

Applying Miles's Seven-Gap Taxonomy to AI-Enabled Cyber Defense: A Framework for Doctoral Research Proposal Development

Mukaila Ishola Adediran^{1*}, Gilbert I.O. Aimufua², Zakari Doko Mohammed³

¹PhD Candidate, Department of Cybersecurity, Nasarawa State University, Keffi (NSUK)

²Director, Center for Cyberspace Studies, Nasarawa State University, Keffi (NSUK)

³Corps Information Technology Office, Federal Road Safety National Headquarters, Abuja

*Corresponding Author: mikkyscho26@gmail.com

ABSTRACT: Artificial intelligence now underpins core cyber-defense capabilities-intrusion detection, malware classification, threat intelligence, and automated response, yet the field still lacks a reliable way to pinpoint where knowledge is missing, disputed, or untested. This paper develops a conceptual and methodological framework for identifying, validating, and prioritizing research gaps in AI-enabled cyber defense by adapting Miles's (2017) seven-component taxonomy: the Evidence, Knowledge, Practical-Knowledge, Methodological, Empirical, Theoretical, and Population Gaps to this domain. Beyond defining the seven components, the paper specifies a synthesis-based population procedure, a primary study validation protocol, data quality and analytical rigor controls, a prioritization scoring method with sensitivity analysis, a design-science evaluation against Nickerson et al.'s (2013) criteria, and an instrument validation analysis covering content, construct, criterion, and reliability validity. The framework is applied to nine primary studies across four AI-cyber-defense subdomains and compared against four existing gap-identification models. The paper closes with templates for converting an identified gap into a citable gap statement, alongside a discussion of limitations, bias, and responsible use.

KEYWORDS: Conceptual Framework, Methodological Framework, Artificial Intelligence, Cyber Defense, Cybersecurity, Research Gaps, Taxonomy, Framework Validation.

I. INTRODUCTION

The escalating scale, complexity, and velocity of cyber threats have rendered traditional, signature-reliant security measures increasingly obsolete. Artificial intelligence from classical machine learning to modern large language models (LLM) floods the entire cybersecurity lifecycle from threat identification to post incident recovery. Yet the sheer volume of publications does not equate to research maturity nor does it give doctoral researchers a reliable way to distinguish solved from unresolved problems. Echoing Miles (2017) research gaps are not self-evident and require a structured method to surface. This paper introduces a conceptual and methodological framework tailored to AI-driven cyber defense for identifying, validating, and ranking research gaps. Its conceptual foundation evolves Miles's (2017) seven-point taxonomy, integrating the five-gap model of Robinson et. al. (2011) with the six-gap model of Müller-Bloch & Kranz (2015) into seven domain-focused components. The accompanying methodology specifies how these components are derived from existing literature, validated against primary studies, prioritized, and

updated as the field evolves. This paper makes four contributions: (1) a seven component conceptual framework with domain-specific definitions, the first such adaptation for AI-enabled cyber defense; (2) a concrete operational methodology spanning specification, validation, quality assurance, analytical rigor, interaction analysis, prioritization, and design-science evaluation; (3) a validation exercise in which nine primary studies across four subdomains test six of the seven components with the resulting prioritization stress-tested via three alternative weighting schemes; and (4) practical templates that convert an identified gap into a research proposal statement.

II. RELATED WORKS

A. Conceptual Foundations of the Framework

The framework's conceptual core is Miles's (2017) seven research-gap categories: Evidence, Knowledge, Practical-Knowledge, Methodological, Empirical, Theoretical, and Population Gap. Six were adapted from Müller-Bloch and Kranz (2015), which built on Jacobs's (2011) taxonomy of research problems; the seventh, the Population Gap, was drawn from Robinson et al.'s (2011) systematic-review framework.

