

Decision-Making Dynamics in Secure Software Development Teams: The Role of AI-Supported Systems

Boris Gramchev

Department of Industrial Management
Technical University of Varna
Varna, Bulgaria
gramchev@gmail.com

Marina Marinova-Stoyanova

Department of Industrial Management
Technical University of Varna
Varna, Bulgaria
mpmstoyanova@tu-varna.bg

Abstract—The increasing integration of AI into secure software development is reshaping technical processes and decision-making within teams. While existing research has focused on performance and detection capabilities, less attention has been given to how AI influences the human aspects of decision formation. This study examines how AI-supported security tools affect decision-making dynamics among software developers and security specialists. Using a quantitative survey-based approach, the study explores relationships between trust in AI, transparency, perceived usefulness, team coordination, and decision confidence. The findings show that AI does not directly determine decision outcomes, but operates through socio-cognitive mechanisms, with trust emerging as the central driver of decision confidence. Transparency contributes indirectly by supporting trust, while team coordination strengthens the integration of AI-generated insights into collective decisions. Perceived usefulness functions mainly as a baseline condition for adoption rather than a determinant of outcomes. The study highlights a key implication: AI-supported systems may increase decision confidence without improving decision quality. By framing AI-supported decision-making as a socio-technical process, this research clarifies how AI reshapes not only what decisions are made, but how they are constructed within teams.

Keywords— *Artificial Intelligence; Decision Confidence; Secure Software Development; Team Coordination; Trust.*

I. INTRODUCTION

The growing integration of artificial intelligence into software development has begun to reshape not only how systems are built, but also how decisions are made within development teams. Nowhere is this shift more visible than in security-critical environments, where the identification and prioritization of vulnerabilities increasingly rely on AI-supported tools. These systems analyze large volumes of code, detect anomalies, and generate recommendations [1], [2] that directly influence how risks are assessed and addressed. While their technical capabilities are well documented, less attention has been given to how they affect the human side of decision-making, particularly in

team-based settings where security decisions are rarely made in isolation [3], [4].

Secure software development has traditionally depended on a combination of expertise, judgment, and coordination among developers and security specialists. Decisions about vulnerabilities are often complex, time-sensitive, and context-dependent, requiring not only technical analysis but also interpretation and negotiation within the team. The introduction of AI-supported systems alters this process by inserting algorithmically generated recommendations into the decision flow. These recommendations do not simply provide additional information; they shape how problems are framed and how attention is allocated. As a result, decision-making becomes increasingly mediated, raising important questions about how teams interpret, trust, and act upon AI-generated insights [5], [6].

Despite the rapid adoption of AI in software security, existing research has largely focused on performance improvements, detection accuracy, and system capabilities. Comparatively little empirical work has examined how AI-supported tools influence decision-making dynamics within teams [7], [8]. This represents a critical gap, as the effectiveness of AI in practice depends not only on its technical performance but also on how it is embedded within social and organizational processes. In particular, factors such as trust in AI, perceived transparency of recommendations, and the quality of team coordination may significantly shape how AI outputs are interpreted and used [7], [9].

This study addresses this gap by examining how AI-supported security tools influence decision-making within secure software development teams. The focus is not on whether AI improves detection accuracy, but on how it affects the process through which decisions are formed. Specifically, the study investigates three interrelated questions: how AI-supported tools influence decision-making processes within teams, how developers and security specialists interpret AI-generated

Online ISSN 3044-7771

<https://doi.org/10.68302/std2026.vol3.132>

© 2026 The Author(s). Published by Green Future Technologies.

This is an open-access article under the [Creative Commons Attribution 4.0 International License](#).

recommendations, and which factors determine the effective integration of these recommendations into collective decision-making. By focusing on trust, transparency, perceived usefulness, and team coordination, the study seeks to capture the key mechanisms through which AI reshapes decision dynamics.

Conceptually, the study is grounded in the view that AI-supported decision-making is a socio-technical process, where technological outputs interact with human judgment and team interaction. Prior research on human-AI collaboration suggests that AI systems do not replace decision-making but reconfigure it by altering how information is structured and how authority is distributed [10], [11]. At the same time, research on trust in AI highlights that users' willingness to rely on algorithmic recommendations depends less on technical accuracy alone and more on perceived reliability and contextual understanding [7]. These perspectives suggest that the impact of AI on decision-making is mediated through human and relational factors rather than being directly determined by system capabilities.

