



This is a peer-reviewed, final published version of the following document, © The Author(s) 2026. <http://creativecommons.org/licenses/by-nc-nd/4.0/> and is licensed under Creative Commons: Attribution-Noncommercial-No Derivative Works 4.0 license:

**Awan, Mujtaba ORCID logoORCID: <https://orcid.org/0000-0001-9055-2869>, Alam, Abu S, Khan, Rafiq Ahmad, Alwageed, Hathal Salamah, Ayouni, Sarra and Almagrabi, Alaa Omran (2026) A generative AI-driven cybersecurity framework for small and medium enterprises software development: an ANN-ISM approach. Scientific Reports, 16. art 9813. doi:10.1038/s41598-026-37614-8**

Official URL: <https://doi.org/10.1038/s41598-026-37614-8>  
DOI: <http://dx.doi.org/10.1038/s41598-026-37614-8>  
EPrint URI: <https://eprints.glos.ac.uk/id/eprint/15873>

#### **Disclaimer**

The University of Gloucestershire has obtained warranties from all depositors as to their title in the material deposited and as to their right to deposit such material.

The University of Gloucestershire makes no representation or warranties of commercial utility, title, or fitness for a particular purpose or any other warranty, express or implied in respect of any material deposited.

The University of Gloucestershire makes no representation that the use of the materials will not infringe any patent, copyright, trademark or other property or proprietary rights.

The University of Gloucestershire accepts no liability for any infringement of intellectual property rights in any material deposited but will remove such material from public view pending investigation in the event of an allegation of any such infringement.

PLEASE SCROLL DOWN FOR TEXT.



OPEN

## A generative AI-driven cybersecurity framework for small and medium enterprises software development: an ANN-ISM approach

Mujtaba Awan<sup>1</sup>, Abu Alam<sup>1</sup>, Rafiq Ahmad Khan<sup>2</sup>✉, Hathal Salamah Alwageed<sup>3</sup>, Sarra Ayouni<sup>4</sup>✉ & Alaa Omran Almagrabi<sup>5,6</sup>

This paper presents an AI-based generative model to address the cybersecurity threats in software development for Small and Medium Enterprises (SMEs). The model aims to address the unique challenges SMEs face in implementing effective cybersecurity practices by leveraging generative AI to enhance threat detection, prevention, and response. Initially, we conducted a multivocal literature review (MLR) and an empirical survey to identify and validate cybersecurity threats and the generative AI practices used in secure software development for SMEs. An expert panel review was then assigned for the process of artificial neural network (ANN) and interpretive structural model (ISM). The ANN model can predict potential cybersecurity threats by learning from historical data and software development patterns. ISM is used to (1) structure and visualize (2) relations between identified threats and mitigation approaches and (3) offer a combined, multi-layered risk management methodology. A case study was conducted to evaluate the effectiveness of the proposed model. The evaluation has shown that the model significantly enhances SME online security and enables rapid adoption of sophisticated AI-based practices for detecting and responding to primary and advanced cyber threats. Phishing and ransomware received high assessments (Advanced), whereas some advanced techniques, e.g., AI-guided evasion and zero-day attacks, were at early stages of development (Understanding and Development). The general results indicated that generative AI can help organizations enhance SME cybersecurity, and some efforts are underway to develop use cases for advanced threats further. The AI-based generative model is a viable and scalable approach to the cybersecurity of SME software development. Such AI-based practices will enable SMEs to effectively protect themselves against various cyber threats systematically. Future studies should focus on developing contemporary threat strategies and on the impediments to global implementation, particularly in less resource-rich settings.

**Keywords** Software development, Small and medium enterprises (SMEs), Cybersecurity threats, Generative AI practices, Multivocal literature review, Empirical survey, ANN-ISM, Case study

The dynamic nature of cybersecurity has created a major challenge to the Small and Medium Enterprise (SME) in dealing with sophisticated cyber threats. SMEs tend to be affected by cyber-attacks as they have few resources, experience, and infrastructure. The increase in the prevalence of cybercrime and the rise of more advanced and enduring threats (APTs) make the transition to scalable, affordable, and resilient cybersecurity designs a

<sup>1</sup>School of Computing and Engineering, University of Gloucestershire, Cheltenham, UK. <sup>2</sup>Software Engineering Research Group, Department of Computer Science and IT, University of Malakand, Chakdara, Pakistan. <sup>3</sup>College of Computer and Information Sciences, Jouf University, Sakaka, Saudi Arabia. <sup>4</sup>Department of Information Systems, College of Computer and Information Sciences, Princess Nourah bint Abdulrahman University, P.O. Box 84428, 11671 Riyadh, Saudi Arabia. <sup>5</sup>Department of Information Systems, Faculty of Computing and Information Technology, King Abdulaziz University, Jeddah, Saudi Arabia. <sup>6</sup>Future Networks and Cyber-Physical Systems (FuN-CPS)-Research Group, Department of Information Systems, Faculty of Computing and Information Technology, King Abdulaziz University, 21589 Jeddah, Saudi Arabia. ✉email: rafiqahmadk@gmail.com; saayouni@pnu.edu.sa

necessity to SMEs. However, most of the solutions that are in place are sophisticated, expensive, and enterprise-focused, putting SMEs at risk.

This study aims to solve this weakness by using Generative AI models, including Generative Adversarial Networks (GANs), to improve threat detection and threat prevention. The proposed paper offers a Generative AI-based cybersecurity framework, which integrates ANN and ISM framework and is specific to the needs of SME in the field of software development. It is hoped that the proposed research will offer a scalable and inexpensive solution to this issue so that, without stretching their own resources, SMEs would be able to avoid cybersecurity risks.

However, as the digital economy is growing, cybersecurity risks are on the rise with some being predatory to SMEs who is resource poor and less secure<sup>1</sup>. Such risks are phishing attacks, ransomware and malware, software/language vulnerabilities and elaborate social engineering attacks, which can all lead to drastic business continuity plans, data integrity problems, and client trust<sup>2–5</sup>.

The increasing severity and frequency of cyber-attacks have prompted the necessity to develop more dynamic, scaled, and creative solutions to cybersecurity over the past few years<sup>6</sup>. The contemporary cyber landscape is both dynamic and adaptive and the old methodology of signature-based or heuristic approaches are failing to keep up with the changes that are happening with threats<sup>7</sup>. Therefore, one of the most critical issues of most SMEs is the inability to maintain their businesses because of a lack of cash and knowledge<sup>7</sup>. As a result, the demand to develop new models of cybersecurity is increasing, and it should be resilient and scalable, and be able to thwart threats on the spot<sup>8</sup>.

One of the solutions that are new to the creative side is generative AI in cybersecurity<sup>9</sup>. ANNs among other generative AIs can be immensely used to create adaptive security systems that can detect, predict, and counter emerging cyber threats<sup>10</sup>. The concept of generative AI can also be evolved into a continuous learning system that evolves in response to constantly changing threats, which will defend the SMEs with intelligence against a variety of cyberattacks through simulating different attack scenarios and creating solutions<sup>11</sup>. Moreover, the work with the ISM methods will allow organizing the extensive analysis of the interdependence between the parameters of cybersecurity as a whole, which leads to the overall approach to the mitigation of threats<sup>12</sup>.

The paper proposes a Generative AI-based framework to curb the cybersecurity threats of software development among the SMEs. The model is an amalgamation of the ANN and ISM models which provide a steady, modular, and reusable threat detection, prevention and response solution. ANNs—another type of machine learning that imitates the neural networks of the human brain are particularly suitable in finding patterns in big data sets, which explains the fact that the model can be used in cybersecurity<sup>13</sup>. Instead, the ISM approach allows focusing on the interdependence and relations between the types of cybersecurity practices and determining the most suitable synergy of practices in response to a threat<sup>14</sup>.

The model will also provide the SMEs with a holistic and flexible cybersecurity environment that is cost-effective even to the resource-strained organizations. With generative AI in the form of automated threat hunting and response, and a policy to establish the most essential cybersecurity practices, the framework is a novel roadmap to the SMEs in improving their health of security and prevent potential cyberattacks. The first is to suggest a solution that has the technical soundness and takes into consideration the operational limitations and issues that SMEs deal with and is, therefore, covered and applicable to a broad spectrum of sectors.

In this research paper, an AI-generated cybersecurity architecture will be identified and developed, where an ANN is used to identify threats and an ISM to evaluate risks, as the needs of SMEs operating in the software development domain. The framework will provide an efficient, scalable solution to the cybersecurity problems affecting SMEs that will allow them to respond to the constantly changing cyber threats efficiently at a low cost.

### Contribution of the study

In this research paper, we present a novel system of using ANNs and ISM to ensure better cybersecurity in Small and Medium Enterprise (SMEs) in the context of software development by the means of generative AI. The greatest contributions of this work are as follows:

- **Creation of a Generative AI-based Framework:** We suggest an innovative framework, which builds on generative AI practices and is used to improve cybersecurity among SMEs involved in software development. Our framework is dynamic in relation to emergent security risks unlike the conventional ones and therefore offers a dynamism in the provision of defence mechanism that conforms to the changes in the cybersecurity environment. This architecture involves the benefit of combining an ANN to detect anomalies with ISM to aid in strategic decision-making, which provides a complete, customizable cybersecurity solution to SMEs.
- **ANN Use in Cybersecurity:** Implementation of ANN in our system will be a contribution, as it will be able to detect superior threats with the help of patterns and anomalies. We show how ANNs can be used to forecast possible failures during the software development cycle so that SMEs will create a proactive approach to avoid cybersecurity attacks. The application contributes to the growing body of literature about AI in cybersecurity, in this case, when SMEs are involved, the resources are often limited.
- **ISM to Structure Cybersecurity Strategies:** The application of ISM in organizing and prioritizing cybersecurity strategies against SMEs is also another significant contribution. Using ISM can assist us in developing a clear hierarchical image and allow SMEs to recognize the main areas where security should be improved and make correct decisions regarding the use of the resources. It is an emerging application of ISM to cybersecurity strategy development that enables SMEs to make sure they secure against vulnerabilities in a systematic way, depending on the circumstances inherent to their operations.
- **Experimental validation and performance evaluation:** The article is thoroughly experimentally validated by a range of real-life questionnaire surveys and case studies, which prove the proposed framework. Our comparison of the efficacy of the AI-based system of generative systems with the existing cybersecurity mech-

- anisms shows that it is more effective in identifying and containing security threats. Its results indicate the feasibility of the framework and its ability to reduce cyber risk in the case of SMEs.
- **Ablation Study to Test Framework Components:** To further test the soundness of the framework, we conduct an ablation study that isolates and tests the components of the ANN-ISM approach. The comparative importance of each of these aspects to the success of the entire cybersecurity model is drawn in the current paper, which throws light on the most vital aspects that have made it effective. Ablation test proves the benefit of hybridisation of ANN and ISM in the provision of smooth functioning of software development of SMEs.
  - **Practical Implications for SMEs:** The design and experimental findings have the potential to assist the SME to enhance their cybersecurity posture. The work is a response to an acute shortage in cybersecurity needs among SMEs, offering a scalable service that is cost-effective, which uses AI-based approaches to enable the latter to defend their software development seniors without the expensive security infrastructure required.

Overall, the current paper can be described as a valuable contribution to cybersecurity because of the proposal of a generative AI-based model that involves an ANN and an ISM to improve the safety of SMEs. It is reasonable to base the framework on the validation of the experiment, ablation study, and practical implications of the experiment, as it will be a good starting point to the further research and application in the cybersecurity world. The paper is organized in such a way that Section "[Background and related work](#)" includes a review of the related works in the field of cybersecurity in software development among SMEs and their relevant approaches and frameworks, along with the significance of generative AI in cybersecurity. Section "[Hybrid research methodology](#)" expounds on the research methodology and the aspects of the Hybrid ANN-ISM Framework. Section "[Results and analysis](#)" reports the findings on how cybersecurity threats are hierarchically organized and AI mitigation practices are generated. The fifth section provides the development of the Hybrid ANN-ISM Framework to address cybersecurity threats in software development among SMEs taking advantage of generative AI practices. Section "Implications of the study" contains the implications of the study. The limitations of the research are presented in section "[Research limitations](#)". Section "Conclusion and future research directions" gives the conclusion and direction of future research.

### Background and related work

The cybersecurity threats grow more daunting to small and medium-sized enterprises (SMEs) as the software industry develops at a breakneck speed. Cyberattacks can target SMEs which are not always armed and equipped with resources and skills to defend themselves against them<sup>7</sup>. These organizations do not have the tools, personnel, and procedures to defend their software development which is a simple target of opponents who wish to compromise application servers, customer records, and supply chain vulnerabilities<sup>8,15</sup>. The security threats found in the threat landscape of software development are: The use of defective coding patterns, lack of testing processes, unprotected security policies and inability to implement patches in time<sup>1</sup>. Such limitations are likely to be caused by the complexity and dynamism of the software development process, which is constantly changing with the threats and the systems being abused<sup>16</sup>. Although the traditional method of cybersecurity is vital, it may not be sufficient and sufficiently fast to cope with the ever-changing nature of cyber threats particularly the small and medium-sized enterprises (SMEs) who have minimal budgetary capacity to invest in cybersecurity mechanisms<sup>17</sup>. The intersection of artificial intelligence (AI) and cybersecurity has been actively discussed over the past few years, particularly in the context of dealing with the unique issues of SME in software development<sup>18</sup>. According to this view, a variety of scholarly publications suggest integrating the AI models (e.g., adversarial and deep learning and intelligent system-based approaches) to identify and mitigate cyber threats during the software development process<sup>10,19</sup> Table 1.

### AI in cybersecurity

In the area of cybersecurity, there has been a growing use of AI techniques to conduct better detection and prevention of cyber threats<sup>9</sup>. ANN and ML-based models have been explored in case of automatic detection

| Study | Method                               | Context                        | Main Findings                                                                             | Strengths                                     | Limitations                                                                   | Relevance to Current Study                                                                       |
|-------|--------------------------------------|--------------------------------|-------------------------------------------------------------------------------------------|-----------------------------------------------|-------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------|
| 20    | Quantitative (survey, data analysis) | Large enterprises              | Found that GANs can generate synthetic data for cybersecurity, improving attack detection | Empirical, real-world data                    | Focuses on large enterprises, lacks SME applicability                         | Similar approach to use GANs, but tailored for the SME software development context              |
| 21    | Qualitative (interviews, case study) | SMEs                           | SMEs struggle with cost and resource constraints in adopting AI-driven cybersecurity      | In-depth understanding of SME challenges      | No AI-specific solutions offered                                              | Directly related to understanding resource limitations in SMEs                                   |
| 22    | Systematic review                    | AI and cybersecurity           | Identified gaps in using AI/ML models for threat detection in software development        | Comprehensive, covers multiple AI models      | Does not focus on generative AI or hybrid approaches like ANN-ISM             | Theoretical background aligns with the generative AI-driven model discussed in the current study |
| 23    | Case study                           | Software development companies | Demonstrated that hybrid AI models can enhance threat detection accuracy                  | Provides a clear framework for implementation | Focuses only on detection, lacks scalability, and cost-effectiveness analysis | Helpful in understanding hybrid model frameworks, but lacks SME focus                            |
| 24    | Quantitative (survey, data modeling) | SMEs                           | AI adoption in SMEs faces barriers in cost, training, and technical infrastructure        | Empirical with SME-specific focus             | Limited discussion on AI-driven cybersecurity solutions                       | Critical for highlighting the SME context in the current study's framework                       |

Table 1. Related Studies.

of threats, anomaly detection, and attack pattern identification. Khan et al.<sup>13</sup> introduced a piece of work that was based on an ANN to identify software security in design through a behavioural pattern analysis with a higher rate of accuracy compared to the traditional methods. AI models have been used in software development to forecast vulnerabilities through the analysis of large codebases with success<sup>4</sup>. These systems facilitate the automatic detection of vulnerabilities that can be used by cybercriminals and implement preventive measures.

Among the possible opportunities is the application of Generative AI techniques to enhance cybersecurity. GANs have also been employed recently to generate artificial attack data to train cybersecurity models<sup>25</sup>. The generative models could mimic various patterns of attacks and generate useful data sets to be used in the training of the detection systems more effectively. Sharma et al.<sup>26</sup> showed that GANs could generate adversarial samples for training anomaly detection systems, thereby enhancing their robustness against sophisticated cyber-attacks.

Nonetheless, AI-based models and techniques, particularly for software development SMEs, seem to lag far behind these developments. It is common for SMEs to have limited resources to allocate to the development and management of classical cybersecurity solutions, which often require expertise and significant infrastructure. As such, custom-made AI models designed for SMEs that prioritize cost efficiency, automation, and scalability are necessary to meet SME cybersecurity needs<sup>8</sup>.

### ANNs in cybersecurity

ANNs are widely used in machine learning-based cybersecurity solutions for their ability to learn from data and generalize well to new, unseen instances. Various ANN structures have been employed to detect threats in software development environments, such as feedforward networks<sup>27</sup>, recurrent neural networks (RNNs)<sup>28</sup>, deep learning models<sup>29</sup>, etc. Georgios et al.<sup>30</sup> studied deep learning methods in identifying malicious code patterns. Das et al.<sup>31</sup> demonstrated that recurrent neural networks can be applied to real-time detection of vulnerabilities in web applications.

ANNs could also be used in cybersecurity for software development to monitor the behaviour of software systems during the development process<sup>32</sup>. These models can be trained on large amounts of data generated from testing, vulnerability scanning, and code reviews. This would be a great asset to SMEs, as they can use continuous threat monitoring services that require minimal manual effort. In the meantime, organisations and SMEs can leverage the combination of ANNs and automated vulnerability scanning to detect and prevent security vulnerabilities in their software products early<sup>33</sup>. There are some difficulties with ANN-based techniques when used in small-scale software development projects. Their high demand for the labelled data sets, training data, and high computing power limit their implementation to the majority of SMEs. As such, interest has increased in developing more effective ANN architectures, such as lightweight architectures that can be executed on and transferred to low-resource devices<sup>34,35</sup>.

### ISM in cybersecurity

ISM is a widely used methodology for exploring complex systems and identifying relationships among elements or components<sup>12</sup>. In cybersecurity, ISM is increasingly recognized as a valuable tool for understanding complex interdependencies among multiple cybersecurity factors, threats, responses, and mitigations<sup>36</sup>. In this review, we investigate the use of ISM in cybersecurity, reporting on and discussing its efficiency, application, limitations, and potential future developments.

Rajan et al.<sup>37</sup> developed a Modified Total Interpretive Structural Modeling (M-TISM) technique to identify relationships among different factors. They studied the factors affecting cybersecurity management effectiveness, focusing on collaboration, training and resources, capabilities, information sharing, technology knowledge and awareness, and technological infrastructure. It also examines the relationships among the identified factors using the M-TISM approach. Etemadi et al.<sup>38</sup> conducted research to identify and prioritize the barriers to the adoption of blockchain technology in cyber supply chain risk management (CSCRM). To understand the interrelationships among such barriers, the ISM is applied to develop a structural hierarchy for deeper analysis of the interrelated cybersecurity factors involved in implementing blockchain in CSCRM. Khan et al.<sup>36</sup> use the ISM methodology to examine the relationships among the core elements of the requirement engineering practice. They extracted 70 best practices and organized them into 11 core categories to help software development organizations define secure software development requirements. The results of ISM shows that the category awareness of secure requirement engineering has the strongest driving influence among the other 10 core categories of requirement engineering practices. By using ISM evidence, they try to recognize security best practices, which are applicable and can be considered for enhancing the software system's security.

Although many studies have been conducted on AI-based cybersecurity models, the ANN framework, and ISM for large-scale organizations, there is an apparent lack of research on this approach for addressing security threats in a small-scale software development environment. Most available solutions concentrate on large corporations with sufficient resources and lack integrated, automated platforms that are affordable and scalable for SMEs. Moreover, very few investigations were available on the synergy between generative AI and classical ANN and ISM, offering an inclusive, adaptable, and affordable solution for SMEs. We intend to plug these gaps by providing a generative AI-driven model based on the ANN-ISM approach to mitigate the cybersecurity risks in software development in SMEs.

### Hybrid research methodology

In this study, we follow a comprehensive six-phase approach (see Fig. 1) to verify and validate our proposed Hybrid ANN-ISM Framework to reduce the cybersecurity threats in software development for SMEs. The first phase comprises a multivocal literature review (MLR) that draws on perspectives from various sources of knowledge, laying a strong foundation for the study. Phase 2 is a field experiment (online questionnaire) to collect practitioners' opinions and understand the problems and perspectives on the matter. The third phase is



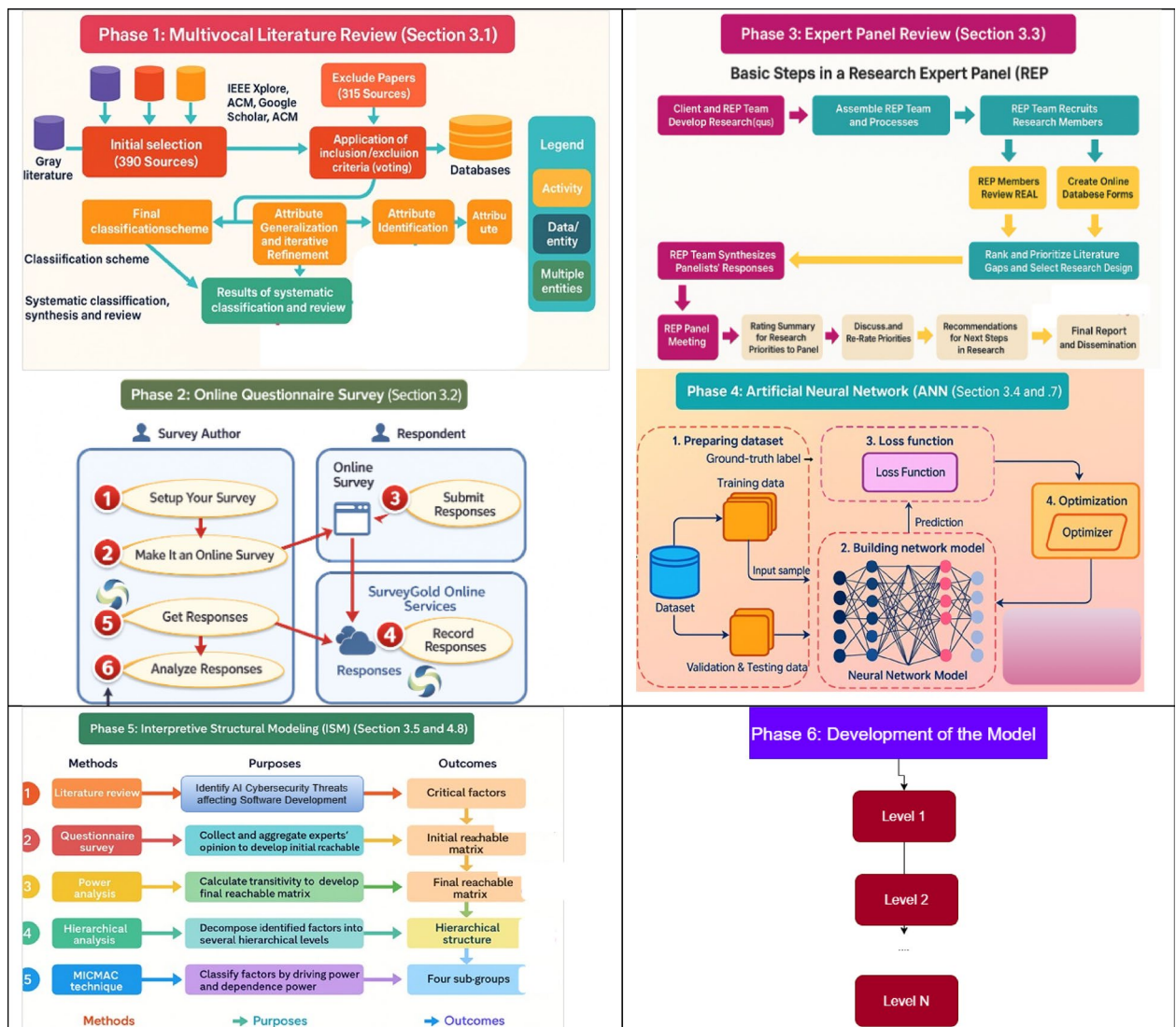


Fig. 1. Research Flow Framework.

an expert panel review to optimize the draft framework through their collective professional wisdom. In the fourth stage, a model is proposed for predicting cybersecurity threats using an ANN. The fifth stage uses ISM for deeper analysis and structuring the relationships among risk factors. Finally, on the sixth stage, a case study will be used to determine whether the suggested approach is feasible and effective in the actual circumstances. The systematic procedure will ensure the all-encompassing and regular examination of the framework capabilities to curb cybersecurity threats in software development among SMEs.