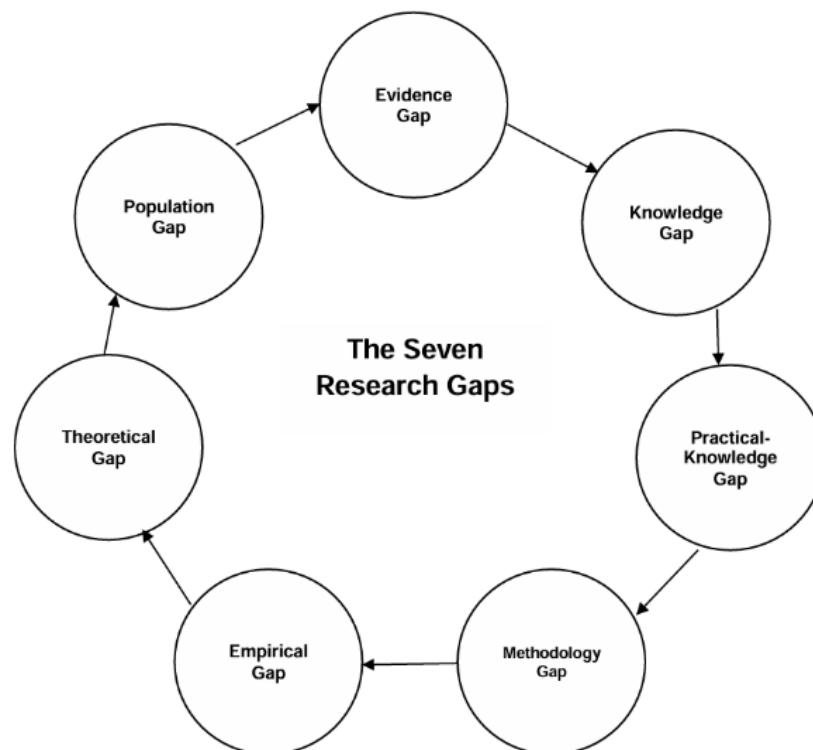


Figure 1. Conceptual model of the seven framework components (adapted from Miles, 2017).

Each component captures a distinct manifestation of incompleteness in the AI-cyber-defense literature: the Evidence Gap (contradictory accuracy and false-positive/false-negative findings across studies; Zhang et al., 2025), the Knowledge Gap (underdeveloped research on organizational trust and governance; Haryanto et al., 2024; Achuthan et al., 2024), the Practical-Knowledge Gap (a documented divergence between advocated human-in-the-loop oversight and actual analyst practice; Haryanto et al., 2024), the Methodological Gap (convergence on a narrow set of techniques and the absence of standardized benchmarks; Achuthan et al., 2024; Alqahtani & Kumar, 2026), the Empirical Gap (rarely tested claims that AI outperforms rule-based defenses; Haryanto et al., 2024), the Theoretical Gap (a scarcity of unifying theory explaining why AI

architectures succeed or fail across threat contexts; Merlano, 2024; Malatji & Tolah, 2024), and the Population Gap (geographic and organizational concentration of the evidence base; Achuthan et al., 2024).

B. Theoretical Justification

Four theoretical strands warrant treating a 'gap' as something a systematic instrument can detect rather than a rhetorical device. Epistemologically, Popper's (1959) falsificationist account and Kuhn's (1962) account of paradigm-bound normal science together justify the premise that a literature can be simultaneously large in volume and immature in coverage: publication volume reflects puzzle-solving activity, not how completely a field has tested its own falsifiable claims. Sociologically, Merton's (1973) norm of organized skepticism explains why a field can fail to notice a gap without individual dishonesty, simply through publication incentives that reward novel positive results over replication; Zhang's (2026) finding that AI-security research has structurally under-incentivized defense-focused relative to attack-focused work for over a decade illustrates this mechanism directly. Classificatory, Bailey's (1994) distinction between a typology (built conceptually from prior theory) and a taxonomy (induced empirically) clarifies that Miles's (2017) seven categories function as a typology appropriate strategy when a domain's dimensions are already reasonably well understood, and one that keeps the design-science evaluation in Section 4 methodologically consistent. Practically, treating an evidence gap as though settled can harden into unjustified operational security choices, while failing to recognize a real knowledge gap can misdirect scarce doctoral research resources—a fragmentation pattern corroborated by recent large-scale automated syntheses of the AI-in-cybersecurity literature.

C. Positioning Relative to Alternative Models

Robinson et al.'s (2011) five-gap framework was built to structure systematic-review output rather than classify the full range of literature incompleteness, and does not distinguish an evidence gap from a practical-knowledge gap; the present framework retains its population-gap coverage while adding components it does not attempt. Müller-Bloch & Kranz's (2015) six-gap framework for qualitative information-systems reviews is the closer comparator, since five of the present components derive from it directly; the addition of the population-gap component is substantive, since a six-component version would have no dedicated category for the geographic and organizational concentration of the AI-cyber-defense evidence base. Naqvi et al.'s (2024) 14-category Naqvi-Gabr framework offers finer granularity but exceeds Nickerson et al.'s (2013) seven-plus-or-minus-two usability ceiling for a single-researcher dissertation-scoping exercise. Gourisetti et al.'s (2020) CyFEr framework & Auwal's (2026) gap-prioritized model assume gaps are already identified and instead score and rank them-complementary to, rather than competing with the present framework's classificatory function, and the basis for Procedure 3's prioritization scoring (Section 3.3). Rodriguez et al.'s (2025) frontier-model capability-evaluation framework addresses attacker-side capability rather than defender-side literature gaps. Finally, Alvesson & Sandberg (2011) argue that gap-spotting approaches such as this one risk reproducing a field's existing assumptions, and propose problematization as an alternative route to research questions; the present framework is suited to scoping a bounded, defensible dissertation topic within an existing research program, while problematization targets the rarer case of challenging that program's premises—the two are complementary rather than mutually exclusive. Unlike these alternatives, the present framework is additionally checked against nine primary studies (Section