Based on this reasoning, the study adopts a model in which transparency contributes to the formation of trust, trust influences decision confidence, and team coordination shapes how individual judgments are translated into collective decisions. Perceived usefulness is treated as a baseline condition that supports adoption but does not directly determine decision outcomes. Rather than formalizing these relationships as strict hypotheses, the study takes an exploratory empirical approach, aiming to identify patterns and relationships within real-world data. This choice reflects the emerging nature of the research area, where rigid hypothesis testing may be less appropriate than capturing the structure of interactions between variables.

The contribution of this study is twofold. First, it provides empirical evidence on how AI-supported systems influence decision-making dynamics in secure software development, moving beyond a purely technical perspective. Second, it offers a conceptual framework that highlights the role of trust, transparency, and coordination as key mechanisms in AI-mediated decision processes. By doing so, the study contributes to a more nuanced understanding of how AI reshapes not only what decisions are made, but how they are formed within teams.

II. AI IN SECURE SOFTWARE DEVELOPMENT

The integration of artificial intelligence into secure software development has accelerated significantly in recent years, reshaping how vulnerabilities are detected, assessed, and prioritized [2], [1]. What was once a largely manual and experience-driven process is now increasingly supported by systems capable of analyzing large volumes of code, identifying patterns, and generating security recommendations in real time. These tools promise not only efficiency gains but also a shift toward more proactive and continuous security practices embedded directly within the development lifecycle [4].

At a functional level, AI is most visibly applied in vulnerability detection. Machine learning models are used to identify known patterns of insecure code, detect anomalies, and flag potential risks that might be overlooked by human reviewers [1]. Unlike traditional rule-based systems, these models can adapt to new data and improve over time, making them particularly valuable in environments where threats evolve rapidly. However, this capability also introduces new forms of uncertainty, as the logic behind detections is often probabilistic rather than deterministic, making it less transparent to users [9].

Beyond detection, AI increasingly plays a role in risk assessment and prioritization. Modern security tools do not simply list vulnerabilities; they rank them, estimate their potential impact, and suggest remediation strategies [1]. In doing so, they move from being passive tools to active participants in the decision-making process. This shift is subtle but important. When a system not only identifies a vulnerability but also recommends what should be addressed first, it begins to shape the decision space itself. Teams are no longer evaluating raw information but interacting with pre-structured judgments generated by the system.

This evolution aligns with broader developments in decision-support systems, where AI is positioned as an augmentation layer rather than a replacement for human expertise [6]. In theory, this augmentation allows teams to combine computational power with human contextual understanding, leading to more informed decisions. In practice, however, the balance between automation and human control is not always stable. As AI systems become more capable, there is a tendency for users to rely more heavily on their outputs, sometimes without fully questioning them. This creates a dynamic where efficiency increases, but critical scrutiny may decrease.

One of the defining characteristics of AI-supported security tools is their ability to operate at scale. They can process vast amounts of code and security data far beyond human capacity, enabling continuous monitoring and near real-time feedback. This is particularly valuable in agile and DevOps environments, where rapid iteration cycles require equally rapid security assessments. AI makes it possible to integrate security checks directly into development pipelines, supporting the idea of "security by design" rather than security as a separate phase [3].

At the same time, the integration of AI into development workflows introduces new challenges related to interpretability and trust. Developers and security specialists must decide how much to rely on AI-generated recommendations, often without full visibility into how those recommendations were produced [12]. This issue becomes more pronounced in complex models, where the reasoning process is not easily accessible. As a result, the effectiveness of AI tools depends not only on their technical accuracy but also on how they are perceived and understood by users [8].

Another important dimension is the interaction between AI systems and team-based work structures. Secure

software development is rarely an individual activity; it involves coordination between multiple roles, including developers, security engineers, and managers. AI tools introduce shared reference points in the form of alerts, scores, and recommendations, which can facilitate communication but can also create new dependencies. Teams may begin to rely on AI outputs as common ground for discussion, which can streamline coordination but may also limit the diversity of viewpoints if those outputs are not critically examined [13].