### Phase 1: multivocal literature review (MLR)

A multivocal literature review (MLR) is a comprehensive and systematic review that draws on multiple perspectives, voices, and sources<sup>39–41</sup>. It represents a spectrum of perspectives, approaches, and results across a field. For this paper, an MLR would entail accessing information from various sources, including peer-reviewed papers, conference papers, industry reports, white papers, and expert opinions. The MLR would focus on cybersecurity threats and on Generative AI practices for software development among SMEs.

Here are the specific steps of this paper to perform an MLR<sup>39,42</sup>:

#### Defining the research questions and scope

- **Establish key research questions:** The first step in MLR is to define the study's principal questions. Here, the main research questions are:
  - o What are the primary cybersecurity threats associated with software development?
  - o What are the best Generative AI practices and strategies we should adopt to mitigate these risks?

- **Determine the scope:** This will involve deciding on the boundaries of the review, by defining the special generative AI technologies (for example, input validation and sanitization, GANs, etc.) and the scope of the particular cybersecurity threats (possibly, injection attacks, code quality and logic errors, and malicious code).

#### Searching for sources

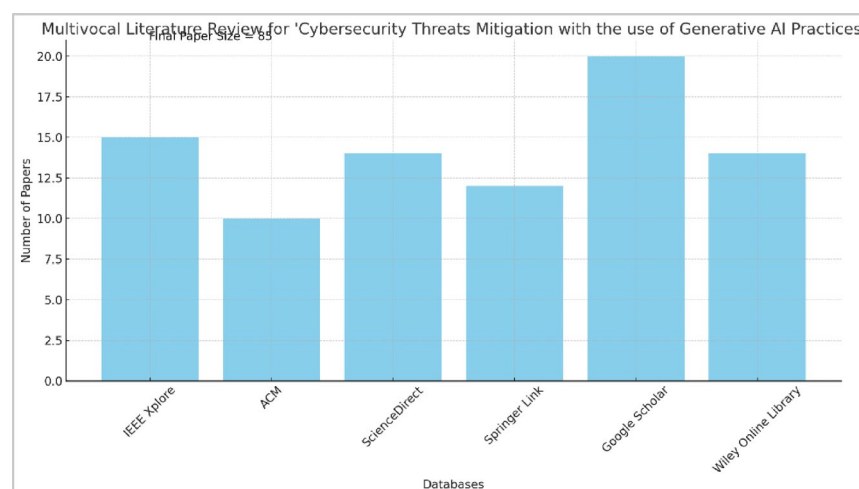
**Find sources:** We look for a diverse range of sources:

- Academic References: Papers on cybersecurity, software development, AI, and Generative Models from high-impact journals and conferences, such as:
  - *IEEE Transactions on Cybersecurity*
  - *Journal of Experimental and Theoretical Artificial Intelligence (JETAI)*
  - *ACM Computing Surveys*
  - *International Journal of Information Security*
  - *Security and Privacy: (Wiley).*
- Industry reports: Announcements from cybersecurity firms, technology companies, and research institutions. Research papers, reports, and white papers from cybersecurity companies, think-tanks, and organizations like:
  - *Gartner*
  - *McKinsey and Company*
  - *OWASP (Open Web Application Security Project)*
  - *ISACA (Information Systems Audit and Control Association)*
  - *National Institute of Standards and Technology (NIST)*
- Government and Regulatory Source: Documents from government departments or standards companies, such as:
  - *EU GDPR Reports*
  - *U.S. Cybersecurity and Infrastructure Security Agency (CISA) Advisories*
- Employ databases: Widely used academic databases such as:
  - *Google Scholar, IEEE Xplore, SpringerLink, ACM, Scopus, etc.*
- Search criteria: query syntax: in specific search term:
  - "cybersecurity risks in software development", "Generative AI practices", "generative models and vulnerabilities", "risk mitigation in software development"

The final sample size (n = 85) is shown in Fig. 2.

#### Screening and selecting sources

- First-level screening: We screen abstracts and titles to include relevant and reliable sources.
- Inclusion criteria:



**Fig. 2.** Final Sample Size.

- o Literature regarding cybersecurity and AI in software development, papers that present solutions to mitigate the identified risks.
- o Recent research papers—Within the past 5–10 years.
- o Consider both cybersecurity threats and AI-specific mitigations.
- o Studies in high-quality peer-reviewed journals and conference proceedings.
- o Updates from reputable cybersecurity firms.
- o Resources and references about generative AI techniques, approaches, models, practices, etc.
- Exclusion criteria:
  - o Non-relevant content about cybersecurity risks or generative AI in software development.
  - o Papers over 10 years old (unless they are seminal).
  - o Non-peer-reviewed sources or opinion pieces.

#### *Data extraction and synthesis*

We extract the following information from the selected papers:

- Cybersecurity risks found: What are the primary cybersecurity risks mentioned in association with software development?
- Mitigation approaches: What are the proposed generative AI practices, methods, or technologies to mitigate risks?
- Emerging Trends: We seek new or novel approaches to secure software development against cybersecurity threats.
- Challenges and gaps: We discuss areas where our literature review highlighted potential gaps or limitations in the current literature.
- Categorize results: We divide our results into several groups, e.g., Cybersecurity risks (e.g., data poisoning, adversarial attacks), Model robustness, and security protocols, best practices, and recommendations in generative AI to protect software development.

#### *Analysis and thematic clustering*

Identify themes and variations within themes: After organizing the data, we looked for themes and variations within themes across sources. For example:

- Security threats in software development.
- Threats of the abuse of generative AI in generating deepfakes or counterfeit content.
- Proactive strategies for mitigating bias include adversarial training, model verification, and an AI ethical framework.
- Cross-source comparisons: We contrasted the sources' conclusions on cybersecurity risks and mitigation strategies.

#### *Synthesizing results and presenting findings*

- Provide a holistic view: We consolidated the most-mentioned cybersecurity risks and generative AI mitigation practices across all sources. Explain the significance of these observations to cybersecurity issues and solutions in generative AI.
- Draw attention to research gaps: We identify unexplored and under-researched topics useful for incoming research. It may be more evidence, new mitigation approaches, or a joint academia-industry partnership.
- Discussion of limitations: We have addressed the weaknesses of the present review (e.g., possible sources bias, insufficient access to databases, or research deficit).

#### *Formulating implications and recommendations*

- Develop implications: Based on the synthesis, we offer practical implications of the research and practice. As an illustration, it may signify where more efforts are needed till we can safely trust software development.
- Practice implications: We provide practical recommendations regarding cybersecurity threats in software development, such as application of certain security standards, regulatory systems or audits of generative AI.

#### *Writing and structuring the literature review*

- Introduction: The paper presents a discussion on the issue of cybersecurity risks and the importance and limitation of generative AI practices in the software development field.
- Methods: We describe the methods of the conduction of the MLR and why various voices and resources were utilized.
- Main body: We present the results under three main themes (exposure, mitigation, and challenges).
- Conclusions: We summarize the main findings, identify research gaps, and make recommendations for future research and practice.

By taking these steps, the MLR contributes to a well-informed, balanced discussion of cybersecurity risks and generative AI methods for addressing them, in which diverse voices and perspectives are heard. This will



be beneficial both academically and by transferring knowledge between academia, industry, and practical application.

### Phase 2: online questionnaire survey

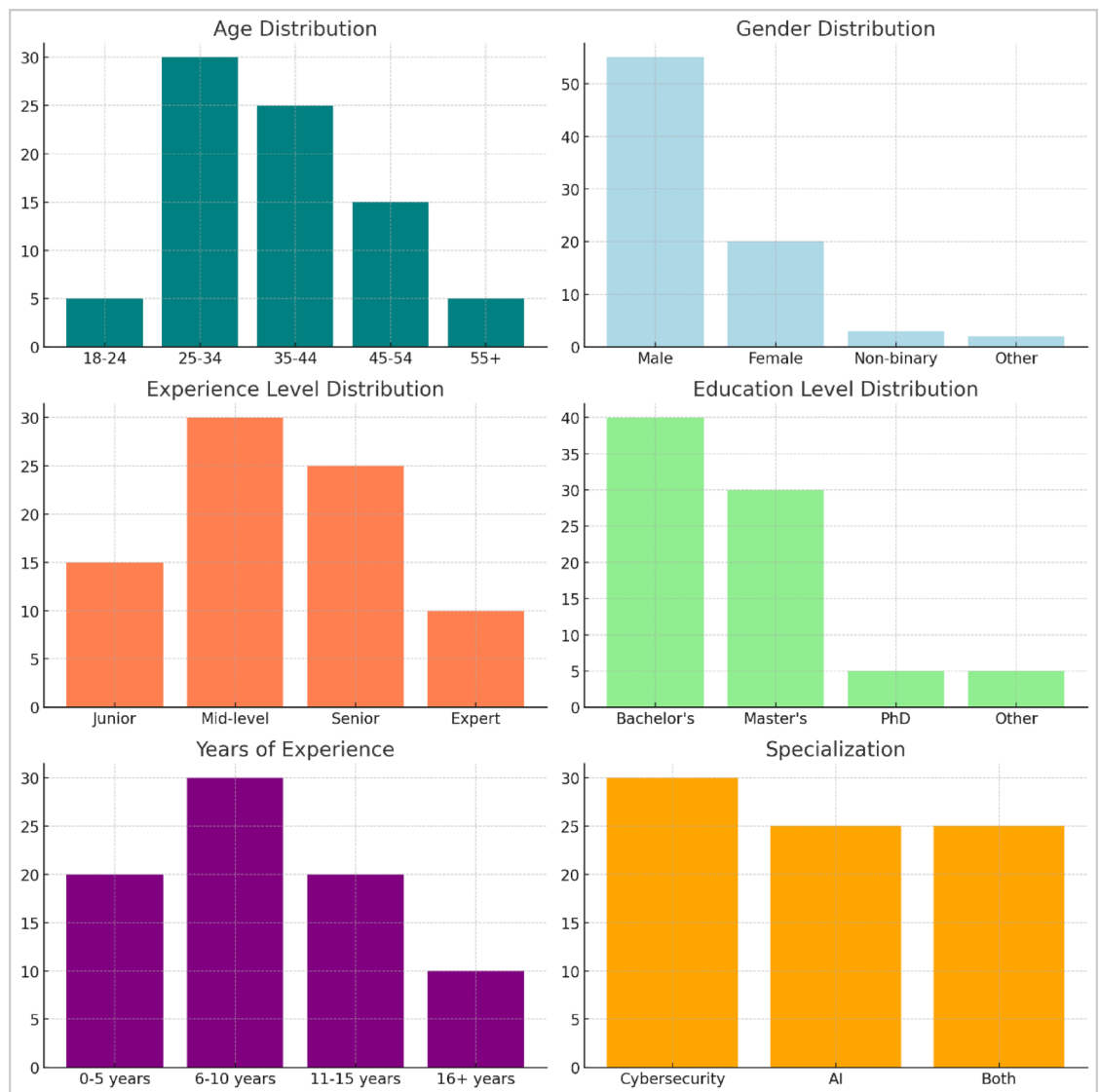
The second step of this research was to design a questionnaire, and several essential elements were considered to ensure a comprehensive and valuable study. The main objective of the survey was to identify the types of cybersecurity threats in software development for SMEs and, as a secondary task, to understand how generative AI can be applied to address them. The following steps were followed in this survey<sup>43–47</sup>:

- This study was conducted in accordance with the ethics guidelines of the University of Gloucestershire Institution Review Board (IRB)/Ethics Committee. The experimental protocols of the study were all in line with the respective ethical guidelines, laws and codes of practice.
- First, we define our intended audience in this survey. We select cybersecurity researchers, software developers, AI researchers, as well as software users and developers in SMEs. The interviewees are partially acquainted with AI, code writing and security. This will make the answers to be based on knowledge and their answers to the objective of the study. The final sample size for this survey is 85 participants.
- The next stage was the development of the questionnaire. The survey covers various categories of questions, including demographic questions, cybersecurity risk identification questions, AI risk mitigation questions, and technology and tool questions. Demographic background questions, how to get basic information about the sample, like their role in the industry, and their working experience. For example, we inquire:
  - o How would you describe your role in the organization?”
  - o How many years of experience in software development/AI/cybersecurity?
- Furthermore, in our questionnaire, questions on identifying cybersecurity risks evaluated the participant's knowledge of typical (security) risks in software development, e.g., code injection, data leakage, and insecure APIs. e.g., a question of concern is,
- What are the most common cybersecurity threats that exist in software development?
- Fig. 3 presents descriptive statistics for the questionnaire respondents. AI practices for risk reduction are indispensable for anyone considering how generative AI is employed to counter threats. We included further questions such as,
- Have you heard of any AI-driven methods for identifying code vulnerabilities?
- What are your thoughts on how generative AI can help counter the cybersecurity challenges associated with software development?
- In addition, technology- and tooling-related questions were used to determine which platforms support automatic code generation and whether they incorporate AI-related security capabilities.
- The survey format consisted of an introduction and a short study profile that provided a précis of the survey's intent, explained how the responses would be used, and indicated the time required to complete the instrument. Also included is an informed consent with a statement explaining that the information provided is confidential and that participants' identities will remain anonymous. The survey was structured according to the same topics: demographics, cybersecurity risks, AI practices, and technology/tools. At the end, participants were thanked, and, if relevant, we mentioned any subsequent events (e.g., a presentation) and the sharing of results.
- Then, finding the right survey tool is also significant. We use online tools like Google Forms to produce and circulate the survey. This product enables us to develop, distribute, and analyze surveys to fit different levels of functionality. When the survey tool is identified, some subjects should be involved in a pilot test. The test was carried out to check the knowledge on the issues of the questions, the length of the survey, and the usability of the tool and to be sure that the survey is functional once it is given to a wider audience.
- The survey was thereafter given to the bigger population of interest after pilot testing and correction of the problems. Here we can appeal to professional circles such as LinkedIn, GitHub, Stack Overflow, or AI/cybersecurity forums and identify the appropriate responders. We also guarantee you a certain date of completion of the survey and there are no worries that the data will be collected in time.
- Seizing control of the information and then going through it. We used the feedback to follow up so that there was nothing going wrong when the survey was being conducted. Then we analysed tendencies, patterns and nuances and cross-read the data. Quantitative data were analysed with the help of statistical packages like SPSS and Excel and data in the open-ended section were analysed through thematic analysis to get themes. This enables us to interpret the results of the survey.
- Lastly, the results were discussed in an organized manner within this paper. The reports deal with the key cybersecurity concerns expressed by the respondents and give an idea of the generative AI techniques that are the most widely used to solve them. To solve the ethical concerns, the privacy and confidentiality of the respondents was also taken into account. The study is carried out based on the ethical guidelines of conducting a research involving human beings, such as informed consent and confidentiality.

### Phase 3: expert panel review

An expert panel reviewed the research presented in this paper. The committee included 23 specialists across cybersecurity, AI, and software development. They represented multiple industries: academia, industry, and research labs, and a combination of professionals with experience in:

- Requirements on cybersecurity and threat management practice.



**Fig. 3.** Demographic Details of Survey Experts.

- More innovative products, and especially products powered by generative AI technologies.
- Secure software development experience.
- AI and cybersecurity: ethical implications.

These individuals have more than 10 years of experience between them; most hold advanced degrees and have worked in leadership roles in their fields. The rigorous process is the study design, which involves rounds of the Delphi. Professionals at each round review the MLR, and the researchers are extensively criticised for the research design and potential improvements. The ANN-ISM framework has a specific scope, as the relevant input from the expert panel is fully incorporated into the research, and the research questions are sharpened.

The researchers evaluated the cybersecurity risk on a risk prism. Perceived lower (approximately 5 percent) and medium perceived importance (approximately 45–50 percent) risks were penalised at 1 and 10, respectively. The other risk scores were rated on a 5-point scale, creating a response scale. These professional opinions are used to develop pairwise matrices that represent the interconnections among various cybersecurity threats in software development within SMEs, as shown in Table 2.

In addition to other analytical methods, ANNs and ISM were applied to ensure the reliability and validity of the research model. These analyses allow for a deeper interpretation of the findings, enhancing their face validity and, ultimately, the strength of the research process.

#### Phase 4: ANN

The fourth stage of this study utilised an ANN process. New data sets are also easy to adapt to. ANNs can handle incomplete or missing input data<sup>48</sup>. The predictions from ANNs are generally better than those from other

methods, such as SEM, multiple linear regression, MDA, and binary logistic regression. ISM is often used to determine the implications of predictors on a predictor variable.

Nevertheless, linear algorithms such as ISM remain restricted to linear mappings of the human decision-making process, as higher-level relations are ignored<sup>49,50</sup>. As a well-known AI model, ANN can address this limitation by learning to model decision-making situations and nonlinear relationships, as stressed by Leong et al.<sup>49</sup>. The multi-layer perceptron structure in an ANN approximates the correlation between inputs and outputs, as the human brain does. One of the significant features of ANNs is their ability to model nonlinear, non-compensatory relationships between attributes<sup>51</sup>.

Overall, ANN models are superior to classical linear procedures and offer greater flexibility and robustness<sup>64</sup>. Nevertheless, ANN cannot be used in attributive analysis or hypothesis testing<sup>49,52</sup>. To address this issue, a two-stage method for integrating the ISM with an ANN has been proposed.

#### ANN training

As an ANN is being trained, we make model intrinsic relationships between the inputs and the outputs by modulating the internal weights, and the input/output pair is represented as:

$$S = (d1, x1), (d2, x2), \dots, (dNi, xNi) \quad (1)$$

The independent data, known as  $x_i$ , and the associated responses,  $d_i$ , are a random sample. These data sets exhibit the non-linear nature of the correspondence between the inputs and the outputs. The goal is to develop an ANN model that can learn this type of invariant relationship on its own. The ANN production is generally in the form of  $w_{ij}y_i + b_i$

$$y = y(x, w) \quad (2)$$

$y$ ,  $x$  and  $w$  are the ANN output, the input parameters, and the unknown weights, respectively. An optimization problem can be solved to determine the best weights that reduces the difference between the predicted output and the actual label. This can be optimised as follows:

$$w^* = \min_x ET = \min_x \|d_i - \gamma(x_i, w)\| \quad (3)$$

where  $ET$  denotes the standard error of the sample. This problem can be approached in many different ways, and the most popular is backpropagation, introduced by Hertz et al.<sup>53</sup>. It is an algorithm that trims the weights of the network by calculating the approximate gradient of the error function concerning the weights, resulting in improved predictions:

$$\omega_{next} = \omega_{now} - \eta \alpha E\tau / \alpha \omega \quad (4)$$

Hertz et al.<sup>53</sup> set the learning rate to  $h$ . Initially, the weights are chosen randomly, and the algorithm is repeated until the optimization condition of Eq. (3) is satisfied. Weights and biases are updated during this process, minimizing mean squared error and allowing the model to achieve the target accuracy.

The weights ( $W_i$ ) and biases ( $b_i$ ) are adjusted until the model achieves the desired accuracy. Alnaizy et al.<sup>54</sup> provide a calibration procedure denoted:

$$V_i = \sum_{j=1}^n w_{ij}y_j + b_i \quad (5)$$

The bias  $b_i$  adjusts the weighted sum of the inputs. A transfer or activation function is then applied to the sum  $V_i$ . This transformation produces the:

$$Z_i = f(V_i) \quad (6)$$

#### Performance of ANN training

The performance of the ANN is evaluated using the Root Mean Squared Error (RMSE), the R-squared ( $R^2$ ), and the Average Absolute Deviation (AAD), expressed as:

$$RMSE = \left[ \frac{1}{n} \sum_{i=1}^n (Y_i - Y_{id})^2 \right]^{0.5} \quad (7)$$

$$R^2 = 1 - \frac{\sum_{i=1}^n (Y_i - Y_{id})^2}{\sum_{i=1}^n (Y_i - Y_{in})^2} \quad (8)$$

$$AAD = \left[ \frac{1}{n} \sum_{i=1}^n \frac{(Y_i - Y_{id})}{Y_{id}} \right] * 100 \quad (9)$$

where  $Y_{id}$  is the observed data;  $Y_i$  is the predicted data;  $Y_m$  is the median of the observed data, and  $n$  is the total number of data.

Phase 5: ISM

The ISM approach was applied in the fifth stage to classify and rank the identified cybersecurity threats in software development for SMEs. The concept of the ISM method, as explained in<sup>55</sup>, was presented to analyze and understand complex relationships among systems and subsystems. By organizing a hierarchy, this method contributes to the acquisition of ability by structuring the variations and directions of various elements. Further, ISM can well model the relationships between visual and structured language<sup>56</sup>. This method is compelling for investigating complex multivariate interactions<sup>57–59</sup>. This approach has been used in many studies better to understand complex systems<sup>60–67</sup>. Figure 1 shows how ISM can be used to map and classify cybersecurity threats in software development for SMEs.

Phase 6: development, implementation, and validation of the proposed model

In the final phase of this research, all findings from phases 1–5 were merged to develop the hybrid ANN-ISM framework for mitigating cybersecurity threats in software development through generative AI practices. The proposed model was then implemented in an organization and was validated through a case study. Further details are presented in Sect. **Evaluation of the Proposed Model**.

Results and analysis

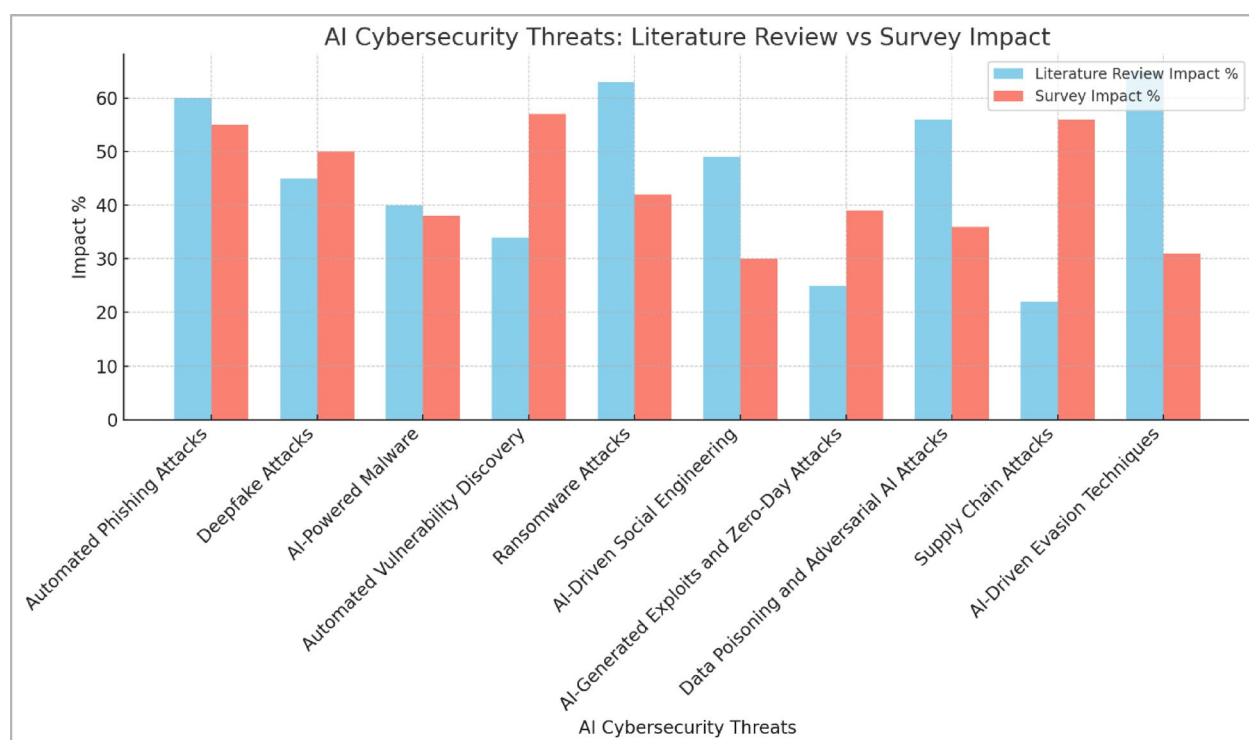
In today’s rapidly changing world of software development, small and mid-sized businesses are increasingly challenged to defend against cybersecurity threats to their operations, intellectual property, and customer data. The rise in advanced cyber threats, such as ransomware, data breaches, and social engineering, has made it imperative that SMEs implement robust security systems. Generative AI methods can provide viable solutions to these issues, especially when it comes to proactive threat detection, automatic vulnerability detection and code security. On the one hand, this paper will discuss the intersection point of cybersecurity threats and, on the other hand, generative AI-based practices and apply the practices to support cybersecurity in the SME sphere. Having AI to predict analytics and pattern recognition capability, SMEs have a chance to enhance their threat mitigation capabilities and safeguard their software development cycles in an environment marked by rapid changes of cyber threats. Table 2 shows AI cybersecurity threats and how they could affect SME software development organizations.

| S. No | AI Cybersecurity Threats                                 |
|-------|----------------------------------------------------------|
| 1     | Automated Phishing Attacks <sup>68,69</sup>              |
| 2     | Deepfake Attacks <sup>70</sup>                           |
| 3     | AI-Powered Malware <sup>71</sup>                         |
| 4     | Automated Vulnerability Discovery <sup>72</sup>          |
| 5     | Ransomware Attacks <sup>3</sup>                          |
| 6     | AI-Driven Social Engineering <sup>73</sup>               |
| 7     | AI-Generated Exploits and Zero-Day Attacks <sup>74</sup> |
| 8     | Data Poisoning and Adversarial AI Attacks <sup>11</sup>  |
| 9     | Supply Chain Attacks <sup>75</sup>                       |
| 10    | AI-Driven Evasion Techniques <sup>76</sup>               |

**Table 2.** AI Cybersecurity Threats in SMEs Software Development Organizations. The information in Fig. 4 is a synthesis of two significant sources, namely, the literature review and survey outcomes:

- **Literature Review:** In an attempt to underpin how perceived different AI cybersecurity threats influenced the software development in SMEs, a literature review of the recent scholarly articles, white papers, and industry reports was undertaken. The sources were also verified in their relevance to the area, and their findings on the frequency and severity of such threats in small and medium enterprises were taken. The literature gave a background of the way the professionals in the field perceived the threats. Figure 4 shows the percentages of impact caused by the literature review which was performed by averaging or synthesizing the information on the threat severity of the said sources with specific focus to the risks posed in the different studies.
- **Data from surveys:** The survey was distributed among SMEs that deal with software development, and such stakeholders as the developers, IT security professionals, and the management are of significance. The questionnaire was to be filled with information regarding the perceived impact of these AI cybersecurity threats necessitating the respondent to rate the perceived severity of each threat on a Likert scale (e.g., 1 to 5). These responses were further coded in percentages in terms of percentage scores achieved by the entire respondents. The average of the results was then obtained to obtain the Survey Impact percentage of each cybersecurity threat illustrated in the figure.

