of residual risk is acceptable, when deployment should stop, and how incidents will be communicated. Security must therefore connect with accountability, privacy, procurement, legal compliance, and operational continuity. This review develops that integrated perspective.

Review Approach and Evidence Base

The article is a thematic narrative review derived from the supplied systematic review of AI governance risk and cybersecurity frameworks. The underlying synthesis screened a large body of literature and retained 95 studies and frameworks that met defined quality criteria. The period from 2020 to mid-2025 captured a sharp increase in research attention, especially after the expansion of generative AI and high-profile regulatory developments.

The source analysis found that cybersecurity and accountability formed a strong technical axis in the statistical mapping of governance frameworks, while ethics and privacy formed a separate human and social axis. Most frameworks clustered around one side or the other. Truly integrated approaches were relatively uncommon. This separation is problematic for critical infrastructure because safety and legitimacy depend on both. A technically robust system may still be unfair or unlawful, while an ethically framed system may remain vulnerable to attack.

The present review organises the evidence around the AI lifecycle and focuses on operational controls. It also uses the source risk-tiering concept to explain how security requirements can be scaled according to impact.

The AI-Specific Threat Landscape

Data poisoning occurs when an attacker introduces misleading or malicious information into training or updating processes. The objective may be broad degradation, a hidden backdoor, or targeted failure under particular conditions. Poisoning is especially dangerous in systems that learn from external or continuously collected data because the attack may resemble normal variation. Data provenance, source authentication, anomaly detection, and controlled update procedures are therefore core governance controls.

Evasion attacks manipulate inputs at the point of use. An adversarial example may look normal to a person while causing the model to make an incorrect prediction. In physical infrastructure, the manipulation may involve sensor readings, images, network traffic, or operational commands. Robustness testing should include realistic threat scenarios rather than only laboratory benchmarks.

Model inversion and inference attacks use access to a model or its outputs to reveal information about training data. These attacks create a direct link between cybersecurity and privacy. Restricting access, limiting output detail, monitoring unusual queries, and applying privacy-preserving methods can reduce exposure. Model extraction attacks seek to reproduce the behaviour or

intellectual property of a system through repeated queries. Rate limits, authentication, watermarking, and behavioural monitoring may be appropriate depending on risk.

Supply-chain threats arise when organisations rely on external datasets, open-source libraries, pretrained models, cloud platforms, or vendors. A vulnerability can enter before the organisation receives the product. Governance should therefore require component inventories, provenance records, vulnerability management, contractual disclosure, and independent validation. The most dangerous assumption is that a trusted vendor automatically provides a trusted model.

Operational failures may also occur without a malicious attacker. Data drift, model drift, software conflicts, configuration mistakes, sensor degradation, and poorly controlled updates can produce outcomes similar to an attack. Security and resilience programmes should address both deliberate and accidental causes because the consequences for essential services may be identical.

Threat or failure mode	Typical consequence	Priority controls	Assurance evidence
Data poisoning	Corrupted learning, hidden backdoors, targeted misclassification	Trusted data sources, provenance checks, anomaly screening, controlled retraining	Dataset lineage, screening report, update approval, rollback record
Adversarial evasion	Incorrect prediction from crafted or manipulated input	Threat-based testing, robustness evaluation, input validation, fallback mode	Red-team report, test cases, residual-risk decision
Model inversion or inference	Exposure of sensitive training information	Output limitation, access control, privacy testing, query monitoring	Privacy test results, access logs, alert history
Model extraction	Loss of intellectual property or replication of sensitive behaviour	Authentication, rate limiting, monitoring, watermarking where suitable	Usage logs, threshold configuration, investigation records
Supply-chain compromise	Vulnerable or manipulated model, library, dataset, or service	Component inventory, vendor due diligence, integrity checks, contractual controls	Software/model bill of materials, supplier assessment, integrity report
Drift or unsafe update	Declining accuracy, unexpected behaviour, service disruption	Continuous monitoring, change control, staged release, rollback capability	Drift dashboard, change record, deployment test, rollback test

Security Across the AI Lifecycle

Security begins during problem definition. The organisation should identify the essential service being supported, the consequences of error, likely adversaries, and conditions under which automation must stop. Threat modelling at this stage helps teams avoid designs that are inherently difficult to secure. For example, a system that accepts unrestricted external inputs may require a different architecture from one operating in a controlled environment.