Recent research [2] also highlights the importance of human oversight in AI-supported security processes. While AI can identify patterns and suggest actions, it lacks contextual awareness and cannot fully account for organizational priorities, trade-offs, or strategic considerations. Human judgment remains essential in interpreting AI outputs and making final decisions. This reinforces the idea that AI should be understood as part of a broader socio-technical system rather than as a standalone solution [11].

Furthermore, the use of AI in security raises concerns related to bias, reliability, and overreliance. Models trained on historical data may reproduce existing biases or fail to detect novel types of vulnerabilities. False positives and false negatives remain a challenge, and overconfidence in AI outputs can lead to missed risks. These limitations do not negate the value of AI but highlight the need for balanced use, where automation is complemented by critical evaluation and continuous validation [14].

Taken together, these developments suggest that AI is transforming secure software development not only at the level of tools but at the level of processes and decision structures. It changes what information is available, how it is interpreted, and how decisions are made within teams. This transformation sets the stage for examining how these changes influence decision dynamics, particularly in terms of trust, transparency, and coordination, which become central factors in determining how AI-supported systems are actually used in practice.

III. TEAM DECISION DYNAMICS IN AI-MEDIATED ENVIRONMENTS

Decision-making in secure software development has traditionally been a collective and iterative process, shaped by technical expertise, risk awareness, and situational context. The introduction of AI-supported systems changes this process in a more subtle but fundamental way. AI does not simply make analysis faster or detection more accurate. It changes how information is formed, how it is interpreted, and how it is validated within teams. As a result, decision-making becomes mediated, moving from a purely human activity toward a hybrid form of interaction where responsibility and judgment are shared between people and systems [5].

A key shift occurs in the nature of decision inputs. Previously, teams relied on direct interpretation of logs, alerts, and contextual signals. With AI, this information is already filtered, prioritized, and framed before it reaches

the team. What appears as input is, in reality, already an interpretation. AI systems highlight what seems important and implicitly suggest what should be done next [13]. This creates efficiency, but it also subtly shapes how problems are understood. Team members are more likely to anchor their reasoning around AI outputs, which can guide attention while at the same time limiting alternative perspectives [10], [6].

This anchoring effect has visible consequences in team interactions. When AI outputs are seen as credible, they help teams move faster and reduce uncertainty. Discussions become more focused, and decisions can be reached with less friction. At the same time, this speed can come at a cost. Teams may converge too quickly on AI-supported conclusions, not because all alternatives have been explored, but because the recommendation already feels sufficiently justified. In high-pressure environments, this tendency becomes even stronger, creating a tension between efficiency and critical evaluation [8].

Trust becomes the central mechanism through which this tension is managed. In AI-mediated environments, trust extends beyond interpersonal relationships to include the system itself. Whether a recommendation is accepted or challenged depends largely on how reliable the system is perceived to be. In many cases, users cannot fully understand how AI outputs are generated, especially when models are complex or opaque [10]. In such situations, trust fills the gap left by limited understanding. Teams act not because they fully comprehend the reasoning, but because they believe the system is dependable [7].

The role of transparency is therefore more complicated than it might initially appear. While it is often assumed that more explanation leads to better decisions, this is not always the case. Detailed explanations can improve understanding, but they can also increase cognitive load or create confusion, especially when team members have different levels of expertise. In practice, the same explanation may be interpreted differently across individuals, which can introduce misalignment rather than clarity. Transparency, therefore, does not directly improve decisions, but works through trust and interaction within the team [15], [16].

Team coordination is equally important in this process. Security decisions are inherently collaborative, requiring input from developers, security specialists, and other stakeholders. AI systems introduce shared reference points, such as alerts and risk scores, which can support alignment. However, coordination does not happen automatically. The meaning of AI outputs still needs to be discussed, negotiated, and aligned across the team. Different roles, experiences, and expectations can lead to different interpretations of the same recommendation, making communication a critical part of the decision process [13].