3.2, Section 4), a degree of empirical validation not previously combined with a domain-specific gap taxonomy in this literature.

III. METHODOLOGY

Specifying the framework's seven components is only half the contribution; the methodology also defines five procedures for populating, validating, prioritizing, and evaluating those components, plus data-quality and analytical-rigor controls applied throughout.

A. Research Questions

The framework addresses five questions as follows:

RQ1: What conceptual components should a research-gap framework for AI-enabled cyber defense contain, and how does each manifest in the literature?

RQ2: What methodology can validate a synthesis-derived component at the level of individual primary studies rather than only aggregated reviews?

RQ3: How do the seven components interact and take priority over one another, and how can they be ranked under resource constraints?

RQ4: Does the framework itself satisfy Nickerson et al.'s (2013) design-science criteria—conciseness, robustness, comprehensiveness, extendibility, and explanatory power?

RQ5: What methodology translates a framework-identified gap into a defensible, citable gap statement for a doctoral proposal?

B. The Seven-Component Framework

The seven-component framework applied to AI-enabled cyber defense is presented in Table 1:

Table 1: *The Seven-Component Framework Applied to AI-Enabled Cyber Defense (Condensed).*

Component	General Definition (Miles, 2017)	Domain-Specific Manifestation
Evidence Gap	Individually valid findings contradict one another at a broader level of analysis.	Detection accuracy and error rates vary widely across studies, producing conflicting reliability claims (Achuthan et al., 2024; Kaur et al., 2023).
Knowledge Gap	Desired findings simply do not yet exist.	Organizational trust and governance mechanisms remain markedly under-researched relative to technical topics (Haryanto et al., 2024).
Practical-Knowledge Gap	A discrepancy exists between advocated and actual professional behavior.	Human-in-the-loop oversight is recommended, but workforce readiness and actual analyst reliance on automation are inconsistently studied (Haryanto et al., 2024; Malatji & Tolah, 2024).
Methodological Gap	Prior research relies on a narrow set of methods, warranting diversification.	The corpus concentrates on familiar algorithms (SVM, CNN, ANN); standardized benchmarks remain absent (Achuthan et al., 2024).
Empirical Gap	Propositions are conceptually plausible but empirically untested.	Few studies compare AI-enhanced and rule-based defenses under equivalent conditions (Achuthan et al., 2024; Haryanto et al., 2024).

Theoretical Gap	Research lacks an applied or unifying theoretical foundation.	The field is dominated by applied, technique-driven studies with little foundational theory (Merlano, 2024; Malatji & Tolah, 2024).
Population Gap	Certain populations remain under-represented in the evidence base.	Research concentrates in a few countries and large enterprises, leaving SMEs, resource-constrained regions, and frontline staff under-represented (Achuthan et al., 2024).

Source: Miles (2017), Robinson et al. (2011), Müller-Bloch & Kranz (2015).

C. Five-Procedure Methodology

- (a). Procedure 1 (synthesis-based specification): This populates each component's domain manifestation from existing systematic reviews and large-scale automated syntheses (principally Achuthan et al., 2024; Haryanto et al., 2024; Kaur et al., 2023; Merlano, 2024; Malatji & Tolah, 2024).
- (b). Procedure 2 (primary-study validation): This tests whether Procedure 1's components are visible at the level of individual studies rather than only as an aggregation artifact. Candidate studies are identified through targeted searches of Google Scholar, IEEE Xplore, and arXiv (cs.CR); a study qualifies only if it reports original data collection or model evaluation mapping onto at least one component in Table 1, and studies reporting a single in-dataset accuracy metric with no cross-study, cross-dataset, or practitioner-facing element are excluded.
- (c). Procedure 3 (gap-interaction and prioritization scoring): This rates each component on explanatory reach, feasibility within a dissertation's constraints, and its potential to unlock progress on the other six, adapted from Gourisetti et al. (2020).
- (d). Procedure 4 (design-science evaluation): This assesses the framework itself against Nickerson et al.'s (2013) five taxonomy-quality criteria (Section 4.3).
- (e). Procedure 5 (instrument-validation analysis): This treats the framework as a measurement instrument and evaluates its content, construct, criterion-related, and reliability validity. Content validity rests on the pedigree of the source taxonomies rather than an original expert-panel review; construct validity draws on convergent findings across independently conducted studies; criterion-related validity is established where Procedure 2 confirms a synthesis-derived component against primary-study evidence; and reliability-inter-rater agreement on component classification-remains an outstanding gap, since classifications were made by the two co-authors rather than multiple independent coders using a shared codebook.