In this case, the impact is considered to be the perceived degree to which individual AI-based cybersecurity threats affect the functioning, protection, and the overall health of the SMEs involved in software development, and the higher a percentage, the more the perceived impact. The intensity of every risk was evaluated in two aspects:



**Fig. 4.** AI Cybersecurity Threats: Literature Review and Survey Impact Percentages.

- **Literature Review Impact:** Depending on the frequency and the severity of each of the threats, as elaborated across academic and industry literature.
- **Survey Impact:** This depends on direct feedback and the ratings of the surveyed SMEs.

The measure of the impact was as follows:

- **Literature Review:** The impact percentages were calculated by summing the severity ratings (e.g., 1–5 scale) across the reviewed articles and translating them into percentages. Research with higher threat severity was accorded greater weight in arriving at the final percentage.
- **Survey Impact:** Survey respondents were requested to rate the severity of each threat on a scale. An average of the responses was used to determine the overall percentage effect shown in Fig. 4. Regarding the illustration, when most respondents ranked a threat as high impact (e.g., 4 or 5 on a scale of 1 to 5), it contributed to a higher percentage in the survey impact.

Figure 4 was built in the following way:

**Literature Review Data Collection:** A systematic review of peer-reviewed articles, industry reports, and other authoritative sources has been conducted. All the essential themes and findings on the AI cybersecurity threats were identified. The severity rating was assigned, and the percentage was determined by averaging the reported impacts across multiple studies.

**Survey Data Collection:** The survey was designed to collect quantitative data from a sample of SMEs. Respondents rated the severity of various cybersecurity threats based on their opinions and personal experience. The data were then used to provide the percentage effect of each threat after aggregation.

**Data Synthesis:** The results of the impact data from the literature review and survey were synthesized to provide a comparative perspective on the perceived severity of the various AI cyber threats. The outcomes of the literature review were presented in blue, and the survey results were presented in red and could be directly compared.

Figure 4 illustrates the level of impact between the research studies and the survey findings regarding the threats to AI security for Small and Medium-sized Enterprises (SMEs). Threat importance Survey analytics demonstrate varying threat severity, since AI-driven threats identified as necessary by academic research are not always aligned with what SMEs rely on in the survey data. Academic sources report threats at lower rates than SMEs do in their surveys.

The SMEs are equally concerned about automated phishing attacks, as mentioned in the Literature Review (60% of the time) and the Survey (55% agreement). The computerized nature of AI in phishing attacks continuously improves their effectiveness and increases their success rate, leaving SMEs to address security considerations and employee awareness training permanently.

In the literature review, deepfake attacks are reported at 45 percent, and in the survey, respondents indicated the same at 50 percent. The nature of Deepfake attacks appears to be more concerning to SMEs than the existing



literature suggests, as they can be used to alter AI-generated media to reflect the image of the company's chief and employees. Impersonation attempts form the sources of social engineering attacks, and that is of greater concern to SMEs.

Based on the research publications and other surveys, there is a high threat of AI-powered malware, 40–38. It is a perfect case since SMEs do not have proper cybersecurity measures that allow attackers to operate unnoticed malware.

The difference between the 34% Literature Review Impact and 57% Survey Impact implies that SMEs would be informed about the vulnerability of the automated system to detect vulnerabilities of AI due to insufficient resources. The fact that such products have these weaknesses is worrying, considering that they can be exploited by many attackers.

The Ransomware Attacks are not eliminated, as in both research works, it was found that the Literature Review indicated a longer effect (63%) than the Survey one (42%). The new ransomware architecture assists hackers to target a particular weakness of the system and modify ransom requests. Such attacks expose SMEs to the risk of having to recover these attacks.

The 49% Literature Review and 30% Survey data suggest that the AI-based social engineering has become a significant threat, as hackers have employed the manipulation strategies in order to make the internal organizational individuals to access the system unauthorized. The survey data suggests that SMEs might not be aware of this threat or worse still, they are less experienced with this threat.

The Literature Review Impact (25) and Survey Impact (39) show that there is a weakness caused by the AI-assisted development of zero-day exploits of unpatched software in organizations. AI rapidly detects and exploits vulnerabilities in business applications, making smaller businesses more susceptible to cyber threats, particularly when they use outdated software or fail to apply patches.

Data poisoning attacks are a significant threat, and with a 56% impact, as reported in the Literature Review. By contrast, SMEs are less aware of adversarial AI attacks, with 36% of respondents reporting an impact, according to the Survey.

The Survey (56%) also shows that the supply chain attacks are of more concern than the Literature Review (22%) assessment of the threat type. The issue of AI-induced supply chain vulnerabilities is of much greater concern to SMEs than other threats, as many attackers embed malicious code in third-party software patches. The reported threat exposure of SMEs in this industry appears to be higher than that reported by academic researchers.

AI-controlled evasion techniques show a greater variety in research findings than in the academic literature, with a Literature Review Impact of 65% and a Survey Impact of only 31%. Scholars in journal articles emphasize the importance of evading security methods driven by AI. However, SMEs are less experienced with these evasion techniques or may lack sufficient detection mechanisms.

Different perceptions of AI cybersecurity threats can be observed by comparing the Literature Review Impact percentage with the Survey Impact percentage. Small and medium-sized enterprises are particularly concerned about automated vulnerability discovery methods, as well as supply chain attacks and ransomware damage in practice. Small and medium enterprises identify AI-powered malware and AI-driven social engineering as low-impact threats, although they officially acknowledge their potential risks. Small businesses must improve their knowledge base and defense capabilities because available academic research does not meet their practical security needs. SMEs need to enhance their threat perception skills to stop the rising tide of AI-based cyberattacks.

Figure 5 includes two linear regression lines: one for the Literature Review Impact and another for the Survey Impact. Both regression lines follow the general form of a linear equation:

$$y = mx + by$$

where:

- $y$  is the dependent variable (the impact percentage).
- $x$  is the independent variable (the type of AI cybersecurity threat).
- $m$  is the slope of the line, representing the rate of change in the impact percentage for each unit change in  $x$ .
- $b$  is the  $y$ -intercept, the value of  $y$  when  $x = 0$ .

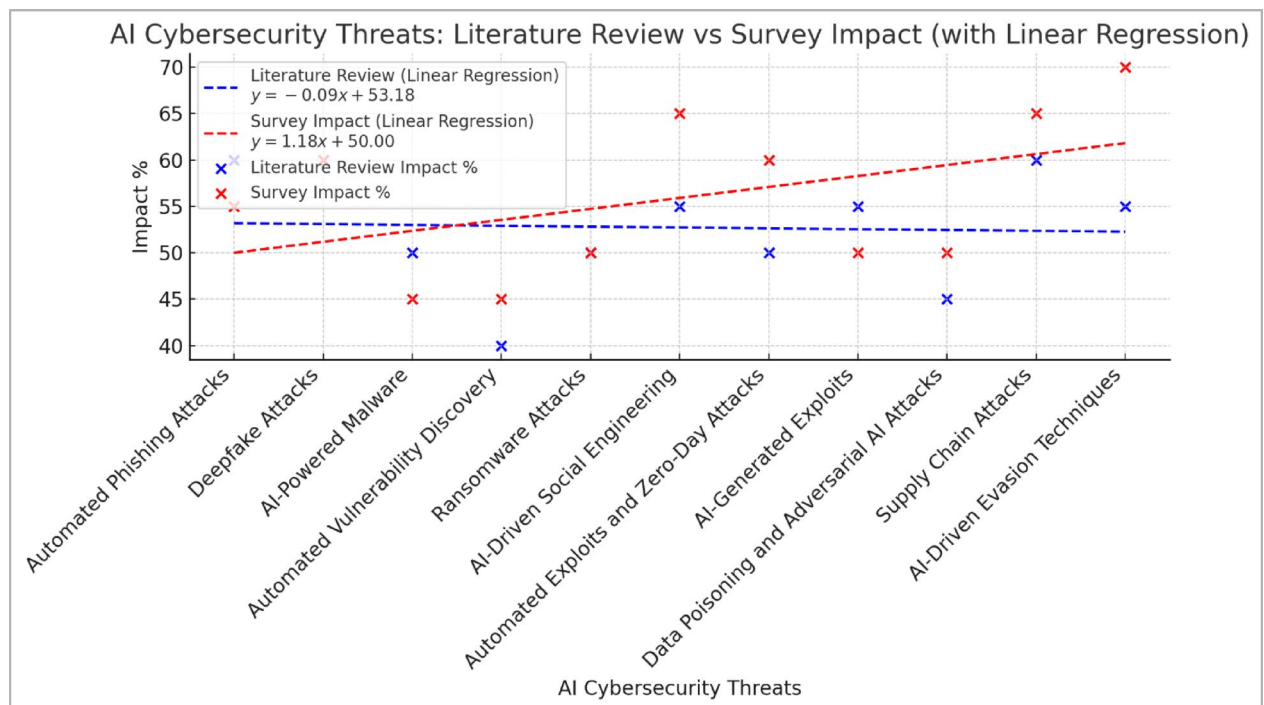
From Fig. 5, the equation for the Literature Review Impact is:

$$y = -0.09x + 53.18$$

- The slope ( $m$ ) of  $-0.09$  indicates that as the value of  $xxx$  increases (moving from left to right on the  $x$ -axis), the impact percentage of the literature review decreases slightly by  $0.09\%$  for each increase in the AI cybersecurity threat type.
- The  $y$ -intercept ( $b$ ) of  $53.18$  suggests that when the AI cybersecurity threat type is at the origin (or if there was a theoretical  $0$  threat), the impact percentage would start at  $53.18\%$ .

The equation for the Survey Impact is:

$$y = 1.18x + 50.00$$



**Fig. 5.** SMEs AI Cybersecurity Threats: Literature Review vs Survey Impact (with Linear Regression).

- The slope (m) of 1.18 indicates that as the value of xxx increases, the impact percentage of the survey response increases significantly by 1.18% for each increase in the AI cybersecurity threat type.
- The y-intercept (b) of 50.00 shows that when the AI cybersecurity threat type is at the origin, the survey impact would start at 50%.

For each AI cybersecurity threat (on the x-axis), the impact percentages (on the y-axis) can be calculated by substituting the specific value of xxx (the type of threat) into each equation.

For example:

For the Automated Phishing Attacks (which would have a specific x value):

Using the Literature Review Regression equation:

$$y = -0.09x + 53.18y$$

$$y = -0.09(1) + 53.18 = -0.09 + 53.18 = 53.09\%$$

Using the Survey Impact Regression equation:

$$y = 1.18x + 50.00y$$

$$y = 1.18(1) + 50.00 = 1.18 + 50.00 = 51.18\%$$

The regression lines help understand the overall trend in how each type of threat affects the cybersecurity field, as measured by the literature and surveys. This is quantified using the equations provided.

### Generative AI practices for mitigating AI cybersecurity threats facing SMEs in secure software development

The relevance of Generative AI Practices to the mitigation of AI Cybersecurity Threats of SMEs in the context of secure Software Development is that the research will reduce the increasing cybersecurity threats faced by SMEs during the development of software. Cyberattacks on SMEs are low hanging fruits because of the low resources and lack of advanced security technologies. As cyber threats become more advanced, the generative AI methods of automatic threat monitoring, predicting vulnerabilities and zero-day security analysis provide SMEs with much-needed solutions in protecting what they hold dear most. The AI-based approaches enable detecting vulnerabilities before any third-party can leverage them, which enhances the safety of the software development process, in general, and minimizes chances of the occurrence of expensive breaches and data loss. Table 3 presents generative AI practices for addressing AI cybersecurity threats in secure software development for SMEs.

| GenAI Practices for Mitigating Automated Phishing Attacks <sup>68,77-81</sup>          |                                                             |                            |                 |
|----------------------------------------------------------------------------------------|-------------------------------------------------------------|----------------------------|-----------------|
| S. No                                                                                  | GenAI Practices                                             | Literature Review Impact % | Survey Impact % |
| 1                                                                                      | AI-Powered Phishing Detection Systems                       | 78                         | 67              |
| 2                                                                                      | Natural Language Processing (NLP) for Email Analysis        | 60                         | 73              |
| 3                                                                                      | Automated URL Analysis and Link Scanning                    | 54                         | 78              |
| 4                                                                                      | AI-Based User Behavior Analysis                             | 67                         | 60              |
| 5                                                                                      | Phishing Simulation and Training                            | 69                         | 56              |
| 6                                                                                      | AI-Based Email Classification                               | 57                         | 58              |
| 7                                                                                      | Multi-Factor Authentication (MFA) Enforcement               | 43                         | 67              |
| 8                                                                                      | Automated Blacklist/Whitelist Maintenance                   | 56                         | 69              |
| 9                                                                                      | AI-Powered Content Filtering                                | 58                         | 48              |
| GenAI Practices for Mitigating Deepfake Attacks <sup>70,82-84</sup>                    |                                                             |                            |                 |
| 10                                                                                     | Deepfake Detection Algorithms                               | 81                         | 67              |
| 11                                                                                     | AI-Powered Content Validation                               | 66                         | 69              |
| 12                                                                                     | AI-Based Watermarking and Metadata Embedding                | 67                         | 57              |
| 13                                                                                     | Real-Time Monitoring of Media                               | 56                         | 41              |
| 14                                                                                     | AI-Powered Secure Communication Channels                    | 73                         | 55              |
| 15                                                                                     | Employee Awareness and Training Programs                    | 78                         | 76              |
| 16                                                                                     | Multi-Factor Authentication (MFA)                           | 65                         | 81              |
| 17                                                                                     | AI-Driven Data Anomalies Detection                          | 41                         | 66              |
| 18                                                                                     | AI-Enhanced Video and Image Integrity Checking              | 55                         | 56              |
| 19                                                                                     | Collaboration with AI Security Solutions Providers          | 76                         | 73              |
| GenAI Practices for Mitigating AI-Powered Malware Attacks <sup>25,30,71,76,85-91</sup> |                                                             |                            |                 |
| 20                                                                                     | AI-Based Malware Detection System                           | 55                         | 84              |
| 21                                                                                     | Behavioral Analysis and Anomaly Detection                   | 76                         | 76              |
| 22                                                                                     | Automated Code Review and Auditing                          | 81                         | 66              |
| 23                                                                                     | Phishing and Social Engineering Attack Prevention           | 67                         | 62              |
| 24                                                                                     | AI-Enhanced Endpoint Security                               | 56                         | 54              |
| 25                                                                                     | Automated Incident Response Systems                         | 73                         | 76              |
| 26                                                                                     | Deep Learning for Predictive Threat Intelligence            | 78                         | 71              |
| 27                                                                                     | Generative Adversarial Networks (GANS) for Malware Analysis | 82                         | 53              |
| 28                                                                                     | AI-driven Threat Hunting                                    | 54                         | 79              |
| 29                                                                                     | Security Patch Management with AI                           | 67                         | 64              |
| GenAI Practices for Mitigating Automated Vulnerability Discovery <sup>92-97</sup>      |                                                             |                            |                 |
| 30                                                                                     | AI-Powered Static Analysis                                  | 67                         | 83              |
| 31                                                                                     | AI-Based Fuzz Testing                                       | 78                         | 68              |
| 32                                                                                     | Automated Code Review and Vulnerability Detection           | 79                         | 78              |
| 33                                                                                     | Machine Learning for Anomaly Detection                      | 80                         | 65              |
| 34                                                                                     | AI-Based Secure Coding Practices                            | 65                         | 67              |
| 35                                                                                     | AI-Driven Threat Intelligence                               | 54                         | 71              |
| 36                                                                                     | AI-Assisted Penetration Testing                             | 78                         | 65              |
| 37                                                                                     | AI-Powered Patch Management                                 | 64                         | 43              |
| 38                                                                                     | Automated Security Audits and Compliance Checking           | 67                         | 56              |
| 39                                                                                     | AI for Security Posture Management                          | 71                         | 73              |
| Generative AI Practices for Mitigating Ransomware Attacks <sup>3,85,98,99</sup>        |                                                             |                            |                 |
| 40                                                                                     | AI-Powered Threat Detection and Prevention                  | 76                         | 87              |
| 41                                                                                     | Automated Vulnerability Scanning                            | 77                         | 48              |
| 42                                                                                     | AI-Based Phishing Detection                                 | 65                         | 56              |
| 43                                                                                     | Behavioral Analytics and User Monitoring                    | 67                         | 76              |
| 44                                                                                     | Predictive Analysis for Threat Intelligence                 | 81                         | 77              |
| 45                                                                                     | Automated Incident Response                                 | 57                         | 65              |
| 46                                                                                     | Data Backup and Recovery Automation                         | 78                         | 39              |
| 47                                                                                     | AI-Driven Encryption Monitoring                             | 65                         | 67              |
| 48                                                                                     | Security Training with AI-Generated Simulations             | 67                         | 81              |
| 49                                                                                     | AI-Enabled Endpoint Protection                              | 71                         | 57              |
| 50                                                                                     | AI for File Integrity Monitoring                            | 80                         | 78              |
| 51                                                                                     | Network Traffic Anomaly Detection                           | 65                         | 76              |
| Continued                                                                              |                                                             |                            |                 |

| GenAI Practices for Mitigating Automated Phishing Attacks <sup>68,77-81</sup>                              |                                                       |                            |                 |
|------------------------------------------------------------------------------------------------------------|-------------------------------------------------------|----------------------------|-----------------|
| S. No                                                                                                      | GenAI Practices                                       | Literature Review Impact % | Survey Impact % |
| 52                                                                                                         | Deep Learning for Malware Analysis                    | 54                         | 66              |
| 53                                                                                                         | AI-Assisted Access Control                            | 67                         | 56              |
| 54                                                                                                         | AI for Blockchain-Based Solutions                     | 81                         | 52              |
| 55                                                                                                         | AI-Driven Deception Technology                        | 57                         | 49              |
| Generative AI Practices for Mitigating Social Engineering Threats <sup>5,100,101</sup>                     |                                                       |                            |                 |
| 56                                                                                                         | Email and Communication Filtering Systems             | 87                         | 81              |
| 57                                                                                                         | Employee Awareness Training                           | 78                         | 67              |
| 58                                                                                                         | Multi-Factor Authentication (MFA)                     | 79                         | 56              |
| 59                                                                                                         | Voice Verification Systems                            | 80                         | 73              |
| 60                                                                                                         | Strict Communication Protocols                        | 65                         | 39              |
| 61                                                                                                         | Employee Education on Social Engineering              | 54                         | 67              |
| 62                                                                                                         | Identity Verification Systems                         | 83                         | 81              |
| 63                                                                                                         | Access Control Policies                               | 68                         | 57              |
| 64                                                                                                         | Continuous Monitoring of Transactions                 | 78                         | 78              |
| 65                                                                                                         | Deepfake Detection Tools                              | 65                         | 66              |
| 66                                                                                                         | Verification Procedures for High-Stakes Communication | 67                         | 62              |
| 67                                                                                                         | Social Media Monitoring                               | 71                         | 54              |
| 68                                                                                                         | Version Control and Code Review Practices             | 56                         | 76              |
| 69                                                                                                         | AI Model Auditing and Logging                         | 73                         | 71              |
| 70                                                                                                         | Access Control on Development Environments            | 39                         | 53              |
| 71                                                                                                         | Personal Data Protection and Encryption               | 67                         | 56              |
| 72                                                                                                         | AI-Based Threat Detection Systems                     | 73                         | 58              |
| 73                                                                                                         | Security Awareness on Tailored Attacks                | 78                         | 67              |
| 74                                                                                                         | Limit Personal Information Shared Online              | 65                         | 69              |
| 75                                                                                                         | Social Media Monitoring Tools                         | 41                         | 48              |
| 76                                                                                                         | Employee Education on Social Media Scams              | 55                         | 78              |
| 77                                                                                                         | AI-Powered Intrusion Detection Systems (IDS)          | 76                         | 66              |
| 78                                                                                                         | Automated Response Systems                            | 54                         | 62              |
| 79                                                                                                         | Behavioral Analysis and Anomaly Detection             | 40                         | 60              |
| Generative AI Practices for Mitigating AI-Generated Exploits and Zero-Day Attacks <sup>25,74,102,103</sup> |                                                       |                            |                 |
| 80                                                                                                         | Threat Modeling with AI                               | 71                         | 67              |
| 81                                                                                                         | Automated Vulnerability Scanning                      | 53                         | 81              |
| 82                                                                                                         | AI-Enhanced Anomaly                                   | 56                         | 41              |
| 83                                                                                                         | AI-Powered Intrusion Detection Systems (IDS)          | 58                         | 55              |
| 84                                                                                                         | AI-Assisted Penetration Testing                       | 67                         | 76              |
| 85                                                                                                         | Machine Learning for Malware Detection                | 87                         | 81              |
| 86                                                                                                         | AI for Patch Management                               | 48                         | 41              |
| 87                                                                                                         | AI-Driven Code Review and Refactoring                 | 56                         | 55              |
| 88                                                                                                         | AI-Powered Threat Intelligence                        | 76                         | 76              |
| 89                                                                                                         | Behavioral Analysis of AI Systems                     | 77                         | 81              |
| 90                                                                                                         | AI-Enhanced Incident Response                         | 65                         | 66              |
| 91                                                                                                         | Generative Adversarial Network (GAN) Detection        | 59                         | 56              |
| 92                                                                                                         | AI-Enhanced Security Logging and Monitoring           | 38                         | 73              |
| Generative AI Practices for Mitigating Data Poisoning and Adversarial AI Attacks <sup>26,104-109</sup>     |                                                       |                            |                 |
| 93                                                                                                         | Data Validation and Preprocessing                     | 81                         | 78              |
| 94                                                                                                         | Robust Model Training                                 | 41                         | 60              |
| 95                                                                                                         | Anomaly Detection                                     | 55                         | 56              |
| 96                                                                                                         | Input Sanitization                                    | 76                         | 58              |
| 97                                                                                                         | Model Explainability and Transparency                 | 81                         | 67              |
| 98                                                                                                         | Data Augmentation                                     | 67                         | 69              |
| 99                                                                                                         | Secure Data Collection                                | 56                         | 48              |
| 100                                                                                                        | Regular Model Auditing                                | 73                         | 60              |
| 101                                                                                                        | Differential Privacy                                  | 39                         | 55              |
| 102                                                                                                        | Adversarial Robustness Evaluation                     | 67                         | 76              |
| 103                                                                                                        | Model Ensembling                                      | 81                         | 56              |
| Continued                                                                                                  |                                                       |                            |                 |

| GenAI Practices for Mitigating Automated Phishing Attacks <sup>68,77–81</sup>             |                                                           |                            |                 |
|-------------------------------------------------------------------------------------------|-----------------------------------------------------------|----------------------------|-----------------|
| S. No                                                                                     | GenAI Practices                                           | Literature Review Impact % | Survey Impact % |
| 104                                                                                       | Redundancy in Data Sources                                | 57                         | 73              |
| 105                                                                                       | Automated Poisoning Attack Detection                      | 54                         | 39              |
| Generative AI Practices for Mitigating Supply Chain Attacks <sup>15,75,110,111</sup>      |                                                           |                            |                 |
| 106                                                                                       | AI-Powered Code Review and Vulnerability Scanning         | 79                         | 66              |
| 107                                                                                       | Automated Dependency Management                           | 80                         | 62              |
| 108                                                                                       | Behavioral Analysis of Third-Party Code                   | 65                         | 54              |
| 109                                                                                       | AI-Driven Continuous Monitoring of Supply Chain           | 54                         | 76              |
| 110                                                                                       | AI-Based Threat Intelligence Sharing                      | 83                         | 71              |
| 111                                                                                       | AI-Powered Incident Response and Automation               | 68                         | 53              |
| 112                                                                                       | Generative AI for Secure Software Design                  | 78                         | 79              |
| 113                                                                                       | AI-Enhanced Risk Prediction and Proactive Threat Modeling | 65                         | 64              |
| 114                                                                                       | Secure Code Generation via AI                             | 48                         | 48              |
| 115                                                                                       | Decentralized Code Signing and AI-Powered Authentication  | 60                         | 60              |
| 116                                                                                       | Supply Chain Risk Assessment Models Using AI              | 55                         | 59              |
| Generative AI Practices for Mitigating AI-Driven Evasion Techniques <sup>10,112–115</sup> |                                                           |                            |                 |
| 117                                                                                       | AI Model Robustness Enhancement                           | 67                         | 69              |
| 118                                                                                       | Anomaly Detection                                         | 81                         | 48              |
| 119                                                                                       | Explainable AI (XAI)                                      | 57                         | 60              |
| 120                                                                                       | Adversarial Training                                      | 78                         | 55              |
| 121                                                                                       | Regular Model Evaluation and Testing                      | 65                         | 76              |
| 122                                                                                       | Robustness Metrics Implementation                         | 67                         | 48              |
| 123                                                                                       | Data Augmentation and Diversification                     | 71                         | 78              |
| 124                                                                                       | Continuous Learning Models                                | 80                         | 66              |
| 125                                                                                       | Secure Data Collection and Labelling Practices            | 65                         | 62              |
| 126                                                                                       | Post-Deployment Monitoring                                | 54                         | 60              |

**Table 3.** Generative AI Practices for Mitigating AI Cybersecurity Threats facing SMEs in Secure Software Development.

| Model Summary |                      |                                                              |
|---------------|----------------------|--------------------------------------------------------------|
| Training      | Sum of Squares Error | 5.263                                                        |
|               | Relative Error       | .957                                                         |
|               | Stopping Rule Used   | 1 consecutive step(s) with no decrease in error <sup>a</sup> |
|               | Training Time        | 0:00:00.02                                                   |
| Testing       | Sum of Squares Error | 2.401                                                        |
|               | Relative Error       | .944                                                         |

**Table 4.** ANN Model Summary. Dependent Variable: AI Cybersecurity Threats a. Error computations are based on the testing sample

ANN modeling

We describe the ANN model’s training and evaluation process in detail to provide clarity and address issues comprehensively. It represents 10 news articles about the 10 AI cybersecurity threats (AICST) relevant to SMEs. To avoid overfitting, the data was split into training and validation sets at 70% and 30%, respectively. Moreover, tenfold cross-validation has been used to ensure the model is valid and not overfit. As shown in Tables 4, 5 and 6, the ANN model has an input layer of 10 features, which are thus included as input elements. The output layer represents the dependent variable and concerns AI cybersecurity risks for SMEs.