Data governance is the next defence layer. Training and operational data should have documented sources, integrity checks, access rules, retention periods, and quality controls. Sensitive datasets require privacy assessment as well as security testing. Changes to data sources should trigger review because a new source can alter both model behaviour and attack exposure.

During development, secure coding and conventional application security should be combined with model-specific testing. Teams should test robustness to plausible perturbations, inspect sensitivity to features, assess failure under incomplete or abnormal data, and evaluate whether outputs reveal too much information. High-impact systems need independent validation rather than relying solely on the development team.

Deployment introduces interfaces, credentials, APIs, devices, networks, and human workflows. The model should operate with the least privilege necessary. Logs must capture important actions without creating new privacy risks. Fallback procedures should be available when the model is unavailable or uncertain. Where the output can affect safety or rights, human reviewers need clear authority to pause or override the system.

Monitoring must continue after release. Performance metrics alone are insufficient. Security monitoring should include unusual query patterns, data-distribution changes, access anomalies, repeated borderline inputs, sensor inconsistency, and unexpected changes in model confidence. Monitoring thresholds should be reviewed because attackers can adapt to static rules. The organisation also needs a schedule for retesting after significant changes.

Retirement deserves equal attention. Models, datasets, credentials, and interfaces should be decommissioned securely. Retained records must support later investigation and accountability. An abandoned endpoint or forgotten model can remain a security risk long after the service has changed.

Risk Tiering and Proportionate Security

Not every AI system requires the same security burden. Risk tiering allows organisations to match controls to impact. A low-risk internal tool using non-sensitive data may require standard access control, logging, and periodic review. A medium-risk system may need formal threat modelling, testing before major updates, and defined incident procedures. A high-risk system affecting essential services or significant rights should undergo independent adversarial testing, continuous monitoring, strict change control, and human oversight. A critical system may require regulator notification, external assurance, redundant operation, and tested fail-safe modes.

Risk classification should consider more than the model itself. Relevant factors include data sensitivity, scale, reversibility of harm, dependence on automation, exposure to external inputs, connectivity to physical systems, affected populations, and availability of alternatives. A simple

model controlling a critical process may deserve more protection than a technically complex model used for a minor administrative task.

Tiering should also determine evidence requirements. High-risk systems need detailed threat assessments, test results, decision records, and incident exercises. The purpose is not paperwork for its own sake. Evidence allows leaders, auditors, regulators, and operators to understand whether controls were applied and whether accepted risks were reasonable.

Incident Response and Operational Resilience

AI incidents can be difficult to recognise because harmful behaviour may appear as ordinary model error. Response plans should therefore define indicators of compromise and routes for technical, operational, legal, privacy, and public communication. Teams need authority to isolate data sources, disable model functions, switch to manual operation, roll back updates, and preserve evidence.

Resilience means more than preventing attacks. It includes the ability to continue essential services, recover safely, and learn from failure. Critical systems should avoid single points of dependence. Alternative procedures, backup models, manual controls, or degraded-service modes may be necessary. Recovery testing should confirm that backups and rollback mechanisms actually work.

Post-incident review should examine governance as well as technical causes. An attack may succeed because testing was weak, but also because ownership was unclear, procurement records were incomplete, warnings were ignored, or staff lacked authority to pause deployment. Lessons should update policies, threat models, training, and risk classification. A mature programme treats incidents as evidence for continuous improvement rather than isolated technical events.