Underlying all five procedures, data-quality controls (source verification against publisher/DOI records, currency screening, duplicate/preprint-versioning control, and extraction-accuracy spot checks) and analytical-rigor controls (deliberate disconfirming-evidence search, triangulation across independent sources, source-to-classification traceability, co-author cross-check, and sensitivity checking of the prioritization ordering) were applied to ensure the evidence base and its interpretation could withstand independent scrutiny.

D. Prioritization and Sensitivity Analysis

Each component was scored 1-5 on explanatory reach, feasibility, and unlocking potential, then combined under three weighting schemes: equal-weight, feasibility-weighted (favoring a resource-constrained doctoral researcher), and reach-weighted (favoring theoretical or field-level contribution).

Table 2: Prioritization Scores under Three Weighting Schemes

Component	Equal-Weight (Rank)	Feasibility-Weighted (Rank)	Reach-Weighted (Rank)
Evidence Gap	4.00 (2)	4.40 (1)	3.70 (3)
Knowledge Gap	3.33 (4)	3.60 (3)	3.10 (5)
Practical-Knowledge Gap	3.33 (4)	3.60 (3)	3.10 (5)
Methodological Gap	4.33 (1)	3.80 (2)	4.80 (1)
Empirical Gap	3.33 (4)	3.60 (3)	3.10 (5)
Theoretical Gap	4.00 (2)	3.20 (6)	4.70 (2)
Population Gap	3.00 (7)	2.60 (7)	3.30 (4)

The Methodological and Evidence Gaps are rank-stable, appearing in the top three under every scheme, consistent with the Methodological Gap functioning as a root cause of the Evidence Gap pattern (Section 4.1). The Theoretical and Population Gaps are highly rank-sensitive-both foundational but comparatively infeasible for a single dissertation-while the Knowledge, Practical-Knowledge, and Empirical Gaps remain consistently mid-ranked, making them comparatively low-risk choices regardless of a researcher's weighting judgment.

E. Operationalizing the Framework

Procedure 5 translation step converts a component identified via Table 1 into language suited to a doctoral proposal. For example, an Evidence Gap statement would note that individual studies report impressive AI-based detection accuracy, but that synthesis-level reviews reveal false-positive and false-negative rates that fluctuate widely by dataset and threat model, undermining vendor-reported figures (Achuthan et al., 2024; Kaur et al., 2023). Analogous templates exist for each of the seven components; researchers are cautioned to adapt the bracketed placeholders to their own sub-topic, population, or method rather than reuse the wording verbatim.

IV. RESULTS AND DISCUSSION

A. Primary-Study Validation Results

Applying Procedure 2, nine qualifying studies were located across four subdomains-network intrusion detection, malware detection, automated threat intelligence extraction, and incident-response practice-with no qualifying study found for the Population component. Schrötter et al. (2024) replicated a deep-learning intrusion detection model that scored above 90% accuracy on its original benchmark but fell to 54% on a newly recorded dataset using identical attack tools, validating the Evidence, Methodological, and Empirical Gaps at once; in a rare head-to-head comparison, a signature-based system (Snort) outperformed the neural network on every attack class. D'hooge et al. (2020) and Elangovan et al. (2026) independently corroborated the same cross-dataset generalization failure using different models and explanatory mechanisms, and Vieira et al. (2026), AlJuhaiman et al. (2026), and Bilal et al. (2026) extended the same

Evidence/Methodological pattern into malware detection, threat-intelligence extraction, and federated learning respectively-evidence that the pattern reflects a structural property of how the field evaluates AI-based defenses rather than an artifact of one subdomain.