Seen in this way, decision-making unfolds across several interconnected layers. At the informational level, AI structures what is visible and what is prioritized. At the cognitive level, individuals interpret these signals based on trust, experience, and perceived usefulness. At the

relational level, teams align their interpretations and arrive at shared decisions. The effectiveness of the overall process depends on how well these layers work together. Weakness in any one of them can disrupt the entire decision dynamic.

Importantly, AI does not replace human judgment but repositions it. The focus shifts from analyzing raw data to interpreting AI-generated representations of that data [1]. This shift carries risks. AI outputs may appear objective, but they are shaped by underlying assumptions, training data, and design choices that are not always visible to users. As a result, teams may rely on recommendations that carry hidden biases or limitations [14].

Another consequence is the emergence of dependency. As teams rely more on AI for prioritization and risk assessment, they may engage less with the underlying data. While this improves efficiency, it can reduce critical thinking and limit the diversity of viewpoints within the team. Over time, this may lead to more uniform decisions, but not necessarily better ones. This is particularly problematic in security contexts, where adaptability and critical judgment are essential [2].

Overall, decision-making in AI-supported environments is best understood as a socio-technical process. Its effectiveness depends not only on the capabilities of AI systems, but also on how they are interpreted, trusted, and integrated into team interactions. By focusing on trust, transparency, and coordination, this study captures the key mechanisms through which AI reshapes decision dynamics and sets the stage for examining their impact on decision outcomes.

IV. RESEARCH METHODOLOGY

This study adopts a quantitative research design to examine how AI-supported security tools influence decision-making dynamics within secure software development teams. Given the exploratory nature of the research questions and the need to capture perceptions, attitudes, and behavioral tendencies across different roles, a structured survey approach was selected as the most appropriate method. Survey-based research has been widely used in studies examining technology adoption, trust in automated systems, and team-level processes, as it allows for the systematic measurement of latent constructs such as trust, perceived usefulness, and decision confidence [16], [17]. This approach is particularly relevant in the context of human-AI collaboration, where decision-making emerges from the interaction between algorithmic recommendations and human judgment [5].

The empirical setting of the study focuses on software developers and security specialists working in security-sensitive development environments. These include professionals involved in vulnerability detection, threat assessment, and secure coding practices, where AI-supported tools are increasingly integrated into daily workflows. A total of 96 participants were included in the study, representing a balanced mix of technical roles. Approximately 62% of respondents identified as software developers, while 38% reported primary responsibilities in cybersecurity or security engineering. The average

professional experience among participants was 5.8 years, indicating a sample with sufficient exposure to both traditional and AI-supported development practices.

Data collection was conducted through an online questionnaire designed to capture key dimensions of AI-supported decision-making. All items were measured using a five-point Likert scale ranging from strongly disagree (1) to strongly agree (5), allowing for consistent quantification of subjective perceptions. The questionnaire was structured to ensure clarity and accessibility, avoiding overly technical language while maintaining conceptual precision. Prior to distribution, the instrument was reviewed and refined to ensure internal consistency and face validity. Each construct was measured using two items adapted to the context of AI-supported security tools, ensuring both conceptual clarity and measurement consistency.

The study examines five core variables that reflect both technological and team-level aspects of decision-making in AI-supported environments. Trust in AI-supported tools captures the extent to which participants perceive AI-generated recommendations as reliable and accurate. In AI-mediated environments, trust plays a critical role in determining whether algorithmic outputs are accepted, questioned, or ignored within team decision processes [7]. Decision confidence reflects the degree to which AI tools enhance individuals' certainty when making security-related decisions. Perceived usefulness assesses the practical value of AI tools in improving vulnerability detection and prioritization processes. AI transparency refers to the extent to which the reasoning behind AI-generated outputs is understandable and interpretable. Finally, team coordination captures how effectively team members align their actions and interpretations when AI-generated insights are introduced into the decision process.