The model was trained using the Adam optimizer with a learning rate of 0.001 for 50 epochs and a batch size of 32. The classification model used the cross-entropy loss function. The model’s performance was based on the Root Mean Square Error (RMSE). As shown in Table 7, the RMSE for the test set mean was 0.957, and for the training set, 0.944. Meanwhile, Table 8 also depicts both the importance and normalized importance of AI cybersecurity risks for SMEs.

Figure 6 shows the normalized importance and sensitivity analysis of AI cybersecurity risks to SMEs according to the ANN model, which captured the nonlinear relationships between the independent variables and their influence on the risks (Fig. 6). To better understand this impact, we extend the analysis in Figs. 7 and 8 to examine how predicted output value variations would affect independent variable values. A simple way to



| Independent Variable Importance            |            |                       |
|--------------------------------------------|------------|-----------------------|
|                                            | Importance | Normalized Importance |
| Automated Phishing Attacks                 | .134       | 69.7%                 |
| Deepfake Attacks                           | .083       | 43.2%                 |
| AI-Powered Malware                         | .073       | 37.6%                 |
| Automated Vulnerability Discovery          | .154       | 79.9%                 |
| Ransomware Attacks                         | .168       | 86.9%                 |
| AI-Driven Social Engineering               | .193       | 100.0%                |
| AI-Generated Exploits and Zero-Day Attacks | .044       | 23.0%                 |
| Data Poisoning and Adversarial AI Attacks  | .100       | 51.9%                 |
| Supply Chain Attacks                       | .034       | 17.8%                 |
| AI-Driven Evasion Techniques               | .016       | 8.5%                  |

**Table 5.** Importance and Normalized Importance of Independent Variables (AI Cybersecurity Risks to SMEs).

| Variables        | T1 | T2 | T3 | T4 | T5 | T6 | T7 | T8 | T9 | T10 | Dependence Power |
|------------------|----|----|----|----|----|----|----|----|----|-----|------------------|
| T1               | 1  | 0  | 0  | 0  | 0  | 0  | 0  | 0  | 0  | 0   | 1                |
| T2               | 1  | 1  | 0  | 1  | 1  | 1* | 0  | 0  | 0  | 0   | 5                |
| T3               | 1* | 1  | 1  | 1* | 1* | 1* | 0  | 0  | 0  | 0   | 6                |
| T4               | 1* | 1  | 0  | 1  | 1  | 1* | 0  | 0  | 0  | 0   | 5                |
| T5               | 1* | 1* | 0  | 1  | 1  | 1  | 0  | 0  | 0  | 0   | 5                |
| T6               | 1  | 1* | 0  | 1* | 1  | 1  | 0  | 0  | 0  | 0   | 5                |
| T7               | 1  | 1* | 1  | 1* | 1  | 1* | 1  | 0  | 0  | 1   | 8                |
| T8               | 1* | 1  | 0  | 1  | 1* | 1  | 0  | 1  | 1  | 0   | 7                |
| T9               | 1* | 1* | 0  | 1  | 1* | 1  | 0  | 1  | 1  | 0   | 7                |
| T10              | 1* | 1* | 1  | 1* | 1  | 1* | 1  | 0  | 0  | 1   | 8                |
| Dependence Power | 10 | 9  | 3  | 9  | 9  | 9  | 2  | 2  | 2  | 2   |                  |

**Table 6.** Final Reachability Matrix.

provide a precise estimate of the importance of these risks is to normalize their significance by how changes in the network model's predicted output values affect the independent variables<sup>52</sup>.

*ANN structure and training process*

The ANN used in this study will be organized into three layers: an input layer, a hidden layer, and an output layer (see Fig. 7). This arrangement has been selected to strike a compromise between complexity and computation efficiency.

**Input Layer:** The input layer contains 10 neurons, each corresponding to a specific cybersecurity threat detected through the assessment framework. These input variables, T1-T10, are seen to grasp both the quantitative and qualitative nature of the risks of automatic code generation.

**Hidden Layer:** The hidden layer has four neurons (H1-H4) in it. The non-linear activation function of each neuron helps the model learn complex relationships between the input features and ensures computational stability. The hidden layer is fully connected to the input layer, making all input variables contribute to the learning process.

**Output Layer:** The output layer consists of one neuron that generates the final prediction, the level of risk related to cybersecurity threats in automatic code generation. A good activation function converts the output to a probability score between 0 and 1.

Summary of Model Parameters:

- Input variables: 10 (T1–T10)
- Hidden layers: 1
- Neurons in hidden layer: 4 (H1–H4)
- Activation functions: Non-linear function of the hidden layer, probability function of the output layer.

**Variables: AI Cybersecurity Threats Level**

**Training and Optimization-** The network was optimized using a standard learning algorithm and trained with a method that varies the learning rate and adds momentum to speed up learning. The difference between the predicted and actual risk levels was minimized using the binary cross-entropy function. The model was trained for 100 iterations with a batch size of 16 and an 80:20 test: train split. To avoid overfitting, early stopping was used, and training was terminated when validation loss stopped improving. Its implementation and training were done in Python, and an appropriate deep learning framework was used on a standard computational system.

| Level Partitions |                          |                                |                                       |        |
|------------------|--------------------------|--------------------------------|---------------------------------------|--------|
| Iteration One    |                          |                                |                                       |        |
| Elements (Mi)    | Reachability Set R (Mi)  | Antecedent Set A (Ni)          | Intersection Set R (Mi) $\cap$ A (Ni) | Levels |
| T1               | 1,                       | 1, 2, 3, 4, 5, 6, 7, 8, 9, 10, | 1                                     | 1      |
| T2               | 1, 2, 4, 5, 6,           | 2, 3, 4, 5, 6, 7, 8, 9, 10,    | 2, 4, 5, 6,                           | –      |
| T3               | 1, 2, 3, 4, 5, 6,        | 3, 7, 10,                      | 3,                                    | –      |
| T4               | 1, 2, 4, 5, 6,           | 2, 3, 4, 5, 6, 7, 8, 9, 10,    | 2, 4, 5, 6,                           | –      |
| T5               | 1, 2, 4, 5, 6,           | 2, 3, 4, 5, 6, 7, 8, 9, 10,    | 2, 4, 5, 6,                           | –      |
| T6               | 1, 2, 4, 5, 6,           | 2, 3, 4, 5, 6, 7, 8, 9, 10,    | 2, 4, 5, 6,                           | –      |
| T7               | 1, 2, 3, 4, 5, 6, 7, 10, | 7, 10,                         | 7, 10,                                | –      |
| T8               | 1, 2, 4, 5, 6, 8, 9,     | 8, 9,                          | 8, 9,                                 | –      |
| T9               | 1, 2, 4, 5, 6, 8, 9,     | 8, 9,                          | 8, 9,                                 | –      |
| T10              | 1, 2, 3, 4, 5, 6, 7, 10, | 7, 10,                         | 7, 10,                                | –      |
| Iteration Two    |                          |                                |                                       |        |
| T1               | –                        | 2, 3, 4, 5, 6, 7, 8, 9, 10,    | –                                     | 1      |
| T2               | 2, 4, 5, 6,              | 2, 3, 4, 5, 6, 7, 8, 9, 10,    | 2, 4, 5, 6,                           | 2      |
| T3               | 2, 3, 4, 5, 6,           | 3, 7, 10,                      | 3,                                    | –      |
| T4               | 2, 4, 5, 6,              | 2, 3, 4, 5, 6, 7, 8, 9, 10,    | 2, 4, 5, 6,                           | 2      |
| T5               | 2, 4, 5, 6,              | 2, 3, 4, 5, 6, 7, 8, 9, 10,    | 2, 4, 5, 6,                           | 2      |
| T6               | 2, 4, 5, 6,              | 2, 3, 4, 5, 6, 7, 8, 9, 10,    | 2, 4, 5, 6,                           | 2      |
| T7               | 2, 3, 4, 5, 6, 7, 10,    | 7, 10,                         | 7, 10,                                | –      |
| T8               | 2, 4, 5, 6, 8, 9,        | 8, 9,                          | 8, 9,                                 | –      |
| T9               | 2, 4, 5, 6, 8, 9,        | 8, 9,                          | 8, 9,                                 | –      |
| T10              | 2, 3, 4, 5, 6, 7, 10,    | 7, 10,                         | 7, 10,                                | –      |
| Iteration Three  |                          |                                |                                       |        |
| T1               | –                        | 3, 7, 8, 9, 10,                | –                                     | 1      |
| T2               | –                        | 3, 7, 8, 9, 10,                | –                                     | 2      |
| T3               | 3,                       | 3, 7, 10,                      | 3,                                    | 3      |
| T4               |                          | 3, 7, 8, 9, 10,                | –                                     | 2      |
| T5               |                          | 3, 7, 8, 9, 10,                | –                                     | 2      |
| T6               |                          | 3, 7, 8, 9, 10,                | –                                     | 2      |
| T7               | 3, 7, 10,                | 7, 10,                         | 7, 10,                                | –      |
| T8               | 8, 9,                    | 8, 9,                          | 8, 9,                                 | 3      |
| T9               | 8, 9,                    | 8, 9,                          | 8, 9,                                 | 3      |
| T10              | 3, 7, 10,                | 7, 10,                         | 7, 10,                                | –      |
| Iteration Four   |                          |                                |                                       |        |
| T1               | –                        | 7, 10,                         | –                                     | 1      |
| T2               | –                        | 7, 10,                         | –                                     | 2      |
| T3               | –                        | 7, 10,                         | –                                     | 3      |
| T4               | –                        | 7, 10,                         | –                                     | 2      |
| T5               | –                        | 7, 10,                         | –                                     | 2      |
| T6               | –                        | 7, 10,                         | –                                     | 2      |
| T7               | 7, 10,                   | 7, 10,                         | 7, 10,                                | 4      |
| T8               | –                        | –                              | –                                     | 3      |
| T9               | –                        | –                              | –                                     | 3      |
| T10              | 7, 10,                   | 7, 10,                         | 7, 10,                                | 4      |

**Table 7.** Level Partitioning of Final Reachability Matrix (FRM).

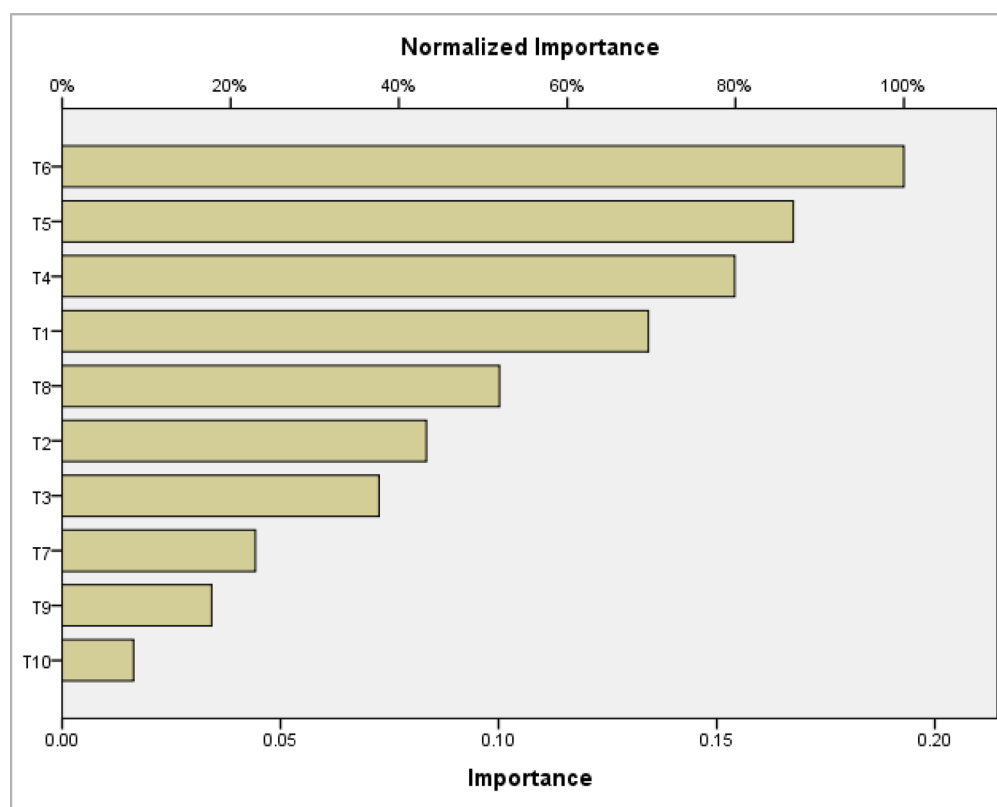
Convergence and high predictive accuracy were consistent across this type of network design. Thus, it can be used to evaluate cybersecurity risks in software development processes—Appendices A and B present detailed views of the data and input variables used by the network.

### ISM findings

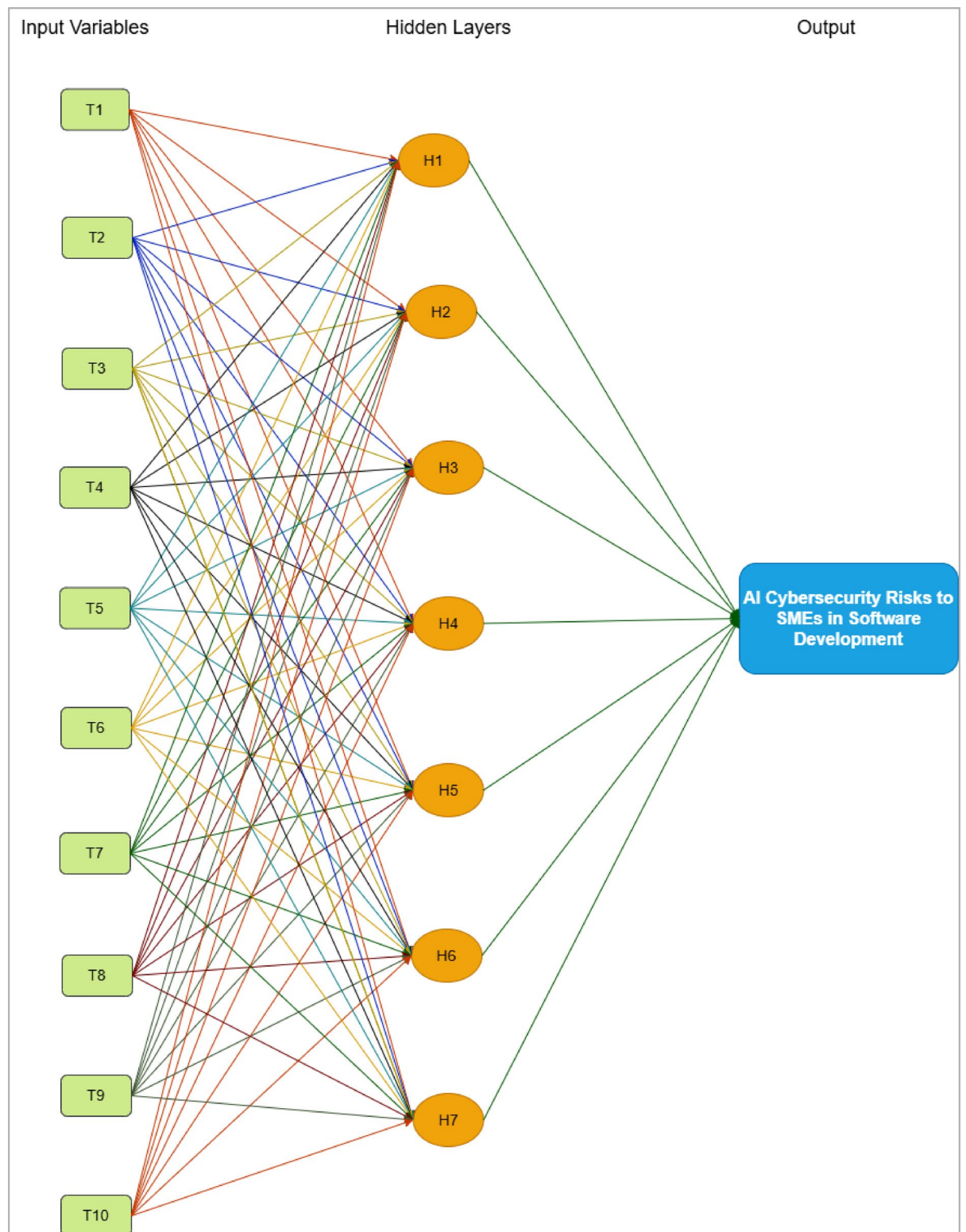
ISM is an empirical method for examining and conceptualizing multifaceted systems, where it is challenging to cognize and locate the association among two or more variables<sup>57,60,116–118</sup>. This is an essential process, particularly for SMEs, to consider cybersecurity threats during software development, as it enables us to trace and chart interdependencies and connections among several factors, and to gain a better understanding of how different factors interrelate to influence the cybersecurity posture as a whole<sup>13,36,63,64</sup>. ISM can also be applied

| Variables        | 1  | 2  | 4  | 5  | 6  | 3 | 8 | 9 | 7 | 10 | Driving Power | Level |
|------------------|----|----|----|----|----|---|---|---|---|----|---------------|-------|
| 1                | 1  | 0  | 0  | 0  | 0  | 0 | 0 | 0 | 0 | 0  | 1             | 1     |
| 2                | 1  | 1  | 1  | 1  | 1* | 0 | 0 | 0 | 0 | 0  | 5             | 2     |
| 4                | 1* | 1  | 1  | 1  | 1* | 0 | 0 | 0 | 0 | 0  | 5             | 2     |
| 5                | 1* | 1* | 1  | 1  | 1  | 0 | 0 | 0 | 0 | 0  | 5             | 2     |
| 6                | 1  | 1* | 1* | 1  | 1  | 0 | 0 | 0 | 0 | 0  | 5             | 2     |
| 3                | 1* | 1  | 1* | 1* | 1* | 1 | 0 | 0 | 0 | 0  | 6             | 3     |
| 8                | 1* | 1  | 1  | 1* | 1  | 0 | 1 | 1 | 0 | 0  | 7             | 3     |
| 9                | 1* | 1* | 1  | 1* | 1  | 0 | 1 | 1 | 0 | 0  | 7             | 3     |
| 7                | 1  | 1* | 1* | 1  | 1* | 1 | 0 | 0 | 1 | 1  | 8             | 4     |
| 10               | 1* | 1* | 1* | 1  | 1* | 1 | 0 | 0 | 1 | 1  | 8             | 4     |
| Dependence Power | 10 | 9  | 9  | 9  | 9  | 3 | 2 | 2 | 2 | 2  |               |       |
| Level            | 1  | 2  | 2  | 2  | 2  | 3 | 3 | 3 | 4 | 4  |               |       |

**Table 8.** Canonical Matrix.



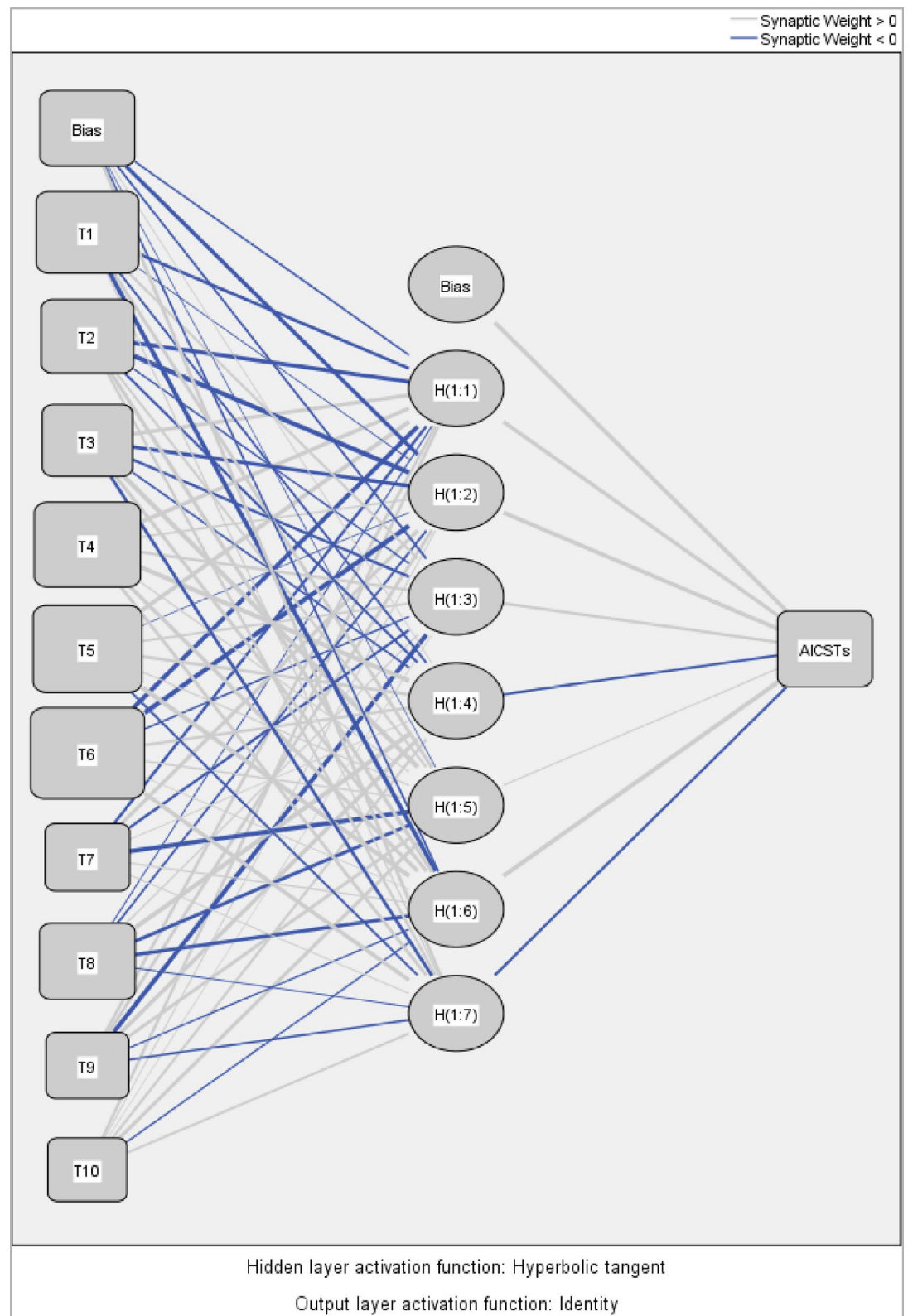
**Fig. 6.** Sensitivity Analysis and Normalized Importance of AI Cybersecurity Risks to SMEs.



