



Open Threat Database Integration and Volunteer-Driven AI Safety Initiatives as a Model for Hybrid Threat Intelligence Architecture in Commercial Cybersecurity Platforms

Anton Kulyk

CEO and Founder, Cyber Trust Innovations LLC, Montenegro.

Abstract

The article examines a hybrid threat intelligence architecture for commercial cybersecurity platforms that integrates open threat databases, industry sources, and volunteer-driven AI safety initiatives. The purpose of the study is to substantiate a three-tier model that combines governmental vulnerability databases, community-based standards, and sources addressing artificial intelligence incidents. The relevance of the research is determined by the expansion of the contemporary threat surface, where CVE-based vulnerabilities, web application risks, supply chain weaknesses, and AI-related incidents require a unified analytical environment. The novelty of the article lies in its interpretation of volunteer-driven AI safety initiatives as an independent component of commercial threat intelligence infrastructure. Using VULNWatch as a case study, the article demonstrates that the combined use of NVD, OWASP, the AI Incident Database, and fourteen integrated tools expands risk coverage, reduces dependence on single-source feeds, and supports alignment with NIST CSF, NIST AI RMF, ISO/IEC 27001, and ISO/IEC 42001. The article will be useful for researchers, developers of cybersecurity platforms, AI governance specialists, and information security auditors.

Keywords: Threat Intelligence, Open Data Integration, AI Safety, National Vulnerability Database, OWASP, AI Incident Database, Hybrid Architecture, AI Governance, Vulnerability Assessment.

INTRODUCTION

The set of risks that commercial cybersecurity platforms must mitigate has widened over the last several years. Classical web application vulnerabilities such as injection flaws, broken access control and security misconfigurations remain a constant operational concern, and the volume of disclosed Common Vulnerabilities and Exposures entries continues to grow year on year (European Union Agency for Cybersecurity, 2024). A second risk surface has emerged in parallel, centred on artificial intelligence systems themselves. It includes model misuse, prompt injection, training data poisoning and downstream harms produced by deployed models. AI risks of this kind sit outside the scope of traditional vulnerability databases, yet they shape the security posture of any organisation that embeds AI into customer-facing or internal workflows (McGregor, 2021).

Threat intelligence at the commercial platform level has

not kept pace with this expansion. The dominant pattern in industry is to assemble intelligence around a proprietary feed bought from a single vendor, or around one or two open repositories such as NVD. Proprietary feeds carry a high price and opaque curation, while single-source open ingestion narrows coverage. AI risks sit outside this picture, since the most active repositories of AI incidents and AI safety guidance are maintained by academic and non-profit projects of the type represented by the AI Incident Database, or by volunteer-driven initiatives oriented toward AI governance frameworks. The institutional separation between governmental vulnerability databases, community standards bodies and AI safety communities leaves a vendor that relies on a narrow source set unable to deliver complete threat intelligence.

The literature contains substantial work on each source class taken individually. NVD has been studied as a corpus for vulnerability analytics, OWASP Top 10 and similar

Citation: Anton Kulyk, "Open Threat Database Integration and Volunteer-Driven AI Safety Initiatives as a Model for Hybrid Threat Intelligence Architecture in Commercial Cybersecurity Platforms", Universal Library of Innovative Research and Studies, 2026; 3(3): 26-31. DOI: <https://doi.org/10.70315/uloap.ulirs.2026.0303005>.

community catalogues have a long practitioner literature, and the AI Incident Database has begun to attract empirical analysis as a resource for AI safety research (Paeth et al., 2025). What remains underdeveloped is a description of how one commercial platform can integrate all three source classes at the architectural level, and how participation in volunteer-driven AI safety work can extend that architecture beyond passive data ingestion.

An earlier paper by the present author described the internal AI pipeline used in VULNWatch for risk prioritisation and explained how findings are ranked once collected (Kulyk, 2026). The present paper addresses a different layer of the system. It looks at the external ecosystem that decides what data the platform sees in the first place, and at how that ecosystem changes when a commercial vendor participates in volunteer-driven AI safety work.

Three contributions follow. The paper proposes a three-tier hybrid threat intelligence model that separates authoritative governmental sources, community- and industry-standard sources, and emerging AI-safety sources, and outlines the complementary strengths and weaknesses of each tier. The model is illustrated through a case study of the VULNWatch platform, which covers the integration of fourteen open and proprietary tools and a thirty-day audience profile of the platform. The paper then discusses what the model implies for alignment with established cybersecurity and AI governance frameworks, and argues that volunteer-driven initiatives form a layer of commercial threat intelligence infrastructure that has so far gone unrecognised.