Decision confidence was selected as the dependent variable as it represents the most immediate and measurable outcome of AI-supported decision-making at the individual level. In the context of secure software development, objective decision quality is difficult to assess directly within a survey-based design, particularly given the complexity and context-specific nature of security incidents. As a result, decision confidence serves as a proximal indicator of decision outcomes, reflecting the extent to which individuals feel assured in their judgments when interacting with AI-generated recommendations. This choice is consistent with prior research on technology-supported decision-making, where confidence is often used as an operational proxy for decision effectiveness in situations where objective performance measures are not readily observable. However, it is important to acknowledge that decision confidence does not necessarily imply decision accuracy, a distinction that is further addressed in the discussion of the findings.

To ensure the reliability of the measurement model, internal consistency was assessed using Cronbach's alpha. All constructs demonstrated satisfactory reliability, with values ranging between 0.78 and 0.85, exceeding commonly accepted thresholds for exploratory research

[18]. These results indicate that the measurement scales are sufficiently stable for subsequent analysis.

The analytical approach combines descriptive statistics, correlation analysis, and multiple regression modeling to examine relationships between variables. Correlation analysis was used to identify the strength and direction of associations between key constructs, while regression analysis was employed to test the predictive effects of trust, transparency, and coordination on decision confidence. In addition, a moderation analysis was conducted to explore whether team coordination strengthens the relationship between trust in AI and decision confidence. This approach allows for a more nuanced understanding of how individual perceptions and team processes interact in shaping decision outcomes.

Overall, the methodological design provides a structured yet flexible framework for capturing the interplay between AI-supported systems and team-level decision dynamics. By combining established measurement approaches with a focused set of variables tailored to secure software development, the study offers a credible empirical basis for examining how AI influences not only technical processes but also the social and cognitive mechanisms underlying collective decision-making.

While the preceding sections provide the conceptual background for understanding AI-supported decision-making, the primary contribution of this study lies in the empirical examination of how trust, transparency, coordination, and perceived usefulness shape decision confidence in secure software development teams.

V. RESULTS

The empirical analysis provides direct evidence regarding the mechanisms through which AI-supported systems influence decision-making within secure software development teams. Overall, the findings indicate that AI tools are perceived as valuable contributors to security-related processes, but their influence is not uniform across variables. The results reveal a differentiated structure in which trust emerges as the dominant driver of decision confidence, while transparency and team coordination play more supportive roles.

Descriptive statistics are presented in Table 1. Perceived usefulness reports the highest mean, indicating that participants clearly recognize the functional value of AI tools in vulnerability detection and prioritization. Trust in AI and decision confidence also show relatively strong levels, although with greater variability, suggesting differences in how individuals evaluate and rely on AI-generated recommendations. Team coordination demonstrates moderately positive levels, while transparency remains comparatively lower, indicating that understanding the reasoning behind AI outputs is still a partial limitation.

TABLE 1. DESCRIPTIVE STATISTICS OF STUDY VARIABLES

Variable	M	SD
Trust in AI	3.85	0.81
Transparency	3.63	0.77
Decision Confidence	3.82	0.85
Team Coordination	3.76	0.68
Perceived Usefulness	3.95	0.68

To examine the relationships between variables, Pearson correlation analysis was conducted. The results, shown in Table 2, indicate that all variables are positively related, although with varying strength. Trust in AI demonstrates a strong positive association with decision confidence, confirming its central role in shaping how confident individuals feel when making security-related decisions. Transparency is also moderately associated with both trust and decision confidence, suggesting that interpretability contributes to decision-making, but not as a dominant factor.

Team coordination shows a positive but more moderate relationship with decision confidence compared to initial expectations, indicating that while coordination supports decision processes, its influence may be more contextual than direct. Perceived usefulness is positively related to all variables, particularly trust, but the strength of these relationships remains moderate, suggesting that usefulness functions more as a background condition than a primary driver. Overall, the correlation pattern supports the structure of the proposed model, while also indicating differences in the relative weight of each variable.

TABLE 2. PEARSON CORRELATIONS AMONG VARIABLES ($p < .01$)

Variable	1	2	3	4	5
Trust in AI	1.00				
Transparency	0.62**	1.00			
Decision Confidence	0.68**	0.55**	1.00		
Team Coordination	0.24**	0.23**	0.40**	1.00	
Perceived Usefulness	0.41**	0.33**	0.43**	0.23**	1.00

To further test the predictive relationships, a multiple regression analysis was conducted with decision confidence as the dependent variable. The results are presented in Table 3. Trust in AI emerges as the strongest predictor, with a substantial and highly significant effect, indicating that perceived reliability of AI-generated recommendations is the primary factor shaping decision confidence.