**Fig. 7.** Proposed ANN Structure.

to the software development industry to decompose the threats and the associations between them, which can assist in establishing the real causes of a vulnerability or threat, the propagational impact, and how or what can be done to mitigate or control it.

The application of ISM to cybersecurity risk control over automated code generation can be made because it provides a systematic and organized method of defining how different risks combine. The identification of the mutual dependence of different risks will allow ISM to become more rational in risk management, create a more



**Fig. 8.** ANN Model.

systematic risk hierarchy and a working risk-management plan. With the growing popularity of the methods in software development, it is probable that ISM will continue to be a priceless asset in the reliability and safety of the software development systems.



### Structural self-interaction matrix (SSIM)

The outcome of the ISM process is normally a structural self-interaction matrix (SSIM) of the interrelationships between the various threats. The following are some of the questions that such a matrix may be useful in answering: what threats are most effective, what are interdependent and what are the root causes of security problems that have remained within the software development system. Successive refinements of this matrix provide a reachability matrix, which assists in the construction of digraph and thus hierarchical model. The resulting model is a vivid example of how cybersecurity risks interact with each other and the way to help organizations understand which of them to address initially. When you are aware of the cause-and-effect relationships, it is easier to develop more focused mitigation strategies.

Twenty experts with a strong background in generative AI, cybersecurity, and software development were invited to a first-round survey and in-depth discussions. These professionals came from various academic institutions and professional backgrounds. Their ideas were later incorporated into forming the SSIM matrix.

The sample size was small, though, which could limit generalizability; the lack of experts further challenged the comparability of other studies. For example, Kannan et al.<sup>117</sup> received input from at least 5 experts when choosing a reverse logistics provider. Soni et al.<sup>119</sup> on urban rail transit systems, and Attri et al.<sup>120</sup> proposed inviting five specialists to identify pivotal strengths for effective maintenance. Its application to DevSecOps challenge categories, for example, was demonstrated using the ISM method<sup>116</sup>. Other researchers have applied the ISM approach to study DevOps testing<sup>60</sup> and best test practices<sup>57</sup>.

### Analysis of cybersecurity threats and their relationships

The SSIM (Online Appendix A) provides an organized representation of different threats and allows specific risks to be analyzed alongside others to identify causalities and potential threats. As provided in Online Appendix A, the terms nodes and edges refer to distinct components and their relationships within a system of risks and controls.

**Nodes:** Either each node represents **AI Cybersecurity Threats (Ts)** or a **Control** applied to mitigate that threat. Specifically:

**AI Cybersecurity Threats (Ts):** Represented by **T1-T10** in the table, these nodes correspond to specific security threats or vulnerabilities within the system. For example, **T1: Automated Phishing Attacks** and **T2: Deepfake Attacks** are distinct risk nodes that can compromise the system's security.

**Controls:** The letters in the cells (e.g., \*, X, O, A, V) represent various types of power or action associated with each threat. These controls are mechanisms or interventions designed to reduce or address threats. Examples of control types are:

- \* indicates a foundational or inherent threat/control.
- X suggests a control that effectively mitigates the associated threat.
- O implies that there is no direct relevance or application of the control for that threat.
- A signifies an alert or mitigation action.
- V indicates a vulnerability related to the threat in question.

**Edges:** An **edge** represents the relationship between two nodes (either threats or controls), indicating how one node influences the other. The types of relationships include:

**Causal Influence:** An edge denotes a causal influence when a threat or a control directly influences or causes another. For example, the absence of safe code inspection may result in weaknesses in the system. Such a relationship in the table is usually symbolized by X or A, and it means that one threat affects the other, whether it increases or decreases it.

**Prerequisite:** An edge is a prerequisite when one risk or control must exist or occur for another risk or control to be relevant or for action to be taken. This implies that one risk must be addressed before proceeding to the next. This is frequently represented in the table by a star, indicating that the existence or alleviation of one risk is assumed to underlie the assessment of another.

**Amplification:** An amplification occurs when an effect or probability of a risk is increased (or reduced) by the occurrence or alleviation of another risk or control. Solving one problem can expose another. This relationship could be denoted by V, in the sense that managing one risk may reveal or increase other risks, or by A, in the sense that the control increases the mitigation of the associated risks.

Example Relationships:

T2: Deepfake Attacks (Row 2)

**T4 (Column 5: Automated Vulnerability Discovery):** The relationship with the label of X implies that mitigating the deepfake attacks could be achieved through the enhancement of automated vulnerability discovery.

**T3 (Column 4: AI-Powered Malware):** The relation is O, which implies that AI-powered malware is not directly related to preventing an injection attack in this case.

T1: Automated Phishing Attacks (Row 2)

- o **T2 (Column 3: Deepfake Attacks):** This is denoted by an A, indicating that automated phishing may raise concerns about deepfakes.
- Summary of Relationships:
- o **Causal Influence (e.g., X and A):** There is a threat or control that is directly related to the happening, or intensity, of another.
- o **Prerequisite (e.g., \*):** Before one threat or control may be assessed or mitigated, there must be the presence of another threat or control.
- o **Amplification (e.g., V):** The mitigation of one threat can amplify the threat, have an amplifying effect, or be more effective.



This structure helps explain the interaction between threats and controls within the system, serving as a roadmap for identifying successful intervention points and potential threats in software development from the perspective of SMEs.

Initial reachability matrix

To construct the Initial Reachability Matrix (IRM), we followed the standard ISM conversion rules that translate the *directional symbols* (V, A, X, O) into binary form:

| Symbol | Meaning                                                                        | IRM Conversion           |
|--------|--------------------------------------------------------------------------------|--------------------------|
| • V    | $T_{i < sub > i < /sub >}$ influences $T_{j < sub > j < /sub >}$               | $(i, j) = 1; (j, i) = 0$ |
| • A    | $T_{j < sub > j < /sub >}$ influences $T_{i < sub > i < /sub >}$               | $(i, j) = 0; (j, i) = 1$ |
| • X    | $T_{i < sub > i < /sub >}$ and $T_{j < sub > j < /sub >}$ influence each other | $(i, j) = 1; (j, i) = 1$ |
| • O    | No relation between $T_{i < sub > i < /sub >}$ and $T_{j < sub > j < /sub >}$  | $(i, j) = 0; (j, i) = 0$ |

By applying these rules, we designed an initial reachability matrix (Online Appendix B). “1” represents a directional influence from row ( $T_{i < sub > i < /sub >}$ ) → column ( $T_{j < sub > j < /sub >}$ ).

The next step is to compute the Final Reachability Matrix (FRM). We start from the Initial Reachability Matrix (IRM) and apply the transitivity rule of Interpretive Structural Modeling (ISM):

$T_i \rightarrow T_j$  and  $T_j \rightarrow T_k$ , then  $T_i \rightarrow T_k$  (even if no direct link exists in the IRM). The following are the step-by-step conversion:

- Step 1: Recal–IRM Summary (Direct Relationships)  
We already have direct influences between cybersecurity threats (T1-T10)  
All diagonal elements = 1
- Step 2: Apply Transitivity  
Each indirect connection (via another T) is converted to 1. This is typically done using the Boolean Matrix multiplication approach (AND/OR).  
After applying transitivity, every cell (i, j) becomes:  
1 if there is any direct or indirect path from  $T_i \rightarrow T_j$   
0 otherwise
- Step 3: Final Reachability Matrix (Table 6)
- Step 4: Interpretation  
The matrix now includes both direct and transitive (indirect) influences.  
A “1” means  $T_i$  has some influence (direct OR indirect) on  $T_j$ .  
This FRM is used to derive:  
Reachability Set (all threats influenced by a given T).  
Antecedent Set (all threats influencing a given T).  
And hence the hierarchical levels (Level 1, Level 2, Level 3, and Level 4, etc.) of the ISM Model.

Table 8 presents the Cross-impact matrix of different cybersecurity threats and their dependencies, with rows and columns as different open risks (T1 to T10). The table cells indicate that a relationship or dependency exists between two threats (i.e., “1” or “1\*”). If a “1” indicates a direct relationship, then a “1\*” might indicate, by contrast, a less direct relation between the entities. The Dependence Power column on the right side of each row represents the number of connections a threat has with all others.

Partitioning the reachability matrix

According to Warfield<sup>121</sup>, the reachability set of a variable includes the variable itself and any other variables that contribute to achieving its goal. The intersection of such sets is computed component-wise. Elements with the same connectivity and intersection are kept at the top level of the ISM tree. Working this hierarchy down, we address the high-level attributes first. Once these features have been found, they are subtracted, and the process is repeated to isolate the following degree. This loop is repeated until a complete hierarchy of all elements is known. These levels are crucial for constructing the ISM model and the diagram.

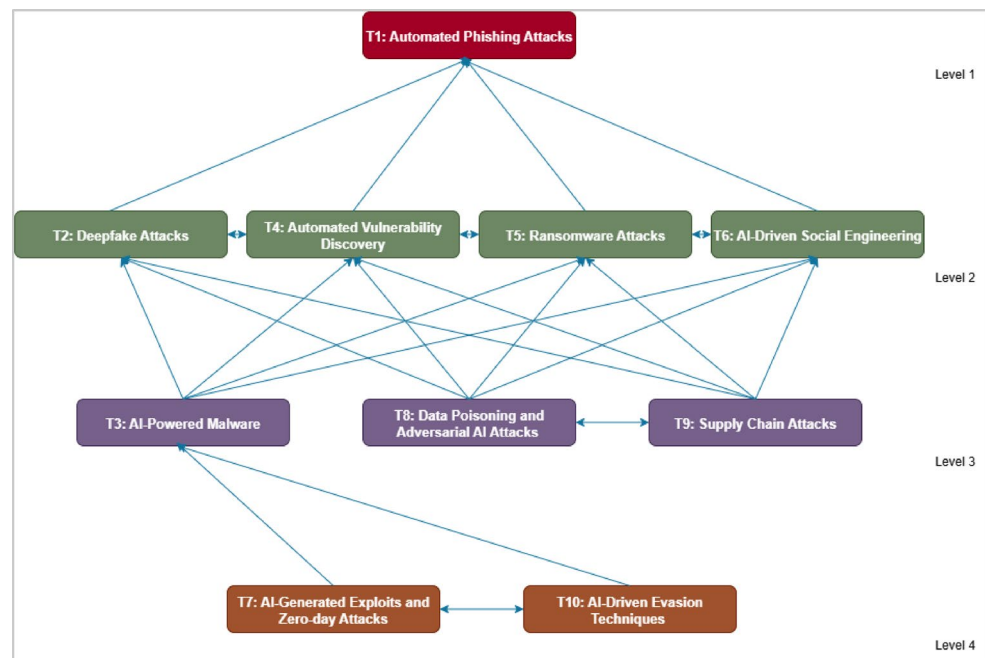
Table 7 defines a four-level Level Partitioning Process that attempts to arrange cybersecurity threats according to their interdependence: reachability set (R), antecedent set (A), and intersection set ( $R \cap A$ ). In each iteration, the matrix further classifies each threat as high, medium, or low based on its interactions with other threats.

Interpretation of the ISM model

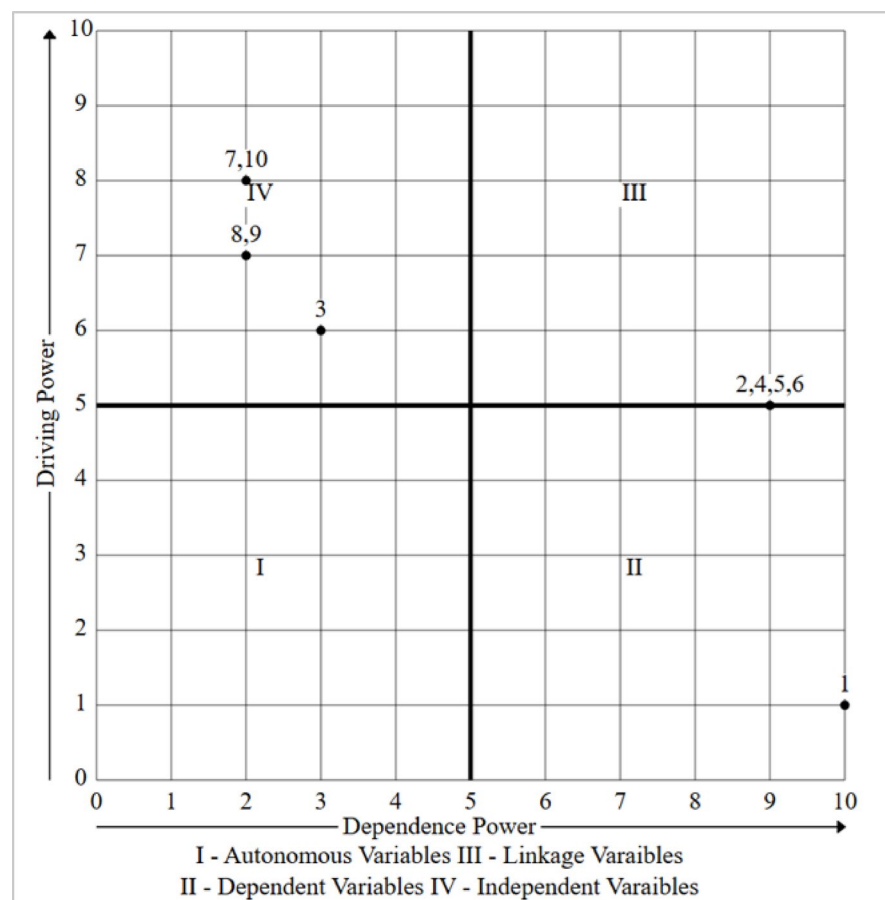
The ISM model was built using the final reachability matrix. Arrows connecting the criteria indicate their interrelatedness. After converting the digraph into the ISM model (Fig. 9), a transitivity analysis was conducted to identify potential ambiguities in the data.

MICMAC analysis

The first one is known as the matrix cross-impact matrix and its name is MICMAC, which is used to examine the key parts and the types of a system. Based on the example of Attri et al.<sup>120</sup>, the technique builds a graph of clumping of factors based on their participation and reliance. The objective of this MICMAC analysis is to cluster these threats as well as to make sure that the results got through interpretative structural modelling are accurate<sup>70</sup>. From this process, the enablers are categorized into four namely, independent, dependent, autonomous as well as linkage variables as illustrated in Fig. 10. Such a classification assists in making clear how each variable fits in the system.



**Fig. 9.** Levels of the Proposed Model.



**Fig. 10.** MICMAC Analysis of the Cybersecurity Threats.

| S. No | Range value in % by IBM | Range of Median value for Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs Using ANN-ISM | Mitigation Level |
|-------|-------------------------|------------------------------------------------------------------------------------------------------------------------------------------|------------------|
| 1     | 0-15%                   | $0 < \text{Median} \leq 0.45$                                                                                                            | Novice           |
| 2     | 15-50%                  | $0.45 < \text{Median} \leq 1.5$                                                                                                          | Comprehension    |
| 3     | 50-85%                  | $1.5 < \text{Median} \leq 2.55$                                                                                                          | Development      |
| 4     | 85-100%                 | $2.55 < \text{Median} \leq 3$                                                                                                            | Advanced         |

Table 9. Cybersecurity Mitigation Levels according to RUP defined by IBM.



Fig. 11. Levels of Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs.

Canonical matrix

The purpose of the MICMAC analysis is to develop a conic matrix. Tables 8 and 9 were used to form the conical matrix in Table 8.

Development of a generative AI-Driven model for mitigating cybersecurity threats in software development for SMEs Using ANN-ISM methodologies

The development of a generative AI-driven model for mitigating cybersecurity threats in software development for SMEs using ANN-ISM methodologies is based on the Secure Software Design Mitigation Model<sup>13</sup>, Sustainable Cloud Computing Model<sup>14</sup>, AI-Driven Cybersecurity Framework<sup>10</sup>, 5G Networks Security Mitigation Model<sup>12</sup>, SAMM<sup>122</sup>, BSIMM<sup>123</sup>, and SCCCMM<sup>14</sup>. The framework stipulates four levels, each comprising several process areas, which are the basis of the model. The methodology developed for the proposed framework is presented in Fig. 11. The model has been built up progressively as follows:

To initiate the framing, data for the ANN-ISM are gathered. Data for the ANN is gathered from questionnaires, academic/field research, and publications, and is used collectively as a knowledge base to aid the system's training. We ensure that the data is processed correctly, measured consistently, and accurately. ISM data is gathered through subject matter experts who give their opinions on the impact of cybersecurity threat in software development on SMEs online interviews.

Then, training of ANN and ISM models is developed. The ANN will deal with qualitative data, which is likely to be able to learn by inputs to alert of possible cybersecurity threats on the auto-code generation. The ISM model is constructed based on the analysis of factors of quality and the comprehension of the effects that cybersecurity threats have on the software development of SMEs.

Using ANN and ISM together is a new step, which allows ANN to play a predictive role and gives a more detailed solution to ISM. ANN predicts the possible threats, ISM to have a systematic description of the activities, and to connect with different cybersecurity threats. The architecture is not only deeply tested on various data sets but also confirmed by cybersecurity and AI professionals, which proves that it can recognise and remove cybersecurity-related threats on any software used. When validation has been successful, then the framework is

prepared to be used. It combines the predictive power of ANNs with the analytical capabilities of ISM to provide a robust approach to mitigating cybersecurity risks in software development for SMEs.

The progress of the hybrid ANN-ISM Framework developed to reduce cybersecurity threats in software development for SMEs using AI generative is depicted in Fig. 11. The model is organized into four layers, each covering certain process areas for establishing the software development for SMEs. In the following, the different levels are detailed and described, along with how they guard against cybersecurity threats in the software development workflow for SMEs.

Figure 11 provides a framework for applying Generative AI to mitigating cybersecurity threats in software development at four levels. The stages reflect different levels of threat sophistication, attack techniques, and how Generative AI powers solutions based on these levels.

- Level 1: Foundation:
  - Automated Phishing Attacks. Phishing can be a technical system compromise to obtain sensitive information from its victims. Also, at this stage, AI could be used to detect phishing attempts, continue to raise your consciousness of what is suspicious, and automatically respond to this type of threat.
- Level 2: Attack Methodologies: At this point, the model begins to consider advanced attacks.
  - Deepfake Attacks: The attacker uses AI to create fake audiovisual content to deceive the victim (e.g., fake videos or audio files).
  - Autonomous vulnerability discovery: With AI, we can automatically detect potential flaws in software systems, allowing attackers to exploit them more quickly.
  - The Ransomware Threat: AI-powered software can threaten industries by exploiting system weaknesses, locking essential data, and demanding a ransom.
  - AI-powered Social Engineering: AI-enabled solutions can create accurate, effective social engineering attacks by mining data to trick victims into revealing sensitive information.
- Level 3: Advanced Threats: The next level is more sophisticated, AI-powered threats:
  - AI-Powered Malware: This malware uses AI to mutate on the fly, making it immune to signature-based defenses.
  - Data Poisoning and Adversarial AI Attack: These attacks involve injecting malicious data into an AI system to disrupt its results or undermine the performance of machine learning models.
  - Through the Supply Chain: With AI-enabled attacks, the supply chain can be attacked, such that network attacks can occur on a much larger scale than was originally possible."
- Level 4: Cutting-Edge Techniques. The last stage includes the most advanced attacks, which use AI in novel ways:
  - AI-Generated and Zero-Day Exploits: This is where AI will create new exploits aimed at unknown weaknesses in a system (zero-day vulnerabilities).
  - AI Bypass Techniques: AI techniques or methods to evade detection, or otherwise bypass, traditional defenses, such as firewalls or intrusion detection systems.

Overall, Fig. 10 illustrates the escalating sophistication of AI-enabled threats and underscores the need for a modern cybersecurity approach that addresses the evolving threats (and risks) in software development for SMEs. Each level builds on the one before it: beginner automated phishing scams give way to top-end AI exploits.

## Evaluation of the proposed model

The structure of the generative AI-driven model for mitigating cybersecurity threats in software development for SMEs using ANN-ISM is divided into four different assessment steps:

- Novice: The organization begins focusing on identifying software cybersecurity threats. The quality of this level is between 0 and 15%.
- Comprehension: This level addresses the documentation and work-to-rule of cybersecurity threat mitigation measures in software development for SMEs. The qualitative score for this stage ranges from 15 to 50%.
- Development: At this stage, the focus is on automating systems and refining software development for SMEs. The qualitative grade for this grade range is 50–85%.
- Advanced: In this stage, the company performs a complete examination, improvement, and elaboration of the security strategy for the SMEs. The qualitative index ranges from 85 to 100%.

To measure the effectiveness of our process domain and practices, we have adopted the SCAMPI<sup>124</sup> approach. The given model employs an evaluation scale based on the IBM Rational Unified Process (RUP), as tabulated in Table 9. The RUP uses a numeric scale in which 0 denotes "no knowledge," and 3 indicates "complete comprehension." Each mitigation measure is given a score, and the median (50) is used to describe the central tendency of the group's scores. This mean is then used to establish the level of development of the respective category in general, ensuring that the scores do not fall within the four levels of RUP and avoiding overlap in levels of mitigation. This approach maintains the distinction between maturity levels and the predictability guarantee, thereby preserving the integrity of the maturity measure the model provides.

| S. No | Structure of the Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs Using ANN-ISM | Agree |     | Disagree |    | Undecided |    |
|-------|---------------------------------------------------------------------------------------------------------------------------------|-------|-----|----------|----|-----------|----|
|       |                                                                                                                                 | N     | %   | N        | %  | N         | %  |
| 1     | Each level of the proposed model is clear and requires no explanation.                                                          | 2     | 50  | 1        | 25 | 1         | 25 |
| 2     | Each level is feasible for cybersecurity in software development for SMEs.                                                      | 4     | 100 | 0        | 0  | 0         | 0  |
| 3     | This model can be used to pinpoint where software development for SMEs has room to improve its cybersecurity risk mitigation.   | 4     | 100 | 0        | 0  | 0         | 0  |
| 4     | It is helpful to divide cybersecurity risks in software development into different levels.                                      | 3     | 75  | 1        | 25 | 0         | 0  |
| 5     | The model's four levels are useful.                                                                                             | 4     | 100 | 0        | 0  | 0         | 0  |

**Table 10.** Evaluation of Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs in Academia.

The Software Engineering Research Group (SERG) at the University of Malakand, Pakistan, conducted a pilot trial of the model. Eight faculty members (three professors, two associate professors, and three assistant professors) were the participants of the trial. These customers were given a document about the proposed model and were asked to provide their feedback. Table 10 summarized their responses and helps assess the model's structure.

Moreover, the article provides a case study of a well-known AI-based software development provider for SMEs to confirm the model's applicability in practice. The members of this case study were the leaders of the company's automatic code generation, cybersecurity, quality assurance, and configuration teams. The researchers were provided with relevant documents and information, and the case study approach used in previous studies<sup>125–130</sup> was employed to collect and analyze data. An Excel checklist was developed to organize the model's categories, processes, and practices across the different levels of mitigation.

The findings, as demonstrated by the evaluation through Table 11, point to the following results as found by the team at the company;

- The company also uses the conventional approach at the moment and is less concentrated on software development security among SMEs.
- The software development configuration is documented. There is a possibility of improving software development for SMEs.
- The firm is working on refining and improving secure configuration procedures for software development with SMEs.