MATERIALS AND METHODS

Research Design

The paper combines a conceptual analysis with an embedded single case study. The conceptual component proposes a three-tier model of threat intelligence sources, and the case study illustrates it in a commercial setting. The case study method follows the principles set out by Yin for analysing present-day phenomena in their real-world context, where the boundaries between the phenomenon and its context are loose (Yin, 2018). The platform under study works as an illustrative case. It is not a representative case; its purpose is to demonstrate the feasibility of the proposed model. The analysis makes no claim to statistical generalisation.

Unit of Analysis: the VULNWatch Platform

The unit of analysis is VULNWatch at vulnwatch.tech, a commercial software-as-a-service platform for automated web vulnerability assessment operated by Cyber Trust Innovations LLC. Its internal scanning and AI prioritisation pipeline appears in the author's earlier work (Kulyk, 2026) and is not repeated here. The platform is examined at the level of its external integrations. That level covers the tools the platform uses, the public sources it ingests, and the volunteer-driven initiatives in which it participates.

Sources Considered

Four classes of external sources are included in the analysis. The first class is the U.S. National Vulnerability Database, maintained by the National Institute of Standards and Technology, which publishes structured CVE records and Common Vulnerability Scoring System values (National Institute of Standards and Technology, n.d.). The second class is the OWASP Top 10, a community catalogue of the most critical web application security risks that the Open Web Application Security Project revises on a recurring schedule (OWASP Foundation, 2021). The third class is the AI Incident Database at incidentdatabase.ai, an international repository that collects real-world incidents involving artificial intelligence systems (McGregor, 2021). The fourth class is Project Cerebellum at projectcerebellum.com, a volunteer-driven organisation that develops a framework for the safe use and regulation of AI, with which Cyber Trust Innovations LLC collaborates on a pro bono basis. The platform's technological layer includes 14 integrated tools, listed in Section 3.2.

Analytical Framework

The architecture is examined against four established governance frameworks. These are the NIST Cybersecurity Framework (Pascoe et al., 2024), the NIST AI Risk Management Framework 1.0 (Tabassi, 2023), ISO/IEC 27001 for information security management, and ISO/IEC 42001 for AI management systems. The four frameworks were chosen because together they span classical cybersecurity governance and the AI-specific governance surface that the proposed model is designed to address.

Empirical Data on the Case

Two forms of empirical data are available on the platform. The first form is the publicly visible set of integrated tools and ingested public sources, as documented on the platform. The second form is internal platform telemetry over a thirty-day observation window, reporting unique visitor counts and the geographic distribution of web traffic. These data describe the audience the platform reaches. They do not describe the security findings it produces. Section 3.4 uses them to characterise the user base. The discussion in Section 4 returns to the observation period and the single vendor nature of these data as limitations of the analysis.

RESULTS

The Three-Tier Hybrid Threat Intelligence Model

Figure 1 presents the proposed three-tier model. Each tier corresponds to a distinct institutional source of threat intelligence, with characteristic strengths and limitations. The integration layer shown at the bottom of the figure is the commercial platform component that normalises records across tiers, correlates findings, applies AI-driven prioritisation and produces user-facing reports.