Team coordination also demonstrates a significant positive effect, although weaker than trust, suggesting that collective processes still play an important role in reinforcing confidence, even if their direct influence is more limited than initially assumed. In contrast, transparency does not reach conventional levels of statistical significance, indicating that while interpretability

contributes to trust, it does not independently drive decision confidence when other variables are considered simultaneously.

The overall model explains a substantial proportion of variance in decision confidence, confirming that the selected variables capture key aspects of AI-supported decision-making dynamics, while also highlighting the dominant role of trust over other factors. The model explains approximately 56% of the variance in decision confidence ($R^2 = .559$).

TABLE 3. MULTIPLE REGRESSION ANALYSIS PREDICTING DECISION CONFIDENCE (DEPENDENT VARIABLE: DECISION CONFIDENCE)

Predictor	β	SE	p
Trust in AI	0.50	0.10	< .001
Transparency	0.17	0.10	0.082
Team Coordination	0.27	0.09	< .01
Perceived Usefulness	0.16	0.10	0.095

To further examine the role of team processes, a moderation analysis was conducted focusing on the interaction between trust in AI and team coordination. The results indicate a statistically significant interaction effect (see Table 4), suggesting that the relationship between trust and decision confidence becomes stronger in teams characterized by higher levels of coordination. In such contexts, AI-generated insights are more effectively integrated into shared decision processes, amplifying their impact. In contrast, in less coordinated teams, trust alone is less sufficient to produce high levels of decision confidence.

TABLE 4. MODERATION ANALYSIS: INTERACTION EFFECT. (DEPENDENT VARIABLE: DECISION CONFIDENCE)

Predictor	β	SE	p
Trust in AI	1.34	0.48	< .001
Team Coordination	1.02	0.50	< .05
Trust $\times$ Coordination	0.18	0.13	< .05

An additional pattern emerges with respect to perceived usefulness. Although it is positively associated with other variables, it does not reach statistical significance in the regression model. This suggests that perceived usefulness may function primarily as a baseline condition for adoption rather than as a direct determinant of decision outcomes. While participants appear to recognize the practical value of AI tools, this recognition does not consistently translate into higher decision confidence or more effective decision-making.

To assess the potential risk of common method bias, Harman's single-factor test was conducted using all measurement items. The results indicate that the first unrotated factor accounts for 45.40% of the total variance, which is below the commonly accepted threshold of 50% (see Table 5). This suggests that common method bias is

unlikely to pose a serious threat to the validity of the findings.

TABLE 5. HARMAN'S SINGLE-FACTOR TEST RESULTS

Component	% of Variance Explained	Cumulative %
1	45.40	45.40
2	12.08	57.48
3	9.60	67.08
4	7.97	75.05
5	5.82	80.87

Taken together, the findings point to a more differentiated structure of AI-supported decision-making than initially assumed. Trust in AI-supported systems emerges as the central mechanism, while coordination plays a supporting role and transparency functions primarily as an indirect contributor. Rather than replacing human judgment, AI appears to reshape how confidence is constructed and how decisions are collectively validated within secure software development teams.

VI. DISCUSSION

The empirical findings provide several insights into the role of AI-supported systems in secure software development teams. Rather than confirming a purely technology-driven view of decision-making, the results suggest that human and team-level factors remain central in determining how AI-generated recommendations are interpreted and used.

The findings of this study offer a more nuanced understanding of how AI-supported systems influence decision-making within secure software development teams. While prior expectations often emphasize the technical capabilities of AI tools, the results here suggest that their impact is primarily shaped by human interpretation, trust, and team interaction rather than by functionality alone.