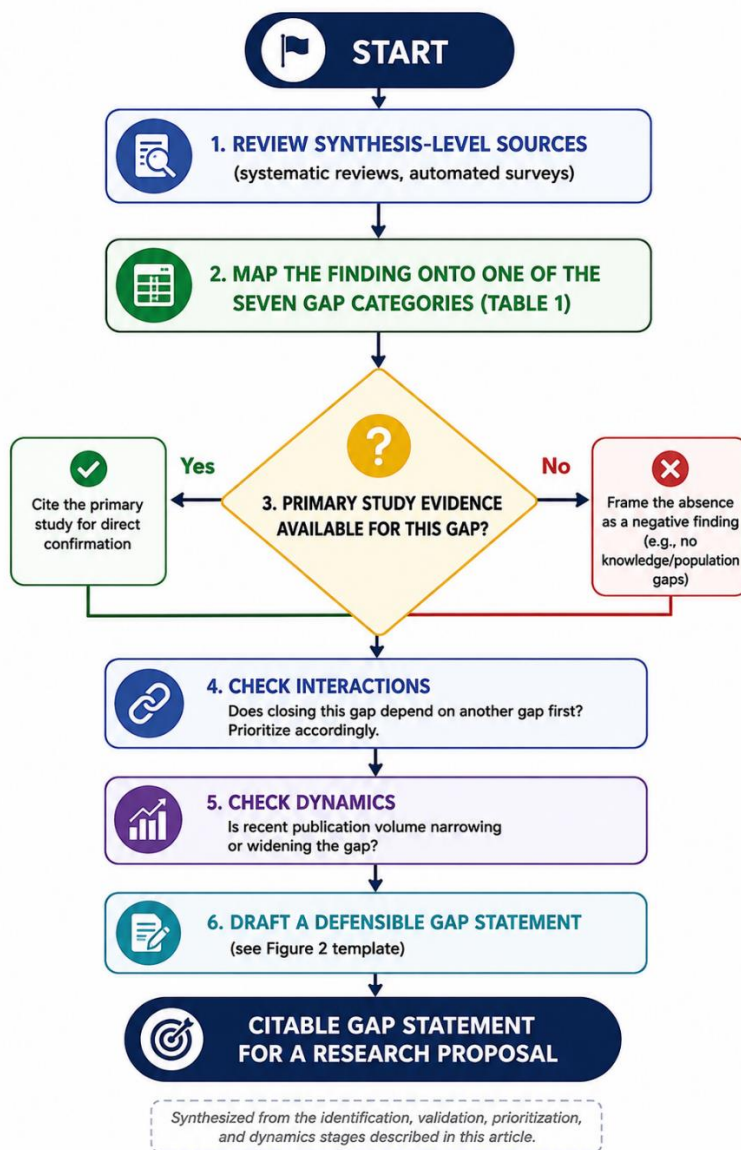


Figure 2: Procedure 2 Applied: The Primary-study Validation Workflow

Alahmadi et al. (2022), through interviews with 41 security operations center practitioners, validated the Practical-Knowledge Gap: AI-generated alerts are overwhelmingly false positives requiring manual validation, and analysts develop informal workarounds rather than following recommended human-in-the-loop procedures. Falowo & Bou Abdo (2026) supplied the first primary-study evidence for the Knowledge Gap: in a survey of 194 U.S.-based practitioners, only 20% of organizations had customized frameworks such as NIST or SANS for AI-driven threats, and only 21% maintained a separate AI threat-modeling process, despite 96% supporting industry-wide revision. The Theoretical Gap was validated negatively: among the primary studies reviewed, only two attempted any explanatory theoretical mechanism in all. No qualifying primary study

was located for the Population Gap despite a targeted, repeated search that identified three genuine near-misses (Kamruzzaman et al., 2026; Ye, 2026; Adabara et al., 2025) an absence consistent with, and arguably strengthening, the claim that this component is under-examined at the primary-data level.

B. Gap Dynamics Over Time

A gap analysis is a snapshot rather than a permanent map. A 2026 UK government-commissioned review of 9,109 AI-security publications found that publication volume grew roughly sixfold from 2021 to 2025 with agentic-AI security research growing 216% between 2024 and 2025 alone—evidence that fast-emerging technology components can narrow quickly. By contrast, Zhang's (2026) finding of a widening imbalance between attack-focused and defense-focused research, attributed to structural publication incentives rather than early-stage neglect, shows that some components persist despite the field's overall maturation. The methodological implication is that applying the framework is not a one-time exercise: a doctoral researcher should re-run an updated version of the search-and-count procedure before committing years of work to a given component.