Sector Applications

Smart grids use AI to forecast demand, detect faults, and balance distributed energy resources. Manipulated forecasts or control signals could affect service continuity and equipment safety. Security controls should combine data validation, network segmentation, real-time anomaly detection, operator oversight, and recovery planning.

Transport systems may use AI for traffic management, maintenance prediction, or autonomous functions. Their exposure to sensors and external environments increases the importance of adversarial testing and fail-safe behaviour. Explanations are also needed so that operators can understand why the system recommends a disruptive action.

Healthcare infrastructure uses sensitive data and may depend on AI for triage, imaging, scheduling, or resource allocation. Security failure can expose personal information or affect care. Controls should integrate privacy protection, clinical validation, access management, and rapid human escalation.

Smart buildings and structural health monitoring rely on networks of sensors and automated analysis. Attackers may alter measurements, suppress warnings, or manipulate energy controls. Data provenance, device security, cross-sensor validation, and engineering review are critical. These examples show that cybersecurity controls must reflect the physical and institutional setting of each system.

Governance Integration and Future Research

Cybersecurity teams cannot manage AI risk alone. Governance committees should connect security findings with privacy, ethics, safety, procurement, and operational decisions. The risk owner must understand both the technical likelihood of failure and the consequences for people and services. Contracts should require vendors to disclose significant changes and cooperate with investigations. Internal audit should test whether security controls are operating, not merely whether policies exist.

The literature still lacks long-term evidence about which AI-specific security controls are most effective in real deployments. Future research should compare adversarial testing methods, develop realistic benchmarks for infrastructure contexts, and measure the cost and benefit of continuous monitoring. There is also a need for shared incident taxonomies so that organisations can learn from events without exposing sensitive details.

Automated security testing and monitoring will expand, but these tools should not create false confidence. Their coverage and limitations need independent evaluation. Human judgement remains necessary to interpret ambiguous signals, understand operational context, and make accountable decisions under uncertainty.

Conclusion

AI-enabled infrastructure changes the meaning of cybersecurity. Protecting networks and software remains essential, but organisations must also defend data pipelines, models, update processes, interfaces, and automated decisions. The reviewed governance landscape shows that these controls are still less mature than broad commitments to privacy and ethics.

A cybersecure-by-design approach begins with threat modelling and continues through data governance, adversarial testing, secure deployment, monitoring, incident response, and retirement. Risk tiering makes this approach proportionate, while the evidence layer makes it auditable. Resilience planning ensures that essential services can continue when prevention fails.

The central lesson is practical: AI security cannot be added after a model has been approved. It must shape the architecture, data, procurement, workflow, and accountability structure from the start. When cybersecurity is integrated with governance, institutions are better able to use AI without sacrificing safety, trust, or the long-term resilience of critical infrastructure.

References

- [1] Ehrlinger, L., & Wöß, W. (2022). A survey of data quality measurement and monitoring tools. *Frontiers in Big Data*, *5*, 850611.
- [2] Aluri, Y. S. (2023). Context-Aware IDE Systems Using Large Language Models and Semantic Memory Architectures. *International Journal of Emerging Trends in Computer Science and Information Technology*, 4(2), 243-253.
- [3] Redman, T. C. (2020). The rise of AI in data quality management. *Harvard Business Review*.
- [4] Batini, C., & Scannapieco, M. (2016). *Data and Information Quality: Dimensions, Principles and Techniques*. Springer.
- [5] Loshin, D. (2011). *The Practitioner's Guide to Data Quality Improvement*. Morgan Kaufmann.