The most prominent finding is the central role of trust in AI-supported systems. Across both correlation and regression analyses, trust emerges as the strongest predictor of decision confidence, indicating that the perceived reliability of AI-generated recommendations is the key mechanism through which these tools influence decision-making. This reinforces the idea that AI does not directly determine decisions but instead operates through human acceptance. When developers trust the outputs, they are more willing to rely on them, integrate them into their reasoning, and act upon them. Without such trust, even technically accurate systems may fail to influence decisions in meaningful ways.

In contrast, transparency plays a more limited role than initially assumed. Although it is positively associated with both trust and decision confidence, its effect becomes statistically weak when other variables are considered simultaneously. This suggests that while understanding AI-generated outputs contributes to the formation of trust, it

does not independently drive decision confidence. In practical terms, teams may rely on AI recommendations even when they do not fully understand how they are produced, as long as those recommendations are perceived as reliable. This finding challenges the common assumption that explainability is always a primary driver of effective AI use and instead positions it as a supporting rather than leading factor.

Team coordination also contributes to decision confidence, although its effect is more moderate than expected. The results indicate that coordination strengthens the integration of AI-generated insights into decision processes, particularly in teams where communication and alignment are already well established. The moderation analysis further supports this interpretation, showing that the relationship between trust and decision confidence becomes stronger in more coordinated teams. This highlights that AI tools do not operate in isolation but are embedded within social systems where their influence depends on how information is shared, discussed, and collectively interpreted.

An important implication of these findings is the distinction between decision confidence and decision quality. While the results clearly show that AI-supported systems increase confidence, this does not necessarily mean that decisions are objectively better. In fact, it is possible that AI tools create a sense of assurance that is not always matched by improved outcomes. This introduces a potential risk in AI-supported environments, where increased confidence may reduce critical evaluation and lead to overreliance on automated recommendations. In this sense, AI may reshape not only how decisions are made, but also how they are perceived and justified within teams.

The role of perceived usefulness further supports this interpretation. Although participants recognize the practical value of AI tools, usefulness does not emerge as a significant predictor of decision confidence. This suggests that functional benefits alone are insufficient to explain how AI influences decision-making. Instead, usefulness appears to function as a baseline condition for adoption, while trust and coordination determine how AI is actually used in practice.

Taken together, these findings point to a layered model of AI-supported decision-making. At the foundational level, usefulness enables adoption. At the interpretive level, transparency contributes to trust. At the cognitive level, trust shapes individual confidence. Finally, at the collective level, team coordination determines how effectively this confidence is translated into shared decisions. This layered structure emphasizes that the impact of AI in secure software development is not purely technological but deeply social and cognitive.

From a practical perspective, the results suggest that organizations should focus not only on improving the technical performance of AI tools, but also on fostering trust and strengthening team coordination. Efforts to increase transparency should be balanced with the

recognition that interpretability alone may not drive effective use. Instead, the challenge lies in integrating AI systems into team processes in a way that supports informed, reflective, and collectively validated decision-making.

A further limitation of the study relates to the sample size, which may limit the statistical power to detect interaction effects. While the sample is sufficient for the main regression analysis, larger samples would be required to more reliably assess moderation relationships in future research.

VII. CONCLUSION

This study set out to examine how AI-supported tools influence decision-making within secure software development teams, shifting the focus from technical performance to the human and team-level dynamics that shape actual outcomes. The findings show that the impact of AI is not direct, but mediated through trust, interpretation, and coordination. Among the examined factors, trust in AI-supported systems emerges as the strongest driver of decision confidence, highlighting that the effectiveness of AI depends less on its capabilities alone and more on how it is perceived and accepted by users.

At the same time, transparency and perceived usefulness, while relevant, do not independently determine decision outcomes. Instead, they contribute indirectly by supporting trust and facilitating interpretation. Team coordination further strengthens the integration of AI-generated insights, emphasizing that decision-making remains a collective process even in highly automated environments. An important implication is that increased confidence does not necessarily imply improved decision quality. AI systems can create a sense of certainty that may not always be justified, introducing potential risks of overreliance. Overall, the study demonstrates that AI in secure software development functions as part of a broader socio-technical system. Its effectiveness depends not only on technological sophistication but on how it is embedded within human judgment and team interaction.