Table 11 summarizes the application of Generative AI practices to mitigate cybersecurity threats in software development for SMEs. A use case is illustrated across levels, process areas, and cybersecurity mitigations to evaluate the efficacy of generative AI approaches in countering ubiquitous threats such as phishing, deepfake attacks, and ransomware. In Table 11:

- ID: The identifier of the process area (PA).
- PA (Process Areas): These are attack types or attack approaches that generative AI can help in mitigating.
- Cybersecurity Mitigation Levels: Objectives are organized into four levels—that is, how much the company has advanced in addressing cybersecurity threats:
  - Novice (0): No, or partial, enforcement of the AI-based mitigation measures.
  - Understanding (1): There are some practices, but they may be minimal or rudimentary.
  - Development (2): Some practices are in place and are currently in implementation and development, or at a development stage.
  - Advanced (3): The best, most established, and effective mitigation actions.
- Methods or techniques (e.g., workflows, AI-enabled phishing-detection systems, or deepfake-detection algorithms) are defined for each process area, and the company's adoption of these methods is scored against one of the mitigation levels.
- Levels of Evaluation: 4 levels of assessment are listed in Table 11:
  - Level 1 (Foundation): Fundamental practices to mitigate typical threats, such as phishing.
  - Level 2 (Attack Methodologies): Advanced risks such as deepfake attacks and automated exploitation discovery.
  - Level 3 (Advanced Threats): Dealing with advanced attack methods such as ransomware, AI-driven malware, and data poisoning.
  - Level 4 (Advanced Techniques) – The most advanced AI techniques are designed to address complex threats, such as zero-day attacks and evasion techniques.
- Each PA (Process Area) at these levels contains multiple P (practices) such as P1, P2, etc. Each practice is assessed based on the extent of implementation and the effectiveness of the generative AI technology.
- **Level 1: Foundation**
  - PA-1: Phishing Attacks: An Important aspect of the proposed model is the AI-driven approach to combatting phishing attacks. They range from the Advanced use of AI-powered phishing detection solutions and

| ID                                                                                                                                                                                  | Process Areas (PA)                                                       | Cybersecurity Mitigation Levels |                   |                 |              |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------|---------------------------------|-------------------|-----------------|--------------|
| Level 1                                                                                                                                                                             | Foundation                                                               | Novice (0)                      | Comprehension (1) | Development (2) | Advanced (3) |
| PA-1: Automated Phishing Attacks                                                                                                                                                    | Generative AI Practices for Addressing Automated Phishing Attacks        |                                 |                   |                 |              |
| P1                                                                                                                                                                                  | AI-Powered Phishing Detection Systems                                    |                                 |                   |                 | 3            |
| P2                                                                                                                                                                                  | Natural Language Processing (NLP) for Email Analysis                     |                                 |                   | 2               |              |
| P3                                                                                                                                                                                  | Automated URL Analysis and Link Scanning                                 |                                 |                   |                 | 3            |
| P4                                                                                                                                                                                  | AI-Based User Behavior Analysis                                          |                                 | 1                 |                 |              |
| P5                                                                                                                                                                                  | Phishing Simulation and Training                                         | 0                               |                   |                 |              |
| P6                                                                                                                                                                                  | AI-Based Email Classification                                            |                                 |                   |                 | 3            |
| P7                                                                                                                                                                                  | Multi-Factor Authentication (MFA) Enforcement                            |                                 |                   | 2               |              |
| P8                                                                                                                                                                                  | Automated Blacklist/Whitelist Maintenance                                |                                 |                   |                 | 3            |
| P9                                                                                                                                                                                  | AI-Powered Content Filtering                                             |                                 | 1                 |                 |              |
| Results of the Evaluation of “Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs” Level-1 (PA-1) Practices implemented by the Company |                                                                          | Score                           |                   | 3               |              |
|                                                                                                                                                                                     |                                                                          | Mitigation Level                |                   | Advanced        |              |
| Level 2                                                                                                                                                                             | Attack Methodologies                                                     |                                 |                   |                 |              |
| PA-1: Deepfake Attacks                                                                                                                                                              | Generative AI Practices for Addressing Deepfake Attacks                  |                                 |                   |                 |              |
| P1                                                                                                                                                                                  | Deepfake Detection Algorithms                                            |                                 |                   | 2               |              |
| P2                                                                                                                                                                                  | AI-Powered Content Validation                                            |                                 |                   |                 | 3            |
| P3                                                                                                                                                                                  | AI-Based Watermarking and Metadata Embedding                             |                                 | 1                 |                 |              |
| P4                                                                                                                                                                                  | Real-Time Monitoring of Media                                            |                                 |                   | 2               |              |
| P5                                                                                                                                                                                  | AI-Powered Secure Communication Channels                                 | 0                               |                   |                 |              |
| P6                                                                                                                                                                                  | Employee Awareness and Training Programs                                 |                                 |                   | 2               |              |
| P7                                                                                                                                                                                  | Multi-Factor Authentication (MFA)                                        |                                 |                   |                 | 3            |
| P8                                                                                                                                                                                  | AI-Driven Data Anomalies Detection                                       |                                 |                   |                 | 3            |
| P9                                                                                                                                                                                  | AI-Enhanced Video and Image Integrity Checking                           |                                 |                   | 2               |              |
| P10                                                                                                                                                                                 | Collaboration with AI Security Solutions Providers                       |                                 | 1                 |                 |              |
| Results of the Evaluation of “Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs” Level-1 (PA-1) Practices implemented by the Company |                                                                          | Score                           |                   | 2               |              |
|                                                                                                                                                                                     |                                                                          | Mitigation Level                |                   | Development     |              |
| Level 2                                                                                                                                                                             | Attack Methodologies                                                     |                                 |                   |                 |              |
| PA-2: Automated Vulnerability Discovery                                                                                                                                             | Generative AI Practices for Addressing Automated Vulnerability Discovery |                                 |                   |                 |              |
| P1                                                                                                                                                                                  | AI-Powered Static Analysis                                               |                                 |                   | 2               |              |
| P2                                                                                                                                                                                  | AI-Based Fuzz Testing                                                    |                                 | 1                 |                 |              |
| P3                                                                                                                                                                                  | Automated Code Review and Vulnerability Detection                        |                                 |                   |                 | 3            |
| P4                                                                                                                                                                                  | Machine Learning for Anomaly Detection                                   |                                 |                   |                 | 3            |
| P5                                                                                                                                                                                  | AI-Based Secure Coding Practices                                         |                                 |                   |                 | 3            |
| P6                                                                                                                                                                                  | AI-Driven Threat Intelligence                                            |                                 |                   | 2               |              |
| P7                                                                                                                                                                                  | AI-Assisted Penetration Testing                                          |                                 |                   |                 | 3            |
| P8                                                                                                                                                                                  | AI-Powered Patch Management                                              |                                 | 1                 |                 |              |
| P9                                                                                                                                                                                  | Automated Security Audits and Compliance Checking                        | 0                               |                   |                 |              |
| P10                                                                                                                                                                                 | AI for Security Posture Management                                       |                                 |                   |                 | 3            |
| Results of the Evaluation of “Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs” Level-1 (PA-2) Practices implemented by the Company |                                                                          | Score                           |                   | 3               |              |
|                                                                                                                                                                                     |                                                                          | Mitigation Level                |                   | Advanced        |              |
| Level 2                                                                                                                                                                             | Attack Methodologies                                                     |                                 |                   |                 |              |
| PA-3: Ransomware Attacks                                                                                                                                                            | Generative AI Practices for Addressing Ransomware Attacks                |                                 |                   |                 |              |
| P1                                                                                                                                                                                  | AI-Powered Threat Detection and Prevention                               |                                 |                   |                 | 3            |
| P2                                                                                                                                                                                  | Automated Vulnerability Scanning                                         |                                 |                   |                 | 3            |
| P3                                                                                                                                                                                  | AI-Based Phishing Detection                                              |                                 |                   | 2               |              |
| P4                                                                                                                                                                                  | Behavioral Analytics and User Monitoring                                 |                                 |                   |                 | 3            |
| P5                                                                                                                                                                                  | Predictive Analysis for Threat Intelligence                              | 0                               |                   |                 |              |
| P6                                                                                                                                                                                  | Automated Incident Response                                              |                                 | 1                 |                 |              |
| P7                                                                                                                                                                                  | Data Backup and Recovery Automation                                      |                                 |                   |                 | 3            |
| Continued                                                                                                                                                                           |                                                                          |                                 |                   |                 |              |



| Level 2                                                                                                                                                                             | Attack Methodologies                                                |                  |   |             |   |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------|------------------|---|-------------|---|
| PA-3: Ransomware Attacks                                                                                                                                                            | Generative AI Practices for Addressing Ransomware Attacks           |                  |   |             |   |
| P8                                                                                                                                                                                  | AI-Driven Encryption Monitoring                                     |                  |   | 2           |   |
| P9                                                                                                                                                                                  | Security Training with AI-Generated Simulations                     |                  |   |             | 3 |
| P10                                                                                                                                                                                 | AI-Enabled Endpoint Protection                                      |                  | 1 |             |   |
| P11                                                                                                                                                                                 | AI for File Integrity Monitoring                                    |                  |   | 2           |   |
| P12                                                                                                                                                                                 | Network Traffic Anomaly Detection                                   |                  |   |             | 3 |
| P13                                                                                                                                                                                 | Deep Learning for Malware Analysis                                  |                  |   |             | 3 |
| P14                                                                                                                                                                                 | AI-Assisted Access Control                                          |                  |   | 2           |   |
| P15                                                                                                                                                                                 | AI for Blockchain-Based Solutions                                   |                  | 1 |             |   |
| P16                                                                                                                                                                                 | AI-Driven Deception Technology                                      | 0                |   |             |   |
| Results of the Evaluation of “Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs” Level-1 (PA-3) Practices implemented by the Company |                                                                     | Score            |   | 3           |   |
|                                                                                                                                                                                     |                                                                     | Mitigation Level |   | Advanced    |   |
| Level 2                                                                                                                                                                             | Attack Methodologies                                                |                  |   |             |   |
| PA-4: AI-Driven Social Engineering                                                                                                                                                  | Generative AI Practices for Addressing AI-Driven Social Engineering |                  |   |             |   |
| P1                                                                                                                                                                                  | Email and Communication Filtering Systems                           |                  |   | 2           |   |
| P2                                                                                                                                                                                  | Employee Awareness Training                                         |                  |   |             | 3 |
| P3                                                                                                                                                                                  | Multi-Factor Authentication (MFA)                                   |                  | 1 |             |   |
| P4                                                                                                                                                                                  | Voice Verification Systems                                          |                  |   | 2           |   |
| P5                                                                                                                                                                                  | Strict Communication Protocols                                      |                  |   | 2           |   |
| P6                                                                                                                                                                                  | Employee Education on Social Engineering                            |                  |   |             | 3 |
| P7                                                                                                                                                                                  | Identity Verification Systems                                       | 0                |   |             |   |
| P8                                                                                                                                                                                  | Access Control Policies                                             |                  | 1 |             |   |
| P9                                                                                                                                                                                  | Continuous Monitoring of Transactions                               |                  |   |             | 3 |
| P10                                                                                                                                                                                 | Deepfake Detection Tools                                            |                  |   |             | 3 |
| P11                                                                                                                                                                                 | Verification Procedures for High-Stakes Communication               |                  |   | 2           |   |
| P12                                                                                                                                                                                 | Social Media Monitoring                                             |                  |   |             | 3 |
| P13                                                                                                                                                                                 | Version Control and Code Review Practices                           |                  |   | 2           |   |
| P14                                                                                                                                                                                 | AI Model Auditing and Logging                                       |                  | 1 |             |   |
| P15                                                                                                                                                                                 | Access Control on Development Environments                          |                  |   |             | 3 |
| P16                                                                                                                                                                                 | Personal Data Protection and Encryption                             |                  | 1 |             |   |
| P17                                                                                                                                                                                 | AI-Based Threat Detection Systems                                   |                  |   | 2           |   |
| P18                                                                                                                                                                                 | Security Awareness on Tailored Attacks                              | 0                |   |             |   |
| P19                                                                                                                                                                                 | Limit Personal Information Shared Online                            |                  |   | 2           |   |
| P20                                                                                                                                                                                 | Social Media Monitoring Tools                                       |                  |   |             | 3 |
| P21                                                                                                                                                                                 | Employee Education on Social Media Scams                            |                  |   |             | 3 |
| P22                                                                                                                                                                                 | AI-Powered Intrusion Detection Systems (IDS)                        |                  | 1 |             |   |
| P23                                                                                                                                                                                 | Automated Response Systems                                          |                  |   |             | 3 |
| P24                                                                                                                                                                                 | Behavioral Analysis and Anomaly Detection                           |                  |   | 2           |   |
| Results of the Evaluation of “Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs” Level-2 (PA-4) Practices implemented by the Company |                                                                     | Score            |   | 3           |   |
|                                                                                                                                                                                     |                                                                     | Mitigation Level |   | Advanced    |   |
| Level 3                                                                                                                                                                             | Advanced Threats                                                    |                  |   |             |   |
| PA-1: AI-Powered Malware                                                                                                                                                            | Generative AI Practices for Addressing AI-Powered Malware           |                  |   |             |   |
| P1                                                                                                                                                                                  | AI-Based Malware Detection System                                   |                  |   |             | 3 |
| P2                                                                                                                                                                                  | Behavioral Analysis and Anomaly Detection                           |                  |   | 2           |   |
| P3                                                                                                                                                                                  | Automated Code Review and Auditing                                  |                  | 1 |             |   |
| P4                                                                                                                                                                                  | Phishing and Social Engineering Attack Prevention                   |                  |   | 2           |   |
| P5                                                                                                                                                                                  | AI-Enhanced Endpoint Security                                       | 0                | 1 |             |   |
| P6                                                                                                                                                                                  | Automated Incident Response Systems                                 |                  |   |             | 3 |
| P7                                                                                                                                                                                  | Deep Learning for Predictive Threat Intelligence                    |                  |   | 2           |   |
| P8                                                                                                                                                                                  | Generative Adversarial Networks (GANS) for Malware Analysis         |                  |   | 2           |   |
| P9                                                                                                                                                                                  | AI-driven Threat Hunting                                            |                  |   |             | 3 |
| P10                                                                                                                                                                                 | Security Patch Management with AI                                   |                  | 1 |             |   |
| Results of the Evaluation of “Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs” Level-3 (PA-1) Practices implemented by the Company |                                                                     | Score            |   | 2           |   |
|                                                                                                                                                                                     |                                                                     | Mitigation Level |   | Development |   |

|                                                                                                                                                                                     |                                                                                   |                  |   |               |   |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------|------------------|---|---------------|---|
| Level 3                                                                                                                                                                             | Advanced Threats                                                                  |                  |   |               |   |
| PA-2: Data Poisoning and Adversarial Attacks                                                                                                                                        | Generative AI Practices for Addressing Data Poisoning and Adversarial Attacks     |                  |   |               |   |
| P1                                                                                                                                                                                  | Data Validation and Preprocessing                                                 |                  |   |               | 3 |
| P2                                                                                                                                                                                  | Robust Model Training                                                             |                  |   |               | 3 |
| P3                                                                                                                                                                                  | Anomaly Detection                                                                 |                  | 1 |               |   |
| P4                                                                                                                                                                                  | Input Sanitization                                                                |                  |   | 2             |   |
| P5                                                                                                                                                                                  | Model Explainability and Transparency                                             |                  |   |               | 3 |
| P6                                                                                                                                                                                  | Data Augmentation                                                                 |                  | 1 |               |   |
| P7                                                                                                                                                                                  | Secure Data Collection                                                            | 0                |   |               |   |
| P8                                                                                                                                                                                  | Regular Model Auditing                                                            |                  |   | 2             |   |
| P9                                                                                                                                                                                  | Differential Privacy                                                              |                  |   |               | 3 |
| P10                                                                                                                                                                                 | Adversarial Robustness Evaluation                                                 |                  |   | 2             |   |
| P11                                                                                                                                                                                 | Model Ensembling                                                                  |                  |   |               | 3 |
| P12                                                                                                                                                                                 | Redundancy in Data Sources                                                        |                  | 1 |               |   |
| P13                                                                                                                                                                                 | Automated Poisoning Attack Detection                                              |                  |   |               | 3 |
| Results of the Evaluation of “Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs” Level-3 (PA-2) Practices implemented by the Company |                                                                                   | Score            |   | 3             |   |
|                                                                                                                                                                                     |                                                                                   | Mitigation Level |   | Advanced      |   |
| Level 3                                                                                                                                                                             | Advanced Threats                                                                  |                  |   |               |   |
| PA-3: Supply Chain Attacks                                                                                                                                                          | Generative AI Practices for Addressing Supply Chain Attacks                       |                  |   |               |   |
| P1                                                                                                                                                                                  | AI-Powered Code Review and Vulnerability Scanning                                 |                  |   | 3             |   |
| P2                                                                                                                                                                                  | Automated Dependency Management                                                   |                  |   |               | 3 |
| P3                                                                                                                                                                                  | Behavioral Analysis of Third-Party Code                                           |                  |   | 2             |   |
| P4                                                                                                                                                                                  | AI-Driven Continuous Monitoring of Supply Chain                                   |                  |   |               | 3 |
| P5                                                                                                                                                                                  | AI-Based Threat Intelligence Sharing                                              |                  | 1 |               |   |
| P6                                                                                                                                                                                  | AI-Powered Incident Response and Automation                                       | 0                |   |               |   |
| P7                                                                                                                                                                                  | Generative AI for Secure Software Design                                          |                  |   |               | 3 |
| P8                                                                                                                                                                                  | AI-Enhanced Risk Prediction and Proactive Threat Modeling                         |                  |   | 2             |   |
| P9                                                                                                                                                                                  | Secure Code Generation via AI                                                     |                  |   |               | 3 |
| P10                                                                                                                                                                                 | Decentralized Code Signing and AI-Powered Authentication                          |                  | 1 |               |   |
| P11                                                                                                                                                                                 | Supply Chain Risk Assessment Models Using AI                                      |                  |   |               | 3 |
| Results of the Evaluation of “Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs” Level-3 (PA-3) Practices implemented by the Company |                                                                                   | Score            |   | 3             |   |
|                                                                                                                                                                                     |                                                                                   | Mitigation Level |   | Advanced      |   |
| Level 4                                                                                                                                                                             | Cutting Edge Techniques                                                           |                  |   |               |   |
| PA-1: AI-Generated Exploits and Zero-day Attacks                                                                                                                                    | Generative AI Practices for Addressing AI-Generated Exploits and Zero-day Attacks |                  |   |               |   |
| P1                                                                                                                                                                                  | Threat Modeling with AI                                                           |                  |   |               | 3 |
| P2                                                                                                                                                                                  | Automated Vulnerability Scanning                                                  |                  |   | 2             |   |
| P3                                                                                                                                                                                  | AI-Enhanced Anomaly                                                               |                  | 1 |               |   |
| P4                                                                                                                                                                                  | AI-Powered Intrusion Detection Systems (IDS)                                      |                  |   | 2             |   |
| P5                                                                                                                                                                                  | AI-Assisted Penetration Testing                                                   |                  |   |               | 3 |
| P6                                                                                                                                                                                  | Machine Learning for Malware Detection                                            |                  | 1 |               |   |
| P7                                                                                                                                                                                  | AI for Patch Management                                                           | 0                |   |               |   |
| P8                                                                                                                                                                                  | AI-Driven Code Review and Refactoring                                             |                  |   |               | 3 |
| P9                                                                                                                                                                                  | AI-Powered Threat Intelligence                                                    |                  | 1 |               |   |
| P10                                                                                                                                                                                 | Behavioral Analysis of AI Systems                                                 |                  |   | 2             |   |
| P11                                                                                                                                                                                 | AI-Enhanced Incident Response                                                     |                  | 1 |               |   |
| P12                                                                                                                                                                                 | Generative Adversarial Network (GAN) Detection                                    | 0                |   |               |   |
| P13                                                                                                                                                                                 | AI-Enhanced Security Logging and Monitoring                                       |                  | 1 |               |   |
| Results of the Evaluation of “Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs” Level-4 (PA-1) Practices implemented by the Company |                                                                                   | Score            |   | 1             |   |
|                                                                                                                                                                                     |                                                                                   | Mitigation Level |   | Comprehension |   |
| Level 4                                                                                                                                                                             | Managed Vulnerabilities and AI Risks                                              |                  |   |               |   |
| PA-2: AI-Driven Evasion Techniques                                                                                                                                                  | Generative AI Practices for Addressing AI-Driven Evasion Techniques               |                  |   |               |   |
| Continued                                                                                                                                                                           |                                                                                   |                  |   |               |   |

| Level 4                                                                                                                                                                             | Managed Vulnerabilities and AI Risks                                |                  |   |             |   |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------|------------------|---|-------------|---|
| PA-2: AI-Driven Evasion Techniques                                                                                                                                                  | Generative AI Practices for Addressing AI-Driven Evasion Techniques |                  |   |             |   |
| P1                                                                                                                                                                                  | AI Model Robustness Enhancement                                     |                  |   | 2           |   |
| P2                                                                                                                                                                                  | Anomaly Detection                                                   |                  | 1 |             |   |
| P3                                                                                                                                                                                  | Explainable AI (XAI)                                                |                  |   |             | 3 |
| P4                                                                                                                                                                                  | Adversarial Training                                                |                  |   | 2           |   |
| P5                                                                                                                                                                                  | Regular Model Evaluation and Testing                                |                  |   |             | 3 |
| P6                                                                                                                                                                                  | Robustness Metrics Implementation                                   |                  |   | 2           |   |
| P7                                                                                                                                                                                  | Data Augmentation and Diversification                               |                  |   |             | 3 |
| P8                                                                                                                                                                                  | Continuous Learning Models                                          | 0                |   |             |   |
| P9                                                                                                                                                                                  | Secure Data Collection and Labelling Practices                      |                  |   | 2           |   |
| P10                                                                                                                                                                                 | Post-Deployment Monitoring                                          |                  | 1 |             |   |
| Results of the Evaluation of “Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs” Level-4 (PA-2) Practices implemented by the Company |                                                                     | Score            |   | 2           |   |
|                                                                                                                                                                                     |                                                                     | Mitigation Level |   | Development |   |

**Table 11.** Example of a Case Study for Evaluation of the Generative AI-Driven Model for Mitigating Cybersecurity Threats in Software Development for SMEs.

AI-based email classification to the Novice use of phishing simulation and training. A 3 rating (Advanced) means the organization has taken the most advanced actions against phishing attacks.

#### • Level 2: Attack Methodologies

- o PA-1 Deepfake Attacks: Generative AI practices respond to the threat of deepfake media. This includes the company's use of Advanced, AI-powered content validation, Development of deepfake detection algorithms, and Development of employee awareness programs. A value of 2 warns that best practices are maturing, but have not yet been perfected.
- o PA-2: Automated Vulnerability Discovery: This PA assesses practices such as automated code review and vulnerability discovery (Advanced) and AI-assisted penetration testing (Advanced) at the level of mitigation 3 (Advanced).
- o PA-3: Ransomware Attacks: The vendor employs AI-driven threat detection and prevention (Advanced), automated vulnerability scanning (Advanced), and behavioral analytics (Advanced). We rated this at an advanced level, with a score of 3, due to a well-implemented AI practice for ransomware mitigation.

#### • Level 3: Advanced Threats

- o PA-1: AI-Powered Malware: As exemplified with AI-based malware detection systems (Advanced) and automated incident response systems (Advanced). The practices are sound evidence, but the average score is 2 (Developing) at this level.
- o PA-2: Data Poisoning and Adversarial Attacks: Stabilize such activities as data validation and preprocessing (Advanced), robust model training (Advanced), and adversarial robustness testing (Advanced). There is also better mechanisms against adversarial attacks at the company, having a score of 3 (Advanced).
- o PA-3: Supply Chain Attacks: The organization applies AI to its processes, including AI-based code inspection and vulnerability scanning (Advanced), automated dependency management (Advanced), and AI-assisted risk prediction (Advanced) with a score of 3 (Advanced).

#### • Level 4: Cutting Edge Techniques

- o PA-1: AI-Generated Exploits and Zero-Day Attacks: The organization has some AI-based practices in the code review (Advanced) and penetration (Advanced) fields, but there is a void in the field of AI-enabled incident response (Comprehension). The total mitigation score of this area stands at 1 (Comprehension), implying that we have covered some advanced capabilities, but we still have efforts to cover to make the company more solid against advanced threats, including zero-days.
- o PA-2: AI-Driven Evasion Techniques: This consists of approaches and practices such as improving AI model robustness, which is a Development task, and explainable AI, which is an advanced task. Overall, this process area is currently at level 2 (Developed), showing some progress, but much more improvement is needed.
- Each process area was rated to determine the overall risk mitigation level for generative AI-driven cybersecurity practices. The final results show:
  - o Practices at Level 1 are largely Advanced, indicating strong application of core practices.
  - o Level 2 practices are a composite of Development and Advanced, and a few are emerging further.
  - o Most over AI-driven threats, for example, ransomware and malware, Level 3 practices are Advanced.

- o Level 4 practices (e.g., responding to zero-day attacks) are similar to understood, but not as mature.

The case study describes a high degree of instantiation of generative AI practices in using software development to address cyber threats for SMEs. The assessment revealed that while the company has successfully leveraged AI to protect against foundational and sophisticated cyber threats, there are still areas for growth, particularly in next-generation technologies. Figure 12 presents the overall evaluation of the proposed model in the case study.