C. Design-Science Evaluation

Assessed against Nickerson et al.'s (2013) five taxonomy-quality criteria, the framework performs well on conciseness (seven components sit within the recommended seven-plus-or-minus-two range) and robustness (six primary studies mapped cleanly onto distinct components without forcing). Comprehensiveness and extendibility are only partially tested: no qualifying study was located for the Population component, and the framework's extendibility has been demonstrated only within a single application. The framework fares well on the explanatory criterion, since each component specifies a testable mechanism rather than merely naming an under-researched topic; however, Nickerson et al. (2013) note that a taxonomy's ultimate usefulness can only be judged by independent adoption, which this paper cannot demonstrate.

D. Limitations

- (a). Component specifications in Table 1 draw primarily on existing systematic reviews and automated syntheses rather than an original primary-source coding of the entire corpus, inheriting their selection and interpretive choices.
- (b). The nine-study validation sample is purposive rather than systematic, concentrated in intrusion detection, and each non-intrusion-detection subdomain rests on a single study; phishing detection remains unrepresented despite a targeted search.
- (c). No qualifying primary study was located for the Population component, leaving the framework's comprehensiveness only partially tested.
- (d). The design-science self-assessment and component classifications were conducted by the paper's own authors, without an independent panel or a formal inter-rater reliability check.
- (e). Because AI-enabled cyber defense is fast-moving, the gap analysis and prioritization scores reflect a snapshot that may shift within a year or two.
- (f). A substantial share of the most recent supporting evidence remains in preprint form and should be treated as more provisional than peer-reviewed sources.
- (g). The validation sample is compositionally skewed toward technical, quantitative detection-accuracy studies; the Practical-Knowledge and Knowledge Gaps each rest on a single primary study, and the Population Gap rests on none.

E. Bias Considerations

Meanwhile, the seven components were fixed in advance from Miles's (2017) taxonomy, both synthesis-level sources and primary studies were read with those categories already in mind, risking a fit-to-structure rather than an emergent classification. The evidence base is geographically and linguistically narrow-concentrated in the United States, India, the United Kingdom, and China (Achuthan et al., 2024; Albahri & AlAmoodi, 2023) a narrowness illustrated concretely by Bade et al.'s (2026) account of Nigeria-specific drivers absent from most literature and Gorsky's (2026) documentation that the most capable cyber-relevant AI models were restricted to a controlled-access program excluding African governments, operators, and universities. Publication bias toward novel attacks over robust defenses (Zhang, 2026) may inflate the apparent severity of defense-focused components, and reliance on an automated, LLM-assisted synthesis (Haryanto et al., 2024) introduces selection and weighting biases on this paper cannot independently detect or correct.

V. CONCLUSION

Artificial intelligence has demonstrably reshaped cyber defense, but the literature documenting this shift remains uneven across the seven components specified here: evidence about real-world reliability is contradictory; knowledge about organizational trust and governance lags technical knowledge; the gap between advocated and actual human oversight is underexplored; methodological diversity and standardized benchmarks are lacking; empirical head-to-head comparisons with traditional defenses are rare; unifying theory is scarce relative to applied technique; and the evidence base under-represents entire regions, organizational sizes, and non-technical populations. This paper's contribution is not only on diagnosis but the framework and methodology that produced it: a seven-component model that is, to the authors' knowledge, the first adaptation of a systematic research-gap taxonomy to AI-enabled cyber defense; a methodology validating that model against nine primary studies across four subdomains; a prioritization procedure stress-tested against three weighting schemes; a procedure for reassessing component severity as the field evolves; a design-science self-evaluation; an explicit positioning against alternative frameworks; and a worked methodology for translating an identified gap into a defensible proposal statement. Responsible use of the framework requires transparency about its reliance on an automated, LLM-assisted synthesis, independent verification of cited claims, adherence to institutional ethical-review and responsible-disclosure practices where dual-use or adversarial techniques are involved, and informed consent and confidentiality protections when collecting data from under-represented populations. Future work should extend Procedure 2's primary-study validation to phishing detection and other unaddressed subdomains and to under-represented regions; test the framework's extendibility and usefulness through independent application by other researchers; and periodically re-apply Procedure 1 to track which components are closing on their own and which require deliberate research investment. With this framework, the paper equips doctoral researchers with a systematic, well-grounded method for identifying a gap substantial enough to anchor a dissertation along with a reusable template for articulating that gap in a research proposal.