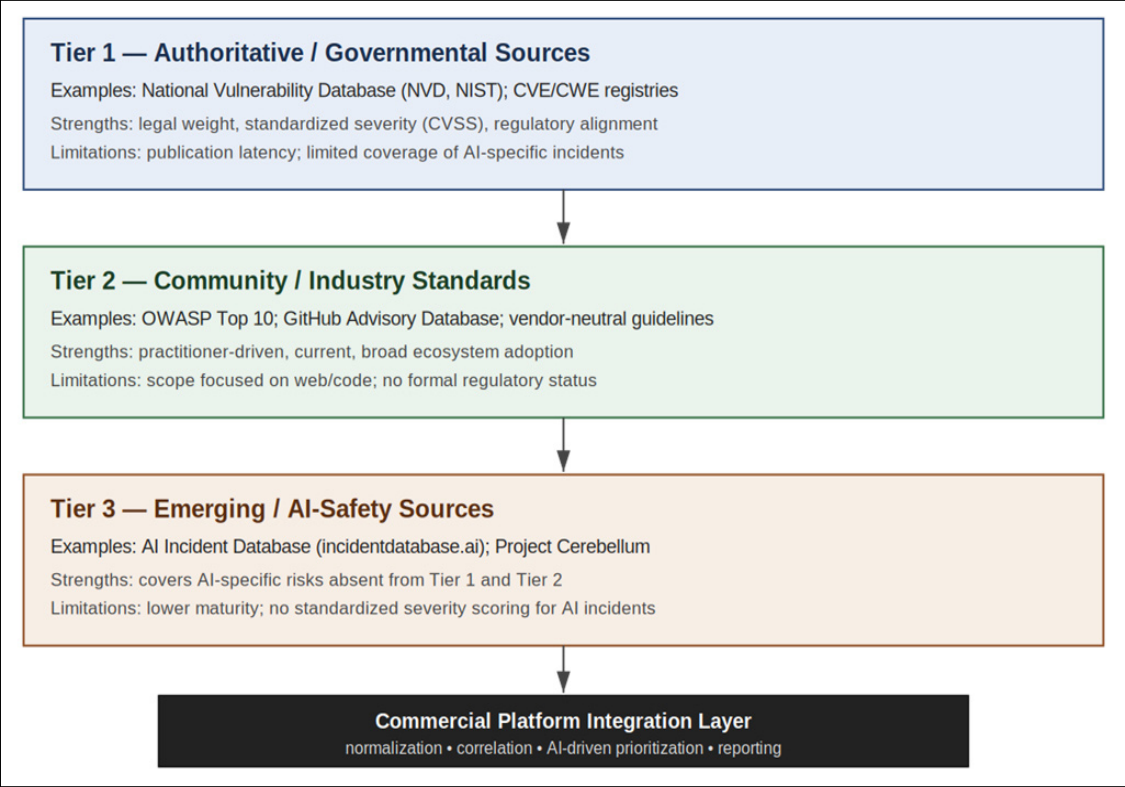


Figure 1. A three-tier hybrid threat intelligence architecture

Tier 1 covers authoritative and governmental sources, of which NVD is the leading example. Sources of this kind carry legal and regulatory weight, use standardised severity metrics such as CVSS, and align with audit and compliance requirements. The limitations of Tier 1 are well documented. Publication latency between disclosure and database entry is one such limitation. The CVE schema covers a narrow band of risks and provides little coverage of AI incidents. Tier 2 covers community- and industry-standard sources, including the OWASP Top 10 and the GitHub Advisory Database. These sources are driven by practitioners, are more current than Tier 1 for new attack classes, and have broad adoption across the development ecosystem. Their scope remains tied to web application and software dependency risks, and they hold no formal regulatory status. Tier 3 covers emerging sources focused on AI safety, including the AI Incident Database and volunteer-driven projects such as Project Cerebellum. Tier 3

encompasses a class of risks that Tiers 1 and 2 leave aside, including model misuse, deployment harm, and governance gaps. The limitations of Tier 3 stem from the lower maturity of its data structures and from the absence of a standardised severity model comparable to CVSS.

The model’s substantive argument is that no single tier provides adequate threat intelligence coverage. Value at the commercial platform level emerges from the correlation between tiers. Deeper investment in any one tier brings less. A vulnerability surfaced by a community tool such as OWASP ZAP, mapped against a Tier 1 CVE record and correlated with a Tier 3 AI incident pattern, yields richer intelligence than any of the three taken in isolation.

Technological Layer: Integrated Tools

Table 1 maps the fourteen tools integrated into VULNWatch to the threat classes they address and to the source tier with which each tool aligns.

Table 1. Integrated tools, addressed threat classes, and source tier

Tool	Primary threat class addressed	Type	Tier
OWASP ZAP	Dynamic application security testing	Open source	Tier 2
Burp Suite	Web traffic interception and testing	Industry standard	Tier 2
Metasploit	Exploit development and red team testing	Open source	Tier 2
GitHub Advisory Database	Dependency vulnerability tracking	Community feed	Tier 2
crt.sh	Certificate transparency and subdomain discovery	Public service	Tier 1
urlscan.io	Phishing and malware URL analysis	Public service	Tier 2
WPScan	WordPress core, plugin and theme vulnerabilities	Open source	Tier 2
Nmap	Network and port reconnaissance	Open source	Tier 2

sqlmap	SQL injection detection and exploitation	Open source	Tier 2
URLhaus	Malware distribution URL intelligence	Community feed	Tier 2
Web Archive	Historical exposure of deprecated assets	Public service	Tier 2
decodo proxy	Cloaking and redirection detection	Custom integration	Tier 2
Tor	Anonymous reachability and abuse testing	Open source	Tier 2
Nuclei	YAML template-based vulnerability scanning and custom detection	Open source	Tier 2
VULNWatch AI Module	AI driven detection, prioritisation, remediation advice	Proprietary	Tier 3