- [6] Wang, R. Y., & Strong, D. M. (1996). Beyond accuracy: What data quality means to data consumers. *Journal of Management Information Systems*, *12*(4), 5–33.
- [7] Pipino, L. L., Lee, Y. W., & Wang, R. Y. (2002). Data quality assessment. *Communications of the ACM*, *45*(4), 211–218.
- [8] Strong, D. M., Lee, Y. W., & Wang, R. Y. (1997). Data quality in context. *Communications of the ACM*, *40*(5), 103–110.
- [9] Haug, A., & Zachariassen, F. (2016). Data quality in CRM systems. *Industrial Management & Data Systems*, *116*(9), 1896–1914.
- [10] Mallemapati, A. (2022). Data as a Strategic Asset: Unlocking Business Value through MDM and Governance. *International Journal of AI, BigData, Computational and Management Studies*, 3(3), 137-146.
- [11] Miller, T. (2023). A hierarchical model for data quality terminology. *Data & Knowledge Engineering*, *143*, 102087.
- [12] Wilkinson, M. D., Dumontier, M., Aalbersberg, I. J., Appleton, G., Axton, M., Baak, A., ... & Mons, B. (2016). The FAIR Guiding Principles for scientific data management and stewardship. *Scientific Data*, *3*, 160018.
- [13] Kumar, M. S., & Yuvaraj, N. (2022). Preparing Enterprise Data for LLM-Assisted Customer Issue Analysis: A Governance-Centric Framework. *International Journal of Artificial Intelligence, Data Science, and Machine Learning*, 3(3), 181-192.
- [14] Stall, S., Yarmey, L., Cutcher-Gershenfeld, J., Hanson, B., Lehnert, K., Nosek, B., ... & Wyborn, L. (2019). Make scientific data FAIR. *Nature*, *570*(7759), 27–29.
- [15] Yuvaraj, N., & Kumar, M. S. (2021). From Governed Data to Customer Health Signals: Integrating Telemetry with Enterprise Data Quality Controls. *International Journal of Emerging Trends in Computer Science and Information Technology*, 2(4), 115-125.
- [16] NIST. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. National Institute of Standards and Technology.
- [17] Mallemapati, A., & Jaladi, D. S. (2022). An API-Driven Master Data Management Framework for Distributed Enterprise Application Integration. *International Journal of Emerging Research in Engineering and Technology*, 3(2), 151-160.
- [18] Herzog, T. N., Scheuren, F. J., & Winkler, W. E. (2007). *Data Quality and Record Linkage Techniques*. Springer.
- [19] Lee, Y. W., Pipino, L. L., Funk, J. D., & Wang, R. Y. (2006). *Journey to Data Quality*. MIT Press.
- [20] Aluri, Y. S. (2021). Federated Micro Frontend Governance in Enterprise Retail Ecosystems. *International Journal of Artificial Intelligence, Data Science, and Machine Learning*, 2(2), 114-125.

- [21] Batini, C., Cappiello, C., Francalanci, C., & Maurino, A. (2009). Methodologies for data quality assessment and improvement. *ACM Computing Surveys*, *41*(3), 1–52.
- [22] Aluri, Y. S. (2022). Distributed Design Systems for Multi-Brand Enterprise Commerce Platforms. *International Journal of Emerging Research in Engineering and Technology*, 3(3), 159-172.
- [23] Redman, T. C. (2008). *Data Driven: Profiting from Your Most Important Business Asset*. Harvard Business Press.
- [24] Jaladi, D. S., & Mallempati, A. (2022). A Scalable Full-Stack. NET Enterprise Application Framework for Data Integration with Master Data Management Systems. *International Journal of AI, BigData, Computational and Management Studies*, 3(2), 169-177.
- [25] Boeckhout, M., Zielhuis, G. A., & Bredenoord, A. L. (2018). The FAIR guiding principles for data stewardship: Fair enough? *European Journal of Human Genetics*, *26*(7), 931–936.
- [26] Kumar, M. S. (2022). An AI-Driven Framework for Data Governance, Quality Management, and Metadata Integration in Enterprise Systems. *International Journal of Artificial Intelligence, Data Science, and Machine Learning*, 3(2), 165-175.