REFERENCES

- [1] Kaur, R., Gabrijelčić, D., & Klobučar, T. (2023). Artificial intelligence for cybersecurity: Literature review and future research directions. *Information Fusion*, 97, 101804. <https://doi.org/10.1016/j.inffus.2023.101804>
- [2] Berente, N., Gu, B., Recker, J., & Santhanam, R. (2021). Managing artificial intelligence. *MIS Quarterly*, 45(3), 1433–1450. <https://doi.org/10.25300/MISQ/2021/16274>
- [3] Binns, R. (2022). Human judgement in algorithmic loops: Individual justice and automated decision-making. *Regulation & Governance*, 16(1), 197–211. <https://doi.org/10.1111/rego.12358>
- [4] Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- [5] Dellermann, D., Ebel, P., Söllner, M., & Leimeister, J. M. (2019). Hybrid intelligence. *Business & Information Systems Engineering*, 61, 637–643. <https://doi.org/10.1007/s12599-019-00595-2>
- [6] Faraj, S., Pachidi, S., & Sayegh, K. (2018). Working and organizing in the age of the learning algorithm. *Information and*

- Organization*, 28(1), 62–70.
<https://doi.org/10.1016/j.infoandorg.2018.02.005>
- [7] Glikson, E., & Woolley, A. W. (2020). Human Trust in Artificial Intelligence: Review of Empirical Research. *Academy of Management Annals*, 14, 627–660.
<https://doi.org/10.5465/annals.2018.0057>
- [8] Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2019). *Multivariate data analysis* (8th ed.). Cengage Learning.
- [9] Jarrahi, M. H. (2018). Artificial intelligence and the future of work: Human–AI symbiosis in organizational decision making. *Business Horizons*, 61(4), 577–586.
<https://doi.org/10.1016/j.bushor.2018.03.007>
- [10] Kaur, D., Uslu, S., Rittichier, K. J., & Durresi, A. (2022). Trustworthy artificial intelligence: A review. *ACM Computing Surveys*, 55(2), 1–38. <https://doi.org/10.1145/3491209>
- [11] McKnight, D. H., Carter, M., Thatcher, J. B., & Clay, P. F. (2011). Trust in a specific technology: An investigation of its components and measures. *ACM Transactions on Management Information Systems*, 2(2), 1–25.
<https://doi.org/10.1145/1985347.1985353>
- [12] Nguyen, T. T., & Reddi, V. J. (2023). Deep reinforcement learning for cyber security. *IEEE Transactions on Neural Networks and Learning Systems*, 34(8), 3779–3795.
<https://doi.org/10.1109/TNNLS.2021.3121870>
- [13] Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation–augmentation paradox. *Academy of Management Review*, 46(1), 192–210.
<https://doi.org/10.5465/amr.2018.0072>
- [14] Sarker, I. H. (2023). Machine learning for intelligent data analysis and automation in cybersecurity: Current and future prospects. *Annals of Data Science*, 10(6), 1473–1498.
<https://doi.org/10.1007/s40745-022-00444-2>
- [15] Ameh, J. E., Otebolaku, A., Shenfield, A., Ikpehai, A., & Sule, D. (2025). *Machine learning approaches in software vulnerability detection: A systematic review and analysis of contemporary methods* [Preprint]. *Research Square*.
<https://doi.org/10.21203/rs.3.rs-5975490/v1>
- [16] Suresh, H., Gomez, S. R., Nam, K. K., & Satyanarayan, A. (2021). Beyond expertise and roles: A framework to characterize the stakeholders of interpretable machine learning and their needs. *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*, 74, 1–16.
<https://doi.org/10.1145/3411764.3445088>
- [17] Tahaei, M., Vaniea, K., & Rashid, A. (2023). Embedding Privacy Into Design Through Software Developers: Challenges and Solutions. *IEEE Security & Privacy* 21(1), 49–57.
<https://doi.org/10.1109/MSEC.2022.3204364>
- [18] Ofusori, L., Bokaba, T., & Mhlongo, S. (2024). Artificial Intelligence in Cybersecurity: A Comprehensive Review and Future Direction. *Applied Artificial Intelligence*, 38(1).
<https://doi.org/10.1080/08839514.2024.2439609>