REFERENCES

- Achuthan, K., Ramanathan, S., Srinivas, S., & Raman, R. (2024). Advancing cybersecurity and privacy with artificial intelligence: Current trends and future research directions. *Frontiers in Big Data*, 7, Article 1497535. <https://doi.org/10.3389/fdata.2024.1497535>
- Adabara, I., Sadiq, B. O., Shuaibu, A. N., Danjuma, Y. I., Maninti, V., & Joe, M. (2025). Agentic artificial intelligence for ethical cybersecurity in Uganda: A reinforcement learning framework for threat detection in resource-constrained environments. *arXiv*. <https://doi.org/10.48550/arXiv.2512.07909>
- Alahmadi, B. A., Axon, L., & Martinovic, I. (2022). 99% false positives: A qualitative study of SOC analysts' perspectives on security alarms. In *31st USENIX Security Symposium (USENIX Security '22)* (pp. 2783–2800). USENIX Association.
- Albahri, O. S., & AlAmoodi, A. H. (2023). Cybersecurity and artificial intelligence applications: A bibliometric analysis based on Scopus database. *Mesopotamian Journal of CyberSecurity*, 2023, 158–169. <https://doi.org/10.58496/MJCSC/2023/018>
- AlJuhaïman, H. A., Emad-ul-Haq, Q., Kim, K., & Lee, S. (2026). Automated cyber threat intelligence extraction from distributed honeypots: A hybrid machine learning approach. *Electronics*, 15(13), Article 2900. <https://doi.org/10.3390/electronics15132900>
- Alqahtani, H., & Kumar, G. (2026). Large language models for cybersecurity intelligence: A systematic review of emerging threats, defensive capabilities, and security evaluation frameworks. *Computers, Materials & Continua*, 87(3), Article 9. <https://doi.org/10.32604/cmc.2026.077367>
- Alvesson, M., & Sandberg, J. (2011). Generating research questions through problematization. *Academy of Management Review*, 36(2), 247–271. <https://doi.org/10.5465/amr.2009.0188>
- Auwal, A. M. (2026). Operationalizing cyber attack prediction: A gap-prioritized framework with dataset and model selection guidelines. *arXiv*. <https://doi.org/10.48550/arXiv.2606.03386>
- Bade, A. M., Babate, A. I., & Bara, M. W. (2026). AI-driven cyber threats and defences in Nigeria's digital economy: A literature review. *International Journal of Research and Innovation in Applied Science*, 11(2), 1505–1511. <https://dx.doi.org/10.51584/IJRIAS.2026.110200140>
- Bailey, K. D. (1994). *Typologies and taxonomies: An introduction to classification techniques*. Sage.
- Bilal, M. A., Ul Islam, I., Idrees, S., et al. (2026). Dataset-centric evaluation of federated intrusion detection models in IoT networks. *Scientific Reports*, 16, Article 2683. <https://doi.org/10.1038/s41598-025-32567-w>
- D'hooge, L., Wauters, T., Volckaert, B., & De Turck, F. (2020). Inter-dataset generalization strength of supervised machine learning methods for intrusion detection. *Journal of Information Security and Applications*, 54, Article 102564. <https://doi.org/10.1016/j.jisa.2020.102564>
- Elangovan, R., Parthasarathy, D. D., Jawahar, M., Kaliyaperumal, P., Balusamy, B., Yogarayan, S., & Venkatesan, V. (2026). Cross-dataset temporal and semantic generalization of intrusion detection models for the future internet. *Future Internet*, 18(4), Article 194. <https://doi.org/10.3390/fi18040194>
- Falowo, O. I., & Bou Abdo, J. (2026). Empirical study on automation, AI trust, and framework readiness in cybersecurity incident response. *Algorithms*, 19(1), 62. <https://doi.org/10.3390/a19010062>