Two observations follow from Table 1. The majority of tools sit at Tier 2, which reflects the practitioner-driven heritage of web vulnerability assessment. A single component, the proprietary VULNWatch AI Module, operates at Tier 3. The pattern mirrors the wider industry, where the AI safety tier of commercial tooling remains thin and often comes from a vendor's in-house development. Mature off-the-shelf components for this layer stay scarce.

Volunteer-Driven Initiatives and the Case of Project Cerebellum

Beyond passive consumption of public data, the platform extends its Tier 3 coverage through volunteer-driven collaboration. Cyber Trust Innovations LLC takes part on a pro bono basis in Project Cerebellum, an organisation that develops a framework for the safe use and regulation of artificial intelligence. The collaboration differs in character from the data feed relationships described in Sections 3.1 and 3.2. The platform contributes engineering capacity to the development of the framework and, in return, gains exposure to evolving normative guidance ahead of its public stabilisation.

This pattern can be called a private volunteer partnership. It

stands apart from the established public-private partnership model in cybersecurity. Under the public-private model, government agencies and commercial actors cooperate around regulated infrastructure. Under the private volunteer model, a commercial actor and a volunteer-driven initiative cooperate around an emerging normative ground that no regulation has yet codified. AI safety offers the clearest current example. Formal regulation is still under construction, with the European Union AI Act as the most developed instrument so far. Private volunteer partnerships fill the gap between practitioner consensus and statutory rule. From the platform's perspective, the partnership functions as an early warning channel for risks that will reach Tier 1 sources several years later.

Audience Profile of the Platform

Figure 2 shows the geographic distribution of platform traffic over a thirty-day observation window. The total number of unique visitors over the window reached 12,220, with a daily maximum of 1,080 and a daily minimum of 344. The leading countries by request volume were the United States (70,427), France (32,147), Hong Kong (13,949), Canada (11,427), and Israel (6,347).

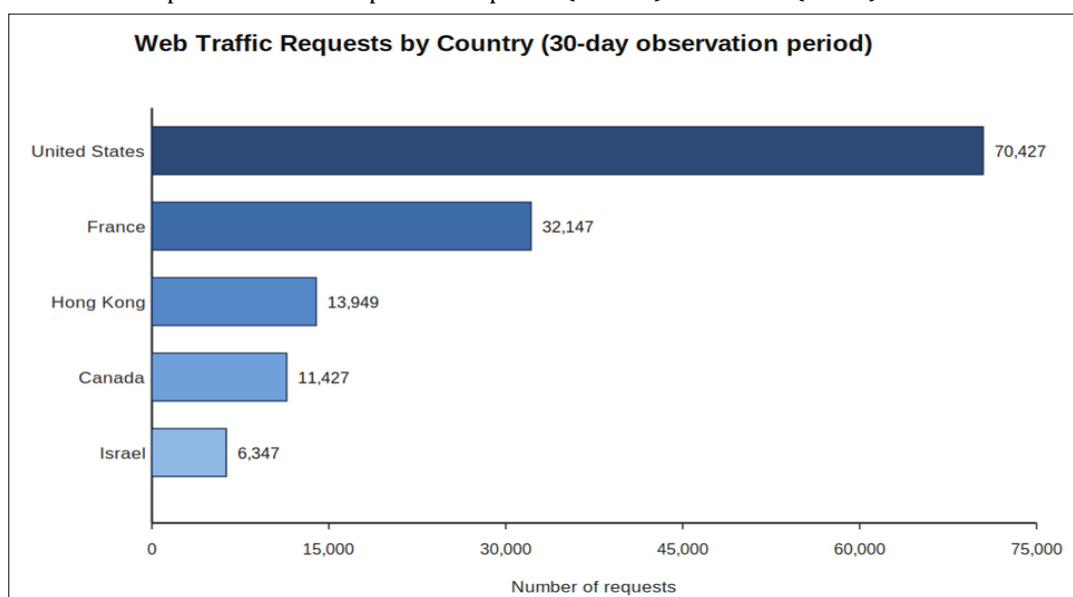


Figure 2. Geographic distribution of platform requests over the thirty-day observation window

Table 2 presents the audience profile in a single view. The distribution looks typical of a business-to-business audience composed of cybersecurity professionals and atypical of a consumer base. The leading countries cluster in jurisdictions with developed cybersecurity industries, and the visitor-to-request ratio fits repeated, work-driven use. Incidental browsing does not match the figures.