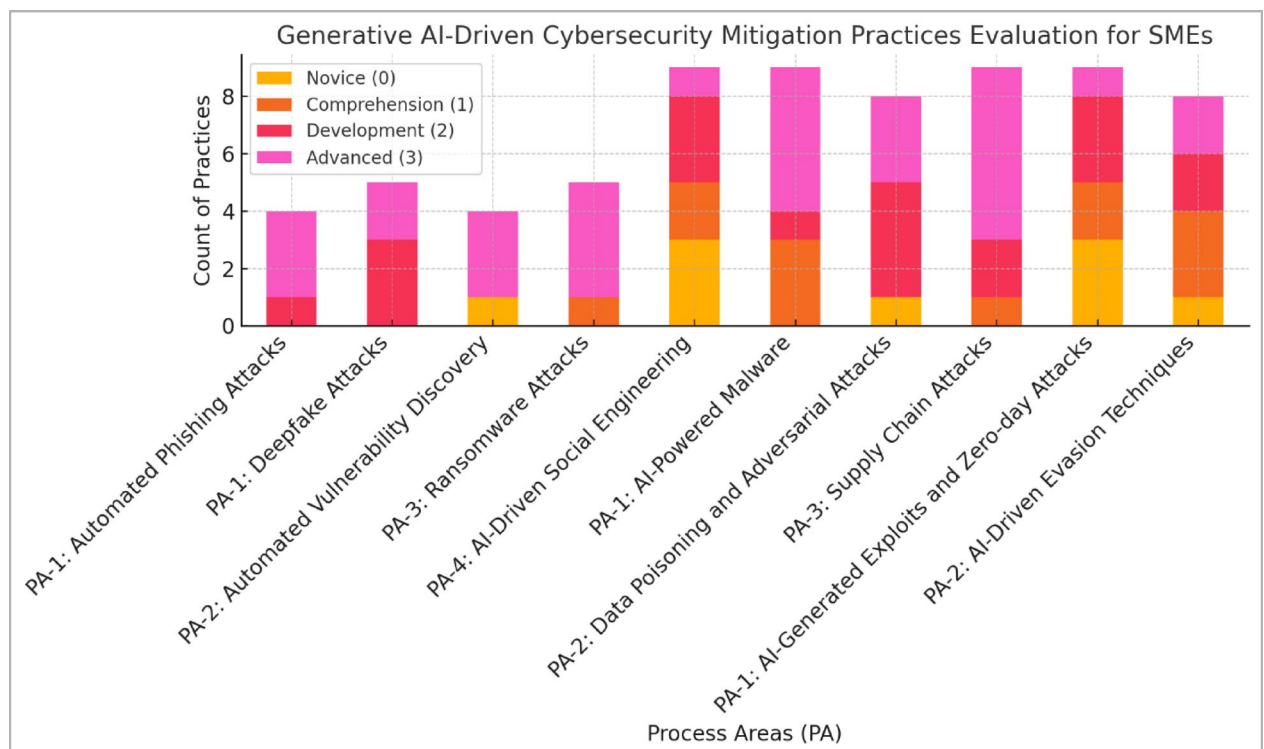
### Framework/mode scenario

#### Scenario 1: data poisoning and adversarial manipulation of LLM inputs

- Threat: External users or developers can implicitly manipulate input to a Large Language Model (LLM) and obtain insecure code.
- Path in ISM:
  - o **Data Poisoning and Adversarial Attack:** The attacker feeds the AI model with modified inputs, e.g., a malformed function description or misleading comments, resulting in the model generating unsafe code. This may result in such vulnerabilities as poor authentication, improper input validation, or poor error handling in the generated code.
- Prioritized Mitigations:
  - o **Input Checking:** Have a way of intercepting and filtering adversarial input before it gets to the LLM. This involves checking of peculiarities or syntax that might result in the generation of insecure code.
  - o **Contextual Analysis:** Simulate methods of contextual analysis to ensure that the code being generated is in conformance with the principles of secure coding and is not engaging in any risky activity, such as the mishandling of sensitive information.
  - o **Model Guardrails:** Add guardrails to the model to guarantee that the model does not generate insecure code, i.e. one that does not conform to the best security practices (e.g. unencrypted sensitive data processing).

#### Scenario 2: code injection via README files or comments

- Threat: README files or code repository comments are likely to contain malicious code.
- Path in ISM:



**Fig. 12.** Case Study Evaluation of the Model.

| Generative-AI Model                                               | Typical use-case(s) in Cybersecurity                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   | Role in Mitigation / Relevance for SMEs and Software Development                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |
|-------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Generative Adversarial Network (GAN)                              | <ul style="list-style-type: none"><li>• Synthetic data generation (network traffic, attack logs, malware traces)</li><li>• Adversarial training / augmentation for intrusion detection / malware detection</li><li>• Simulating varied and rare / zero-day attacks</li><li>• Data-balancing for ML-based IDS / anomaly detectors</li></ul>                                                                                                                                                                                             | <ul style="list-style-type: none"><li>- Synthetic attack / traffic data generation: GANs can produce realistic network-traffic patterns, malware traces, or unusual user behavior logs. This enables SMEs — which often lack large, labeled cybersecurity datasets — to build or train detection models even under data scarcity</li><li>- Improved ML-based intrusion / anomaly detection: By augmenting limited datasets (especially for rare / under-represented attack types), GAN-augmented training helps ANN-based or ML-based IDS detect a broader range of threats with higher accuracy</li><li>- Threat simulation / red-teaming / robustness testing: GANs can simulate polymorphic malware or novel attack vectors, helping developers or security teams in SMEs to test software robustness under varying threat scenarios before release</li></ul>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |
| Diffusion Model (Score-based / diffusion-based generative models) | <ul style="list-style-type: none"><li>• Synthetic data generation (tabular, network traffic, attack logs) as an alternative to GANs</li><li>• Creating diverse variations of data / traffic / threat patterns for training or simulation</li><li>• Augmenting datasets for anomaly detection, IDS, and security testing</li></ul>                                                                                                                                                                                                      | <ul style="list-style-type: none"><li>- Augment training datasets for IDS / anomaly detection: Diffusion models can generate synthetic network-traffic data or attack scenarios that help train machine-learning or ANN-based detection systems when real data is sparse or unavailable</li><li>- Threat scenario generation / security testing: they are able to generate a wide and realistic threat scenarios (e.g., anomalous traffic, payload differences) allowing SMEs to perform stress tests, what-if, and vulnerability testing of their software in advance. This is an aid to the secure by design software development</li><li>- Multiple types of threat data: Diffusion models have the ability to generate a wide range of synthetic samples (including rare/new types of attacks), eliminating issues of class imbalance and enhance generalization of detection models</li></ul>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |
| Large Language Model (LLM — e.g., Transformer / GPT-style models) | <ul style="list-style-type: none"><li>• Cyber-threat intelligence (CTI): analyzing, summarizing, and generating threat reports, code analysis, vulnerability descriptions, exploit summarization</li><li>• Simulating phishing, social-engineering, or adversarial social attacks (text-based)</li><li>• Assisting in automated code review, vulnerability detection, and security-aware code recommendations</li><li>• Supporting incident-response via natural-language analysis, log inspection, alert summarization, etc</li></ul> | <p>Phishing / social engineering detection/ simulation: LLMs are capable of analysing email / text messages to identify suspicious language patterns, phishing messages, or impersonation, a widely used attacker method with SMEs having low levels of security personnel. Can also be applied to create authentic phishing simulation email to train employees</p> <p>Threat intelligence &amp; vulnerability evaluation: LLMs have the capability to process existing vulnerability databases, exploit descriptions, security advisories, summarizing, classifying and highlighting of applicable threats to SMEs. This assists the small dev teams to be aware of the pertinent vulnerabilities and patches</p> <p>Security-conscious code inspection / secure by design support: As part of software development lifecycle (SDLC), LLMs can be used to help developers identify insecure code patterns, or propose secure variants of the code—reducing the risk of adding a vulnerability in the development process (particularly where the security knowledge of an SME user is limited)</p> <p>Automated incident-response, alert triage, and documentation: LLMs can assist in parsing logs, alerts, and generating human readable summaries or remediation recommendations—decreasing the amount of small teams have to do, and reducing response time</p> <p>- Automated incident-response, alert triage, and documentation: LLMs can help parse logs, alerts, and produce human-readable summaries or remediation suggestions — reducing overhead on small teams and enabling faster response times</p> |

Table 12. Generative AI-Driven Models for Mitigating Cybersecurity Threats.

- o **Code Injection:** README files or comments can serve as sources, allowing attackers to inject malicious instructions or code that changes the functionality of the code generation process. This manipulation might lead to vulnerable code, such as a backdoor or data leaks.
- **Prioritized Mitigations:**
  - o **Policy Filters:** Policy filters can be used to filter external inputs, such as readme files or comments, to remove any harmful content.
  - o **Context Isolation:** The code generation process should be isolated from any comments or documentation, so that only valid input code can affect the generated code and not be manipulated by possibly compromised metadata.
  - o **Static Application Security Testing (SAST):** Add to the CI/CD pipeline; use tools of the following type to automatically locate malicious code injections during the generation phase and warn developers before deployment.

These case scenarios depict how the various software development pipeline threats can be mapped with the ISM framework. This will enable them to identify potential pathways for risk spread and implement targeted mitigations to reduce the likelihood of security violations.

Generative AI-driven models for mitigating cybersecurity threats

Table 12 summarizing Generative AI-driven models (e.g., GANs, Diffusion models, LLMs) and how they can be used for mitigating cybersecurity threats in software development / SME contexts.

Implications of the study

The study proposes a generative AI-based model to address security threats in SMEs’ software development, combining ANNs and ISM. The practical implications of the theoretical impacts of the research have many practical advantages to SMEs and intellectual results to the academic literature of cybersecurity, AI, and software development.

- **Practical/managerial implications:** The key value of the research is that it has led to a substantial advance in the cybersecurity practices of SMEs that are frequently rejected in terms of the discussion of innovative security technologies because they have limited resources. With the application of generative AI strategies, SMEs can reinforce their automated (AI-based) defenses: detect fake phishing attacks, ransomware, and stop deepfakes. This does not only minimize chances of being exposed to data breach, but the resilience of the SMEs will also be enhanced in an ever-increasing threat environment.
- **Scalability and Efficiency:** SMEs have limited resources at all times and therefore they find it hard to come up with effective cybersecurity solutions. The proposed model is an Artificial Intelligence-based generative model that could be used as a promising scaling system to be applied to organizations of different sizes and levels

of complexity. Including ANN-ISM techniques, the model allows SMEs to automate major security functions. The model enables even a small business to be able to afford to install sophisticated security technology without incurring high costs in terms of the initial capital necessary to buy specialized equipment and staff.

- **The Role of AI in Threat Detection and Mitigation:** The vast majority of historical definitions of cybersecurity are reactive as they are based on familiar vulnerabilities and threat patterns, which have already been identified. The implementation of generative AI allows preventing threats and responding in advance. The forecasting and prevention of possible security breach will be done correctly by ANN models which will enable the SMEs to forecast and handle the threats before they interfere with normal operations. These security forecast capability will enable the SMEs to be a step further than the cybercriminals and safeguard themselves better.
- **AI Adoption Framework for SMEs:** The lack of skill set and resources necessary to implement AI-based solutions to success is one of the main obstacles to SME sector. The presented paper is a systematic discussion of popularizing AI in cybersecurity, targeting small businesses. The ANN-ISM approaches will give us a framework that will be used to furnish our model with a systematic rule of thumb in utilizing AI in cybersecurity to the SMEs. According to this research study, the use of AI has been adopted, allowing SMEs to access the new technologies which will democratize access to advanced hackers.
- **Policy and Industry Implications:** The implications of the research study are to policy makers and the industry stakeholders. The necessity of helping those units curb cybercrime is on the rise as a result of the importance of SMEs in the world economy. In this study, it is found that the industry standards and regulations have to start embracing new AI-based cybersecurity practices and play a role in enhancing the use of generative AI tools and practices in SMEs. Such models could have a greater effect with government-based programs to raise AI awareness and infrastructure to assist SMEs.
- **Academic Implications:** The study also contributes to the existing literature on AI, cybersecurity and the neural network as a security solution. The ANN and ISM combination is suggested as a novel method of research and application of cybersecurity in SMEs, which opens the new research opportunities in the sphere. Generative AI usage in the sphere of software development and threat mitigation is a comparatively new concept, and the current paper is only the first step into studying the role of AI in the realm of cybersecurity.

In general, the contributions of such a study are immeasurable, as they will introduce academics and practitioners, on the one hand, to SMEs with scalable and efficient solutions to cybersecurity attacks, and, on the other hand, to the following scholarly community on AI technology and its problems. This model has a significant potential to enhance the security of SMEs in that they can satisfy the needs and challenges unique to their businesses and enhance their capacity to fight the escalating attacks.

## Research limitations

While the research provides some valuable guidelines on the potential use of generative AI-empowered models for sector-specific cyber-threats reduction within the SMEs (applying the ANN-ISM procedures in the software development environment), it also has several limitations that should be acknowledged:

- **Small number of case studies:** The work relies heavily on a single case study to validate the efficacy of the generative AI-powered model. These findings may not be relevant to all SMEs, for organizations may have different cybersecurity maturity levels, resources, and technological platforms. Research on a more diverse group of cases would provide a better picture of how the model works in diverse populations.
- **AI Ability Expectation:** The paper presupposes that AI generated by ANN and other AI based on ANN can be used to address the broad spectrum of cybersecurity issues. Nevertheless, this performance of AI in the domain might vary in practice, according to the nature of cyber-attacks, the quality of data, and the constant evolving cyber threats the researcher stated. This is one of the limitations that one should take into account when applying the model to other domains.
- **Data dependency and quality:** The quality of data and its amount that is used to train the AI-based models is the main factor that affects the performance of the models. The cybersecurity data in SMEs may be insufficient or inaccurate, and it may have an impact on the efficiency of generative AI systems. The analysis does not take into account the possibility of SMEs that face problems with the collection, cleaning, and organization of the data that is required to train AI.
- **Absence of Long-Term Analysis:** The study is based on building and primary testing of the model upon the introduction; there is no need to analyze long-term opportunities. But, the effectiveness of the AI model in the long term is not known. Cyber-attacks are dynamic and it is unclear whether the adversarial model would be satisfactory as new vectors and techniques of attack arise.
- **Computational and Resource Constraints:** The usage of generative AI and ANN-based approaches to cybersecurity can be associated with significant computational needs. Despite the beneficial impact of high-level production technologies, discussed paper also mentions that SMEs can experience budgetary limitations, lack of knowledge, and resources, in the introduction of such a solution. These issues could limit the popularity of the proposed model in SMEs that are resource-constrained.
- **AI Methods Generalization:** There is a possibility that the proposed ANN-ISM methods did not specifically serve all the cybersecurity risks or sectors. In particular, we could be required to make the AI model better able to identify particular attack styles, including zero-day exploits or APTs. This can restrict the application of the findings to the context of threats.
- **Complexity/Implementation:** The paper describes the idea of generative AI-based model application and gives a threat mitigation case study. However, it does not elaborate on the problems that confront organizations (or SMEs in this case) in adopting such technologies. Such obstacles can be integration of AI systems



with a current system of a company, unwillingness to introduce AI-based systems among employees and the effective utilization of this tools.

- **Ethical and privacy concerns:** There is no discussion in the paper regarding the ethicality of using AI in cybersecurity, false positives, privacy of data, or the necessity of transparency in the work of AI. It is these that need to be taken into account when introducing AI into such sensitive areas as software development.
- **Scalability of Model:** The model offered by the study can be applied to SMEs, however, the scalability of based AI-generative model of all the sizes and complexities of firms is not fully explored. In bigger companies, bigger data will have to be processed, and security issues will demand additional tailoring or scaling of outputs of AI systems.

In spite of these constraints, the study provides a solid foundation for future research on AI-based cybersecurity systems, as it offers insights into how SMEs could leverage advanced AI practices to expand their options and capabilities in cybersecurity. These can be mitigated in future work to enhance the applicability and soundness of the proposed model.

## Conclusion and future research directions

In this paper, an AI-based generative model has been developed to enhance cybersecurity in the software development process of Small and Medium-sized software Enterprises (SMEs) by applying ANN and ISM methods. Cybersecurity issues in this fast-paced high-speed software-developing world are interpreted in a regimented and flexible combination of generative AI, ANN, and ISM.

According to the case, in such ways, SMEs would transform their organizations into more robust defenses against common cybersecurity threats, such as phishing, ransomware, deep fake attacks, and automated vulnerability detection. The analysis of the different arenas of processes showed that it highly adopted state of the art practices implying that generative AI, especially coupled with neural networks offer advanced, scalable and proactive defenses than the current models of cybersecurity.

Other important functions that need to be automated to ensure the small IT staff of SMEs is relieved of their duties include phishing detection and vulnerability scanning, which helps to respond to the attack more quickly and reduce its effectiveness. Moreover, the application of ISM allowed understanding and visualizing dependencies between cybersecurity measures much better, which subsequently made the model more applicable to SMEs with resource limitations. The overall ability of the model to generalize against various attack targets and the subsequent improvements made to the process of mitigation proves the effectiveness of the suggested model in the SME-appropriate environment.

Even though this research has created the foundation of an AI-based implementation that generates solutions to curb cybersecurity attacks against SMEs, future research can be approached differently:

- **Reinforcement Learning for Adaptive Security:** The reinforcement learning (RL) methods can be employed in the future to offer adaptiveness to new cyber threats. It is possible that RL could also allow the system to refine its defense based on simulated ongoing attacks on the play, which supports the learning of the AI and allows it to adapt dynamically more quickly.
- **Broadening AI for Threat Detection Capabilities:** Although the current model applies generative AI in responding to particular threats, it can be expanded to include the widest possible range of threats, such as insider attacks, zero-day threats, and social engineering. Areas for future work include developing hybrid models combining multiple machine learning methods to address a wide range of attack vectors.
- **Model explainability and interpretability:** One drawback of AI models, especially neural networks, is their decision-making transparency. Subsequent work should further emphasise making these models explainable, so that SMEs can understand why specific mitigation measures are suggested or taken, thereby fostering trust and adoption.
- **Real-Time Cooperation of Threat Intelligence:** As a prospective line of research, this paper may investigate the use of generative AI and mock models with threat intelligence sharing systems. By leveraging information with external parties, such as cybersecurity vendors and governments, SMEs can gain access to live threat data and ensure that the AI-driven model identifies threats promptly and accurately.
- **Scalability and Customization for SMEs:** Given that SMEs can be highly heterogeneous in their resource endowment and industry-specific threat landscape, future work should consider more customizable versions of the generative AI-driven model. Scalable, module-based, industry-specific cybersecurity implementations would enable SMEs to select and customize the most suitable AI-informed practices for their organizations and risk exposure.
- **Validation in Different SME Contexts:** Figures in other places and locations will be employed to prove the usefulness of the model and its applicability. The research will ensure that the model gets updated to fit the requirements of certain settings as the model is tested on different organizations, without compromising its generalizability to SMEs across the globe.
- **Human-AI Collaboration for Improved Decision Making:** The future of cybersecurity may manifest itself in the form of a collaboration of human specialists and artificial intelligence. The next step would be to combine human intelligence with artificial intelligence so that it can analyze danger quickly. Simultaneously, the final decision is made by human professionals, whilst the decision-making process remains to AI suggestions.

The generative AI-enabled cybersecurity threat defense model can be enhanced by taking into consideration both the strengths and the shortcomings of the above research, which is critical to allowing the SMEs to build a robust cybersecurity threat defense mechanism. The constant creation of new technologies and cybersecurity policies will be instrumental in ensuring the protection of SMEs against resilient cyber attacks.

## Data availability

All data generated or analyzed during this study are included in this manuscript. If someone wants to request the data from this study, kindly contact “Rafiq Ahmad Khan” ( rafiqahmadk@gmail.com )(mailto: rafiqahmadk@gmail.com) ).

Received: 26 July 2025; Accepted: 23 January 2026

Published online: 07 February 2026

## References

- Almeida, F. Comparative analysis of EU-based cybersecurity skills frameworks. *Comput. Secur.* **151**, 104329 (2025).
- Birithriya, S. K., Ahlawat, P. & Jain, A. K. Detection and prevention of spear phishing attacks: A comprehensive survey. *Comput. & Security*, **151**, 104317 (2025).
- Yan, P. & Talaei Khoei, T. Securing the internet of things: A comprehensive review of ransomware attacks, detection, countermeasures, and future prospects. *Franklin Open*. **11**, 100256 (2025).
- Chanda, R. C., Vafaei-Zadeh, A., Hanifah, H. & Nikbin, D. Assessing cybersecurity awareness among bank employees: A multi-stage analytical approach using PLS-SEM, ANN, and fsQCA in a developing country context. *Comput. & Security*, **149**, 104208 (2025).
- Waelchli, S. & Walter, Y. Reducing the risk of social engineering attacks using SOAR measures in a real world environment: A case study. *Comput. & Security*, **148**, 104137 (2025).
- Ali, J. et al., A deep dive into cybersecurity solutions for AI-driven IoT-enabled smart cities in advanced communication networks. *Comput. Commun.* **229**, 108000 (2025).
- Al Aamer, A. K. & Hamdan, A. Cyber security awareness and smes’ profitability and continuity: literature review, in *Emerging Trends and Innovation in Business and Finance*, (eds Khoury, R. E. & Nasrallah, N.) ed Singapore: Springer Nature Singapore, 593–604. (2023).
- Chaudhary, S., Gkioulos, V. & Katsikas, S. A quest for research and knowledge gaps in cybersecurity awareness for small and medium-sized enterprises. *Comput. Sci. Rev.* **50**, 100592, (2023).
- Nadella, G. S. et al., Generative AI-Enhanced Cybersecurity Framework for Enterprise Data Privacy Management. *Computers*. **14**, 55 (2025).
- Khan, H. U. et al. AI-driven cybersecurity framework for software development based on the ANN-ISM paradigm. *Sci. Rep.* **15**, 13423 (2025).
- Coppolino, L., D’Antonio, S., Mazzeo, G. & Uccello, F. The good, the bad, and the algorithm: The impact of generative AI on cybersecurity. *Neurocomputing*. **623**, 129406 (2025).
- Khan, R. A., Khan, H. U., Alwageed, H. S., Hashimi, H. A. & Keshta, I. 5G networks security mitigation model: an ANN-ISM hybrid approach. *IEEE Open. J. Commun. Soc.* **6**, 881–925 (2025).
- Alzahrani, A. & Khan, R. A. Secure software design evaluation and decision making model for ubiquitous computing: A two-stage ANN-Fuzzy AHP approach. *Comput. Hum. Behavior*. 108109 (2023).
- Hathal, S., Alwageed & Ghani, A. Ismail. Keshta, Rafiq. Ahmad.Khan, Abdulrahman. Alzahrani, Muhammad. Usman. Tariq, and An empirical study for mitigating sustainable cloud computing challenges using ISM-ANN. *PLOS ONE*. **19**, 1–34, (2024).
- Wong, L. W., Lee, V. H., Tan, G. W. H., Ooi, K. B. & Sohal, A. The role of cybersecurity and policy awareness in shifting employee compliance attitudes: Building supply chain capabilities. *Int. J. Inf. Manag.* **66**, (2022).
- Caniglia, A., Dentamaro, V., Galantucci, S. & Impedovo, D. FOBICS: Assessing project security level through a metrics framework that evaluates DevSecOps performance. *Inf. Software Technol.* **178**, 107605, (2025).
- Rawindaran, N. et al., Enhancing Cyber Security Governance and Policy for SMEs in Industry 5.0: A Comparative Study between Saudi Arabia and the United Kingdom, *Digital*. **3**, 200–231 (2023).
- Kaur, R., Klobučar, T. & Gabrijelčić, D. Harnessing the power of language models in cybersecurity: A comprehensive review. *Int. J. Inf. Manag. Data Insights*, **5**, 100315 (2025).
- Olusanya, O. O., Jimoh, R. G., Misra, S. & Awotunde, J. B. A neuro-fuzzy security risk assessment system for software development life cycle. *Heliyon*, **10**, e33495 (2024).
- Dutta, I. K., Ghosh, B., Carlson, A., Totaro, M. & Bayoumi, M. Generative Adversarial Networks in Security: A Survey, in *11th IEEE Annual Ubiquitous Computing, Electronics & Mobile Communication Conference (UEMCON)*, 0399–0405 (2020).
- ul Haq, F., Suki, N. M., Zaigham, H., Masood, A. & Rajput, A. Exploring AI Adoption and SME Performance in Resource-Constrained Environments: A TOE-RBV Perspective with Mediation and Moderation Effects. *J. Digital Econ.*, (2025).
- Olatunji, A. P., Alozie, E., Olagunju, H. & Udensi, F. A Systematic Review of the Role of Artificial Intelligence in Cybersecurity, in *IEEE 5th International Conference on Electro-Computing Technologies for Humanity (NIGERCON)*, 1–6. (2024).
- Gupta, A. K., Dixit, N., Kumar, S., Rawat, P. & Madhumita A Novel Hybrid Approach for Threat Detection in Cyber Security using AI algorithm, in *International Conference on Computing, Sciences and Communications (ICCS)*, 1–6. (2024).
- Zavodna, L., Ueberwimmer, M. & Frankus, E. Barriers to the implementation of artificial intelligence in small and medium sized enterprises: pilot study. *J. Econ. Management*, **46**, 331–352 (2024).
- Sharma, O., Sharma, A. & Kalia, A. MIGAN: GAN for facilitating malware image synthesis with improved malware classification on novel dataset. *Exp. Syst. Appl.* **241**, 122678 (2024).
- Sharma, P., Kumar, M., Sharma, H. K. & Biju, S. M. Generative adversarial networks (GANs): Introduction, Taxonomy, Variants, Limitations, and applications. *Multimedia Tools Applications*, 2024.
- Morrison, O. M., Pichi, F. & Hesthaven, J. S. GFN: A graph feedforward network for resolution-invariant reduced operator learning in multifidelity applications. *Comput. Methods Appl. Mech. Eng.* **432** 117458, (2024).
- Levshun, D., Levshun, D. & Kotenko, I. ForecaState: Framework for industrial Internet of Things state forecasting using recurrent neural networks with hyperparameters optimization. *Engi. App. Artificial Intelligence*, **159**, 111627 (2025).
- Khadidos, A. O. et al. CyberSentry: Enhancing SCADA security through advanced deep learning and optimization strategies. *Int. J. Critical Inf. Protection*. **50**, 100782 (2025).
- Agrafiotis, G. et al. A Deep Learning-based Malware Traffic Classifier for 5G Networks Employing Protocol-Agnostic and PCAP-to-Embeddings Techniques, presented at the Proceedings of the 2023 European Interdisciplinary Cybersecurity Conference, Stavanger, Norway, (2023).
- Das, S., Tariq, A., Santos, T., Kantareddy, S. S. & Banerjee, I. Recurrent neural networks (RNNs): Architectures, training Tricks, and introduction to influential research, in *Machine Learning for Brain Disorders*, (ed Colliot, O.) ed New York, NY: Springer US, 117–138. (2023).
- Crisostomo, J., Bacao, F. & Lobo, V. Machine learning methods for detecting smart contracts vulnerabilities within Ethereum blockchain – A review, *Exp. syst. Appl.* **268**, 126353 (2025).
- Udas, P. B., Karim, M. E. & Roy, K. S. SPIDER: A shallow PCA based network intrusion detection system with enhanced recurrent neural networks. *J. King Saud University - Computer Inf. Sci.* **34**, 10246–10272 (2022).