- Gorsky, M. A. (2026). Artificial intelligence as game changer in cybersecurity: What we learned in 2025-2026, and how this is relevant for Africa. arXiv. <https://doi.org/10.48550/arXiv.2606.20102>
- Gourisetti, S. N. G., Mylrea, M., & Patangia, H. (2020). Cybersecurity vulnerability mitigation framework through empirical paradigm: Enhanced prioritized gap analysis. *Future Generation Computer Systems*, 105, 410–431. <https://doi.org/10.1016/j.future.2019.12.018>
- Haryanto, C. Y., Hartanto, Y., Elvira, A. M., Lomempow, E., Nguyen, T. D., Arakala, A., & Vu, M. H. (2024). Contextualized AI for cyber defense: An automated survey using LLMs (arXiv:2409.13524). arXiv. <https://doi.org/10.48550/arXiv.2409.13524>
- Jacobs, R. L. (2011). Developing a research problem and purpose statement. In T. S. Rocco & T. Hatcher (Eds.), *The handbook of scholarly writing and publishing* (pp. 125–141). Jossey-Bass.
- Kamruzzaman, M., Siddiki, M. S., Mondal, R. S., Saha, S., & Bhuiyan, M. N. A. (2026). Lightweight AI-driven intrusion detection for SMEs: Comparative analysis of machine learning models on TON_IoT. *Intelligent Decision Technologies*, 20(1), 351–367. <https://doi.org/10.1177/18724981251398354>
- Kaur, R., Gabrijelčič, D., & Klobučar, T. (2023). Artificial intelligence for cybersecurity: Literature review and future research directions. *Information Fusion*, 97, 101804. <https://doi.org/10.1016/j.inffus.2023.101804>
- Kuhn, T. S. (1962). *The structure of scientific revolutions*. University of Chicago Press.
- Malatji, M., & Tolah, A. (2024). Artificial intelligence (AI) cybersecurity dimensions: A comprehensive framework for understanding adversarial and offensive AI. *AI and Ethics*, 5, 883–910. <https://doi.org/10.1007/s43681-024-00427-4>
- Merlano, C. (2024). Enhancing cyber security through artificial intelligence and machine learning: A literature review. *Journal of Cyber Security*, 6(1), 89–116. <https://doi.org/10.32604/jcs.2024.056164>
- Merton, R. K. (1973). *The sociology of science: Theoretical and empirical investigations*. University of Chicago Press.
- Miles, D. A. (2017). A taxonomy of research gaps: Identifying and defining the seven research gaps [Paper presentation]. Doctoral Student Workshop: Finding Research Gaps—Research Methods and Strategies, Dallas, TX, United States.
- Müller-Bloch, C., & Kranz, J. (2015). A framework for rigorously identifying research gaps in qualitative literature reviews. In *Proceedings of the Thirty Sixth International Conference on Information Systems* (pp. 1–19). Association for Information Systems.
- Naqvi, W. M., Gabr, M., Arora, S. P., Mishra, G. V., Pashine, A. A., & Quazi Syed, Z. (2024). Bridging, mapping, and addressing research gaps in health sciences: The Naqvi-Gabr research gap framework. *Cureus*, 16(3), Article e55827. <https://doi.org/10.7759/cureus.55827>
- Nickerson, R. C., Varshney, U., & Muntermann, J. (2013). A method for taxonomy development and its application in information systems. *European Journal of Information Systems*, 22(3), 336–359. <https://doi.org/10.1057/ejis.2012.26>
- Popper, K. R. (1959). *The logic of scientific discovery*. Hutchinson.
- Robinson, K. A., Saldanha, I. J., & McKoy, N. A. (2011). Development of a framework to identify research gaps from systematic reviews. *Journal of Clinical Epidemiology*, 64(12), 1325–1330. <https://doi.org/10.1016/j.jclinepi.2011.06.009>
- Rodriguez, M., Popa, R. A., Liang, L., Wang, A., Rahtz, M., Kaskasoli, A., Dafoe, A., & Flynn, F. (2025). A framework for evaluating emerging cyberattack capabilities of AI. arXiv. <https://doi.org/10.48550/arXiv.2503.11917>

- Schrötter, M., Niemann, A., & Schnor, B. (2024). A comparison of neural-network-based intrusion detection against signature-based detection in IoT networks. *Information*, 15(3), Article 164. <https://doi.org/10.3390/info15030164>
- Vieira, C., Vitorino, J., Maia, E., & Praça, I. (2026). Machine learning transferability for malware detection (arXiv:2603.26632). arXiv. <https://doi.org/10.48550/arXiv.2603.26632>
- Ye, W. (2026). Strengthening cybersecurity in small and medium-sized enterprises: Balancing technology, costs, compliance, and employee awareness. *SAGE Open*, 16(1). <https://doi.org/10.1177/21582440251414951>
- Zhang, J., Bu, H., Wen, H., Liu, Y., Fei, H., Xi, R., Li, L., Yang, Y., Zhu, H., & Meng, D. (2025). When LLMs meet cybersecurity: A systematic literature review. *Cybersecurity*, 8(1), Article 55. <https://doi.org/10.1186/s42400-025-00361-w>
- Zhang, Y. (2026). AI security research should better incentivize defense research. arXiv. <https://doi.org/10.48550/arXiv.2605.23448>