Table 2. Platform audience profile over the thirty-day observation window

Metric	Value
Total unique visitors over 30 days	12,220
Maximum unique visitors per day	1,080
Minimum unique visitors per day	344
Leading country by request volume	The United States, with 70,427 requests
Second by request volume	France, with 32,147 requests
Third by request volume	Hong Kong, with 13,949 requests
Fourth by request volume	Canada, with 11,427 requests
Fifth by request volume	Israel, with 6,347 requests

Two interpretive points follow from the profile. The audience composition provides soft evidence of the platform's positioning. A tool whose users gather in established cybersecurity markets is used in professional contexts, where the marginal cost of adding another intelligence tier is justified by professional demand. The international spread of the audience strengthens the case for cross-jurisdictional sources within the model. A platform whose users are in the United States, France, Canada, and Israel cannot rely on a single national source set, so the multi-tier model also serves as a multi-jurisdictional model.

Alignment with Governance Frameworks

Table 3 maps the three tiers of the proposed model to four established governance frameworks. The mapping is indicative. It is not exhaustive. The table identifies the function within each framework to which each tier brings the strongest contribution.

Table 3. Mapping of model tiers to selected governance frameworks

Framework	Tier 1 contribution	Tier 2 contribution	Tier 3 contribution
NIST CSF 2.0	Identify, Detect	Identify, Protect	Govern for AI
NIST AI RMF 1.0	Limited	Limited	Govern, Map, Measure, Manage
ISO/IEC 27001	Control objectives A.5–A.8	Operational controls	Limited
ISO/IEC 42001	Limited	Limited	AI management system core

The pattern in Table 3 reinforces the case for a multi-tier architecture. Classical cybersecurity frameworks such as NIST CSF and ISO/IEC 27001 are well served by Tiers 1 and 2, yet engage with the AI surface at the margins. AI-specific frameworks such as NIST AI RMF 1.0 and ISO/IEC 42001 exhibit the opposite pattern, as they require Tier 3 inputs that Tiers 1 and 2 cannot provide. A platform that integrates all three tiers can align with both regimes in parallel, which matters more for vendors serving customers facing both information security and AI governance audit requirements.

DISCUSSION

What the Hybrid Model Adds Over a Monolithic Approach

The clearest practical benefit of the proposed model is coverage. A monolithic threat intelligence source, whether a single proprietary feed or a single open repository, omits a sizeable portion of the contemporary threat surface by construction. NVD, on its own, omits most AI incidents. An OWASP-centred approach omits supply-chain CVE details. An AI-centric approach omits the operational backbone of classical web security. The three-tier typology offers nothing radical in itself. The argument here is that a commercial platform should organise around all three tiers as primary inputs of comparable architectural weight. A second benefit of the model is reduced single vendor dependency. A platform

built on public Tier 1 and Tier 2 sources, supported by Tier 3 community work, faces less exposure to pricing changes, deprecation, or altered access conditions in any single feed, which matters more for smaller and mid-sized vendors.

Volunteer-Driven Initiatives as an Under-Recognised Resource

The AI safety field, at the time of writing, develops in non-profit, academic, and volunteer-driven settings, while commercial tooling lags behind. Vendors that ignore this layer miss out on early visibility into emerging AI risk taxonomies. The role now played by AI safety initiatives parallels the role OWASP played for web security fifteen years ago. A community body that began as an advocacy group became the de facto reference point for an entire branch of the industry. If the trajectory repeats, vendors that participate now in volunteer-driven AI safety work will be earlier on the maturity curve than those that wait for the field to consolidate.