34. Alahmari, A. & Duncan, B. Cybersecurity Risk Management in Small and Medium-Sized Enterprises: A Systematic Review of Recent Evidence, in *2020 International Conference on Cyber Situational Awareness, Data Analytics and Assessment (CyberSA)*, 1–5. (2020).
35. Leotta, M., Ricca, F., Marchetto, A. & Olinas, D. An empirical study to compare three web test automation approaches: NLP-based, programmable, and capture and replay. *J. Software: Evol. Process.* **36**, e2606 (2024).
36. Khan, R. A., Akbar, M. A., Rafi, S., Almagrabi, A. O. & Alzahrani, M. Evaluation of requirement engineering best practices for secure software development in GSD: An ISM analysis. *J. Software: Evolution Process.* vol. n/a, 1–19, (2023).
37. Rajan, R. et al. Developing a modified total interpretive structural model (M-TISM) for organizational strategic cybersecurity management. *Technol. Forecasting Soc. Change.* **170**, 120872. (2021).
38. Etemadi, N., Van Gelder, P. & Strozzi, F. An ISM Modeling of Barriers for Blockchain/Distributed Ledger Technology Adoption in Supply Chains towards Cybersecurity. *Sustainability.* **13**, 4672, (2021).
39. Garousi, V., Felderer, M. & Mäntylä, M. V. Guidelines for including grey literature and conducting multivocal literature reviews in software engineering. *Inf. Software Technol.* **106**, 101–121 (2019).
40. Ito, C. & Ozer, M. Multivocal literature review on zero-trust security implementation, *Comput. & Security.* **141**, 103827 (2024).
41. Akbar, M. A., Smolander, K., Mahmood, S. & Alsanad, A. Toward successful DevSecOps in software development organizations: A decision-making framework. *Inf. Software Technol.* **147**, 106894, (2022).
42. Al-Matouq, H., Mahmood, S., Alshayeb, M. & Niazi, M. A maturity model for secure software design: A multivocal study. *IEEE Access.* **8**, 215758–215776 (2020).
43. Wagner, S. et al., Status Quo in Requirements Engineering: A Theory and a Global Family of Surveys. *ACM Trans. Softw. Eng. Methodol.* **28**, 1–48, (2019).
44. M. Humayun, M. Niazi, M. Assiri, and M. Haoes, Secure Global Software Development: A Practitioners' Perspective, *Appl. Sci.* **13**, 2465 (2023).
45. Ilyas, M., Khan, S. U., Khan, H. U. & Rashid, N. Software integration model: An assessment tool for global software development vendors. *J. Software: Evolution Process.* vol. n/a, e2540, (2023).
46. Creswell, J. W. *Research design: qualitative, quantitative and mixed methods approaches*, 3rd edition: Sage, London, (2009).
47. Lethbridge, T. C., Sim, S. E. & Singer, J. Studying software engineers: data collection techniques for software field studies. *Empirical Softw. Eng.*, **10**, 311–341 (2005).
48. Lee, S. C. Prediction of concrete strength using artificial neural networks. *Eng. Struct.* **25**, 849–857 (2003).
49. Leong, L. Y., Hew, T. S., Tan, G. W. H. & Ooi, K. B. Predicting the determinants of the NFC-enabled mobile credit card acceptance: A neural networks approach. *Exp. Syst. Appl.*, **40**, 5604–5620 (2013).
50. Chan, F. T. & Chong, A. Y. A SEM–neural network approach for Understanding determinants of interorganizational system standard adoption and performances. *Decis. Support Syst.* **54**, 621–630 (2012).
51. Zhang, H. et al. A Framework for Designing Fair Ubiquitous Computing Systems, arXiv preprint arXiv:2308.08710, (2023).
52. Chong, A. Y. L. Predicting m-commerce adoption determinants: A neural network approach. *Exp. Syst. Appl.* **40**, 523–530 (2013).
53. Hertz, J., Krogh, A., Palmer, R. G. & Horner, H. *Introduction To the Theory of Neural Computation* Ed (American Institute of Physics, 1991).
54. Alnaizy, R., Aidan, A., Abachi, N. & Jabbar, N. A. Neural network model identification and advanced control of a membrane biological reactor. *J. Membrane Sep. Technol.* **2**, 231 (2013).
55. Sage, A. P. *Interpretive Structural Modeling: Methodology for Large Scale Systems* 1–445 (McGraw-Hill,, 1977).
56. Ravi, V. & Shankar, R. Analysis of interactions among the barriers of reverse logistics. *Technological Forecasting and Social Change*, **72**, 1011–1029 (2005).
57. Rafi, S., Akbar, M. A., Mahmood, S., Alsanad, A. & Alothaim, A. Selection of devops best test practices: A hybrid approach using ISM and fuzzy TOPSIS analysis. *J. Software: Evol. Process.* **34**, e2448 (2022).
58. Qureshi, K. M. et al. Exploring the lean implementation barriers in small and Medium-Sized enterprises using interpretive structure modeling and interpretive ranking process. *Appl. Syst. Innov.* **5**, 84 (2022).
59. Talib, F., Rahman, Z. & Qureshi, M. R. An interpretive structural modeling approach for modeling the practices of total quality management in service sector. *Int J. Modelling Oper. Management, Inderscience*, **1**, 223–250, (2011).
60. Rafi, S. et al. Exploration of DevOps testing process capabilities: An ISM and fuzzy TOPSIS analysis, *Appl. Soft Comput.*, **116**, 108377 (2022).
61. Sakar, C., Koseoglu, B., Toz, A. C. & Buber, M. Analysing the effects of liquefaction on capsizing through integrating interpretive structural modelling (ISM) and fuzzy Bayesian networks (FBN), *Ocean Engineering*, **215**, 107917 (2020).
62. Patel, M. N., Pujara, A. A., Kant, R. & Malviya, R. K. Assessment of circular economy enablers: Hybrid ISM and fuzzy MICMAC approach. *J. Clean. Production*, **317**, 128387 (2021).
63. Ali, S., Huang, J., Khan, S. U. & Li, H. A framework for modelling structural association amongst barriers to software outsourcing partnership formation: an interpretive structural modelling approach. *J. Software: Evol. Process.* **32**, e2243 (2020).
64. Ali, S. et al. Analyzing the interactions among factors affecting cloud adoption for software testing: a two-stage ISM-ANN approach. *Soft Computing*, **26**, 8047–8075 (2022).
65. Qureshi, K. M. et al. Accomplishing Sustainability in Manufacturing System for Small and Medium-Sized Enterprises (SMEs) through Lean Implementation. *Sustainability*, **14**, 9732 (2022).
66. Qureshi, M. R. & Kumar, P. An integrated model to identify and classify the key criteria and their role in the assessment of 3PL services providers. *Asia Pac. J. Mark. Logistics*, **20**, 227–249 (2008).
67. Qureshi, M. R. & Kumar, P. Modeling the logistics outsourcing relationship variables to enhance shippers' productivity and competitiveness in logistical supply chain. *Int. J. Productivity Perform. Management*, **56**, 689–714 (2007).
68. Azeez, N. A. et al. Adopting automated whitelist approach for detecting phishing attacks, *Comput. & Security*, **108**, 102328 (2021).
69. Ferrag, M. A. et al., Generative AI in Cybersecurity: A Comprehensive Review of LLM Applications and Vulnerabilities, Internet of Things and Cyber-Physical Systems, (2025).
70. Ain, Q. U., Javed, A. & Irtaza, A. DeepEvader: An evasion tool for exposing the vulnerability of deepfake detectors using transferable facial distraction blackbox attack. *Eng. App. Artificial Intell.* **145**, 110276 (2025).
71. Patsakis, C., Casino, F. & Lykousas, N. Assessing LLMs in malicious code deobfuscation of real-world malware campaigns. *Exp. Syst. Appl.*, **256**, 124912 (2024).
72. Austin, A., Holmgreen, C. & Williams, L. A comparison of the efficiency and effectiveness of vulnerability discovery techniques. *Inf. Softw. Technol.* **55**, 1279–1288 (2013).
73. Nobles, C. The weaponization of artificial intelligence in cybersecurity: A systematic review. *Procedia Comput. Sci.* **239**, 547–555 (2024).
74. Iturbe, E., Llorente-Vazquez, O., Rego, A., Rios, E. & Toledo, N. Unleashing offensive artificial intelligence: Automated attack technique code generation. *Comput. & Security.* **147**, 104077 (2024).
75. Andreoli, A., Lounis, A., Debbabi, M. & Hanna, A. On the prevalence of software supply chain attacks: Empirical study and investigative framework. *For. Sci. Int.: Digital Investigation.* **44**, 301508, (2023).
76. Arivudainambi, D., Varun Kumar, K. A., & Visu, P. Malware traffic classification using principal component analysis and artificial neural network for extreme surveillance. *Comput. Commun.* **147**, 50–57 (2019).

77. Al-Subaiey, A., Al-Thani, M., Abdullah Alam, N., Antora, K. F. & Khandakar, A. and S. M. A. Uz Zaman, Novel interpretable and robust web-based AI platform for phishing email detection. *Comput. Electrical Eng.* **120**, 109625 (2024).
78. Lin, Z. et al. Tracking phishing on Ethereum: Transaction network embedding approach for accounts representation learning. *Comput. & Security*, **135**, 103479 (2023).
79. Majgave, A. B. & Gavankar, N. L. Automatic phishing website detection and prevention model using transformer deep belief network. *Comput. & Security*, **147**, 104071 (2024).
80. Nanda, M., Saraswat, M. & Sharma, P. K. Enhancing cybersecurity: A review and comparative analysis of convolutional neural network approaches for detecting URL-based phishing attacks, e-Prime-Advances in Electrical Engineering, Electronics and Energy, **8**, 100533 (2024).
81. Sushma, K. S. N., Viji, C., Rajkumar, N., Ravi, J., Stalin, M., & Najmusher, H. Healthcare 4.0: A Review of Phishing Attacks in Cyber Security. *Procedia Computer Sci.*, **230**, 874–878 (2023).
82. Habbal, A., Ali, M. K. & Abuzaraida, M. A. Artificial Intelligence Trust, Risk and Security Management (AI TRISM): Frameworks, applications, challenges and future research directions. *Exp. Syst. Appl.*, **240**, 122442 (2024).
83. Kumar, A. et al. Advances in DeepFake detection algorithms: Exploring fusion techniques in single and multi-modal approach. *Inf. Fusion*. **118** 102993 (2025).
84. Rabhi, M., Bakiras, S. & Di Pietro, R. Audio-deepfake detection: adversarial attacks and countermeasures. *Exp. Syst. Appl.* **250**, 123941 (2024).
85. Djenna, A., Belaoued, M., Lifa, N. & Moualdi, D. E. Proactive Anti-Ransomware cybersecurity approach. *Procedia Comput. Sci.* **238**, 821–826 (2024).
86. Gaber, M. G., Ahmed, M. & Janicke, H. Malware detection with artificial intelligence: A systematic literature review. *ACM Comput. Surv.* **56**, 148 (2024).
87. Jáněz-Martino, F., Alaiz-Rodríguez, R., González-Castro, V., Fidalgo, E. & Alegre, E. Spam email classification based on cybersecurity potential risk using natural language processing. *Knowledge-Based Syst.*, **310**, 112939 (2025).
88. Ma, K. P., Ryu, D. J. & Lee, S. J. Reverse Analysis Method and Process for Improving Malware Detection Based on XAI Model. *Computers, Materials Continua*, **81**, 4485–4502 (2024).
89. Sirbu, A. G. & Czibula, G. Automatic code generation based on Abstract Syntax-based encoding. Application on malware detection code generation based on MITRE Attack techniques. *Exp. Syst. Appl.*, **264**, 125821 (2025).
90. Vouvoutsis, V., Casino, F. & Patsakis, C. Beyond the sandbox: Leveraging symbolic execution for evasive malware classification. *Comput. & Security*, **149**, 104193 (2025).
91. Wang, M., Zhang, Y. & Wen, W. Improved capsule networks based on Nash equilibrium for malicious code classification. *Comput. & Security*, **136** 103503 (2024).
92. Alghawli, A. S. A. & Radivilova, T. Resilient cloud cluster with DevSecOps security model, automates a data analysis, vulnerability search and risk calculation. *Alexandria Eng. J.*, **107**, 136–149 (2024).
93. Althar, R. R., Samanta, D., Kaur, M., Singh, D. & Lee, H. N. Automated Risk Manage. Based Softw. Secur. Vulnerabilities Manage. *IEEE Access*, **10**, 90597–90608, (2022).
94. Casola, V., De Benedictis, A., Mazzocca, C. & Orbinato, V. Secure software development and testing: A model-based methodology. *Comput. & Security*. **137**, 103639 (2024).
95. del-Hoyo-Gabaldon, J. A., Moreno-Cediel, A., Garcia-Lopez, E., Garcia-Cabot, A. & de-Fitero-Dominguez, D. Automatic dataset generation for automated program repair of bugs and vulnerabilities through SonarQube, *SoftwareX*, **26**, 101664 (2024).
96. Kirilchuk, S., Reutov, V., Nalivaychenko, E., Shevchenko, E. & Yaroshenko, A. Ensuring the security of an automated information system in a regional innovation cluster. *Transp. Res. Procedia*, **63**, 607–617 (2022).
97. Schuckert, F., Katt, B. & Langweg, H. Insecurity Refactoring: Automated Injection of Vulnerabilities in Source Code. *Comput. & Security*. **128**, 103121 (2023).
98. Fotis, F. Cyberattacks: economic impacts and risk management strategies. *Procedia Comput. Sci.* **251**, 672–677 (2024).
99. McIntosh, T. et al., Harnessing GPT-4 for generation of cybersecurity GRC policies: A focus on ransomware attack mitigation. *Comput. & Security*. **134**, 103424 (2023).
100. Burriss, S. K. et al. , Redesigning an AI bill of rights with/for young people: Principles for exploring AI ethics with middle and high school students. *Comput. Educ.: Artificial Intell.* **7**, 100317 (2024).
101. Wu, W., Wang, S., Ding, G. & Mo, J. Elucidating trust-building sources in social shopping: A consumer cognitive and emotional trust perspective. *J. Retailing Consumer Services*. **71**, 103217 (2023).
102. Gershfeld, I. & Sturm, A. Evaluating the effectiveness of a security flaws prevention tool, *Inf. Software Technol.* **170**, 107427 (2024).
103. Zheng, W., Deng, P., Gui, K. & Wu, X. An Abstract Syntax Tree based static fuzzing mutation for vulnerability evolution analysis, *Inf. Software Technol.*, **158**, 107194 (2023).
104. Cai, Z. et al. Generative adversarial networks: A survey toward private and secure applications. *ACM Comput. Surv.* **54**, 132, (2021).
105. Chen, P., Du, X., Lu, Z. & Chai, H. Universal adversarial backdoor attacks to fool vertical federated learning. *Comput. & Security*. **137**, 103601 (2024).
106. Guo, X. Towards Automated Software Testing with Generative Adversarial Networks, in 2021 51st Annual IEEE/IFIP International Conference on Dependable Systems and Networks - Supplemental Volume (DSN-S), 21–22. (2021).
107. Huang, L., Liu, H., Liu, Y., Shang, Y. & Li, Z. A Generative Adversarial Imitation Learning Method for Continuous Integration Testing, in 2024 5th International Seminar on Artificial Intelligence, Networking and Information Technology (AINIT), 1084–1089. (2024).
108. Jati, A. et al. Adversarial attack and defense strategies for deep speaker recognition systems. *Comput. Speech & Language*, **68**, 101199 (2021).
109. Tan, H. et al. Adversarial Attack and Defense Strategies of Speaker Recognition Systems: A Survey. *Electronics*. **11**, 2183, (2022).
110. Charles, V., Emrouznejad, A. & Gherman, T. A critical analysis of the integration of blockchain and artificial intelligence for supply chain. *Annals Oper. Res.* **327**, 7–47 (2023).
111. Rahman, M. S., Khalil, I., Atiquzzaman, M. & Bouras, A. A lightweight practical consensus mechanism for supply chain blockchain. *High-Confidence Computing*, **5**, 100253 (2025).
112. Aloupogianni, E. et al., AI-Driven Optimization of Small Cell Deployment for Beyond 5G Networks. *Procedia Comput. Sci.* **238**, 908–913 (2024).
113. Patel, O. AI-Driven smart contracts. *J. Artif. Intell. & Cloud Comput.*, 1–9, (2024).
114. Sarker, I. H., Furhad, M. H. & Nowrozy, R. Cybersecurity: an Overview, security intelligence modeling and research directions. *SN Comput. Sci.* **2**, 18 (2021).
115. Zacharis, A., Katos, V. & Patsakis, C. Integrating AI-driven threat intelligence and forecasting in the cyber security exercise content generation lifecycle. *Int. J. Inform. Security*, **23**, 2691–2710, (2024).
116. Azeem Akbar, M., Mahmood, S., Alsanad, A. & Com, A. Toward successful devsecops in software development organizations: A Decision-Making framework. *Inf. Softw. Technol.* **147**, 0227 (2022).
117. Kannan, G., Pokharel, S. & Sasi Kumar, P. A hybrid approach using ISM and fuzzy TOPSIS for the selection of reverse logistics provider. *Resources, Conservation Recycling*, **54**, 28–36 (2009).
118. Agarwal, A. & Vrat, P. Modeling attributes of human body organization using ISM and AHP. *Jindal J. Bus. Res.* **6**, 44–62 (2017).



119. Soni, M. End to End Automation on Cloud with Build Pipeline: The Case for DevOps in Insurance Industry, Continuous Integration, Continuous Testing, and Continuous Delivery, (2015).
120. Attri, R., Grover, S., Dev, N. & Kumar, D. Analysis of barriers of total productive maintenance (TPM). *Int. J. Syst. Assur. Eng. Manag.* **4**, 365–377 (2013).
121. Warfield, J. N. Developing interconnection matrices in structural modeling. *IEEE Trans. Syst. Man. Cybernetics.* **SMC-4**, 81–87 (1974).
122. Raluca Brasoveanu, Yusuf Karabulut, and Ivan Pashchenko. 2022. Security Maturity Self-Assessment Framework for SoftwareDevelopment Lifecycle. In The 17th International Conference on Availability, Reliability and Security, August 23–26, 2022, Vienna, Austria. ACM, New York, NY, USA 8 Pages (ARES 2022).
123. Gary & McGraw Sammy, Migue, and J. West, Building Security In Maturity Model (BSIMM) Version 6, 1–65, (2015).
124. Team, S. U. Standard CMMI Appraisal Method for Process Improvement (SCAMPI) A, Version 1.3: Method Definition Document, Handbook CMU/SEI-2011-HB-001March (2011).
125. Aldin, N. A. N., Abdellatif, W. S. E., Elbarbary, Z. M. S., Omar, A. I. & Mahmoud, M. M. Robust speed controller for PMSG wind system based on Harris Hawks optimization via wind speed estimation: A real case study. *IEEE Access.* **11**, 5929–5943 (2023).
126. Khan, R. A. et al. An evaluation framework for communication and coordination processes in offshore software development outsourcing relationship: using fuzzy methods. *IEEE Access.* **7**, 112879–112906 (2019).
127. Khan, T. A. et al. Secure IoMT for disease prediction empowered with transfer learning in healthcare 5.0, the concept and case study. *IEEE Access.* **11**, 39418–39430 (2023).
128. Elghanam, E., Ndiaye, M., Hassan, M. S. & Osman, A. H. Location selection for wireless electric vehicle charging lanes using an integrated TOPSIS and binary goal programming method: A UAE case study. *IEEE Access.* **11**, 94521–94535 (2023).
129. Krishnamoorthy, P. et al. Effective scheduling of Multi-Load automated guided vehicle in spinning mill: A case study. *IEEE Access.* **11**, 9389–9402 (2023).
130. Khan, R. A., Khan, S. U., Alzahrani, M. & Ilyas, M. Security assurance model of software development for global software development vendors. *IEEE Access.* **10**, 58458–58487 (2022).

## Acknowledgements

This work was funded by Princess Nourah bint Abdulrahman University Researchers Supporting Project number (PNURSP2026R896), Princess Nourah bint Abdulrahman University, Riyadh, Saudi Arabia.

## Author contributions

All the authors equally contributed to this manuscript:1. Mujtaba Awan contributed to the conceptualization and methodology of the research, conducted the data analysis, and wrote the original draft of the manuscript. He also contributed to the manuscript review and revision.2. Abu Alam led the design of the research framework and was involved in data collection. Hathal also assisted with writing the discussion and managing the manuscript's revision process.3. Rafiq Ahmad Khan contributed to data collection and validation, especially at the level of the technical implementation of the study. Hussein contributed to data analysis and manuscript revision.4. Hathal S. Alwageed contributed professionally by providing the software and tools for data analysis, helping support the statistical analysis, and contributing to the interpretation of the results. Ismail helped in the final editing and formatting of the manuscript.5. Sarra Ayouni participated in designing the methodology and drafting the review of the literature related to the subject. Sarra also critically revised the manuscript and ensured the clarity and quality of the final draft. She was also partially involved in fundraising for this study.6. Alaa Omran Almagrabi contributed to the study planning, offered professional guidance on the results interpretation, and wrote the manuscript. He is also involved in the critical review of the findings and the literature review. All authors have read and approved the final manuscript. The work was carried out with extraordinary contributions from all the authors, as explained above.

## Funding

This work was funded by Princess Nourah bint Abdulrahman University Researchers Supporting Project number (PNURSP2026R896), Princess Nourah bint Abdulrahman University, Riyadh, Saudi Arabia.

## Declarations

## Competing interests

The authors declare no competing interests.

## Additional information

**Supplementary Information** The online version contains supplementary material available at <https://doi.org/10.1038/s41598-026-37614-8>.

**Correspondence** and requests for materials should be addressed to R.A.K. or S.A.

**Reprints and permissions information** is available at [www.nature.com/reprints](http://www.nature.com/reprints).

**Publisher's note** Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

**Open Access** This article is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License, which permits any non-commercial use, sharing, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if you modified the licensed material. You do not have permission under this licence to share adapted material derived from this article or parts of it. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by-nc-nd/4.0/>.

© The Author(s) 2026