Legal and Regulatory Implications

The multi-tier model offers favourable audit properties from a legal standpoint. Tier 1 sources supply regulator-recognised reference data. Tier 2 sources supply community-recognised standards. Tier 3 participation supplies documented engagement with the AI safety community. For a commercial vendor preparing for audits under ISO/IEC 27001 and ISO/

IEC 42001, or for compliance review under the European Union AI Act, the combination forms a coherent evidential trail. The fact that Tiers 1 and 2 remain publicly verifiable lets customers and regulators confirm the data lineage on which the platform operates without depending on the vendor's word, and this serves as a non-trivial differentiator against opaque proprietary feeds.

Limitations

Four limitations of the present work deserve attention. The case study draws on a single platform, serves as an illustration of the model, and does not constitute proof. The model itself needs validation across additional commercial settings before any claim of broader applicability can stand. The audience profile data comes from a thirty-day window at one vendor and cannot be generalised beyond that window. The proposed integration of tiers assumes the continued availability of open source software, and the deprecation of any of them, with NVD as the most acute case, given recent debate around its funding cycles, would shift the model. Severity scoring across tiers remains heterogeneous, since CVSS sits at Tier 1, qualitative practitioner judgement sits at Tier 2, and ad hoc taxonomies sit at Tier 3. Reconciliation between these scoring approaches is itself an open research problem.

CONCLUSION

Commercial cybersecurity platforms operate against a threat surface that no longer fits inside the boundaries of any single intelligence source. Classical CVE-based vulnerability management, community-driven standards work, and emerging AI safety reporting each cover part of that surface, and none of them on its own suffices. The paper has proposed a three-tier hybrid threat intelligence model that integrates authoritative governmental sources, community and industry-standard sources, and emerging AI-safety sources, with volunteer-driven initiatives added as a fourth qualitative input. The model has been illustrated through the case of the VULNWatch platform, including its 14 integrated tools, its participation in the Project Cerebellum initiative, and a 30-day audience profile consistent with a professional user base.

Three practical implications follow for other commercial vendors. The first concerns source diversity, which belongs in the design at the architectural level, since a platform built around a single feed cannot move to multi-tier intelligence without rebuilding its ingestion layer. A second implication concerns volunteer-driven AI safety work, which has become the main source of structured signals on AI risk at the time of writing, and ignoring it leaves a measurable coverage gap.

A third implication concerns alignment, since established cybersecurity and AI governance frameworks already accommodate the model, and alignment with NIST CSF, NIST AI RMF 1.0, ISO/IEC 27001, and ISO/IEC 42001 follows from the tiers themselves. On a broader view, the open public databases and the volunteer-driven initiatives examined here function as critical infrastructure for the cybersecurity and AI safety industries, and the case for commercial participation in their upkeep through funding, engineering contribution or pro bono framework development stands stronger than the prevailing pattern of passive consumption would suggest.

REFERENCES

1. European Union Agency for Cybersecurity. (2024). *ENISA Threat Landscape 2024*. <https://www.enisa.europa.eu/publications/enisa-threat-landscape-2024>
2. Kulyk, A. (2026). Artificial Intelligence-Driven Risk Prioritization in Automated Web Vulnerability Assessment. *International Journal of Science and Research (IJSR)*, 15(4), 167–173. <https://doi.org/10.21275/SR26331122718>
3. McGregor, S. (2021). Preventing Repeated Real World AI Failures by Cataloging Incidents: The AI Incident Database. In *Proceedings of the AAAI Conference on Artificial Intelligence* (pp. 15458–15463). <https://doi.org/10.1609/aaai.v35i17.17817>
4. National Institute of Standards and Technology (n.d.). National Vulnerability Database (NVD). *NVD*. <https://nvd.nist.gov/>
5. OWASP Foundation. (2021). *OWASP Top 10:2021*. <https://owasp.org/Top10/2021/>
6. Paeth, K., Atherton, D., Pittaras, N., Frase, H., & McGregor, S. (2025). Lessons for Editors of AI Incidents from the AI Incident Database. In *Proceedings of the AAAI Conference on Artificial Intelligence* (pp. 28946–28953). <https://doi.org/10.1609/aaai.v39i28.35163>
7. Pascoe, C., Quinn, S., & Scarfone, K. (2024). *The NIST Cybersecurity Framework (CSF) 2.0* (Report No. NIST CSWP 29). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.CSWP.29>
8. Tabassi, E. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)* (Report No. NIST AI 100-1). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.100-1>
9. Yin, R. K. (2018). *Case Study Research and Applications: Design and Methods* (6th ed.). SAGE Publications.