



# Explainable Threat Intelligence for Nontechnical Decision Makers

**Dr. Sophie M. Richardson**

Associate Professor, University of Manchester, United Kingdom

## Abstract

Cyber threat intelligence has become an important component of organizational cybersecurity because it transforms threat-related data into contextual knowledge that can support security decisions. Yet the usefulness of threat intelligence depends not only on technical accuracy but also on whether intended users can understand its significance, uncertainty, organizational consequences, and recommended actions. This requirement becomes particularly important when cyber-risk decisions involve executives, business managers, finance officers, public administrators, compliance leaders, and other stakeholders who do not possess specialist cybersecurity expertise. The present study develops the concept of Explainable Threat Intelligence (ETI) as a human-centered communication layer that converts technically rich cyber threat intelligence into decision-oriented explanations without eliminating evidence, uncertainty, or analytical traceability.

NIST defines cyber threat intelligence as threat information that has been analyzed or enriched to provide context for decision-making, while NIST Cybersecurity Framework 2.0 is deliberately structured so that cybersecurity outcomes can be understood by executives, managers, and practitioners with different levels of technical expertise. A simulation-based comparative design representing 270 hypothetical nontechnical decision makers was developed to compare raw technical CTI, conventional executive summaries, and explainable CTI. The simulated results indicate mean decision accuracies of 60.4%, 76.9%, and 87.3%, respectively. Explainable CTI also produced higher simulated decision confidence and shorter decision time.

The findings are methodological illustrations rather than empirical population estimates, but they demonstrate how explanation quality may be evaluated through measurable decision outcomes. The paper concludes that effective threat intelligence for nontechnical audiences should communicate threat relevance, affected business assets, confidence, supporting evidence, likely consequences, uncertainty, and prioritized response options while retaining links to underlying technical detail. Such an approach can strengthen cyber-risk governance by connecting technical detection capabilities with accountable organizational decision-making.

**Keywords:** cyber threat intelligence; explainable threat intelligence; cybersecurity decision-making; nontechnical decision makers; explainable artificial intelligence; cyber-risk communication; threat-informed defense; executive cybersecurity



## 1. Introduction

Cybersecurity decisions increasingly extend beyond security operations centers and information-technology departments. Major incidents can affect financial continuity, legal obligations, organizational reputation, supply chains, customer services, public safety, and strategic priorities. Consequently, executives, operational managers, compliance officers, public administrators, board members, and other nontechnical stakeholders may be required to decide whether an organization should isolate systems, suspend services, authorize emergency expenditure, communicate with external parties, modify operational plans, or accept residual cyber risk. Cyber threat intelligence can support such decisions by converting observations about threat actors, indicators, tactics, techniques, vulnerabilities, and incidents into contextual knowledge. NIST describes cyber threat information as information that can assist organizations in identifying, assessing, monitoring, and responding to cyber threats, and emphasizes that enrichment and contextualization make such information more actionable for decision-making.

A persistent problem, however, is that the form in which threat intelligence is technically useful is not necessarily the form in which it is strategically understandable. Security analysts may interpret malicious IP addresses, file hashes, detection signatures, exploit chains, ATT&CK technique identifiers, confidence assessments, vulnerability information, command-and-control infrastructure, and telemetry patterns directly. Nontechnical decision makers commonly require a different abstraction: what is happening, why it matters to the organization, which services or objectives may be affected, how confident analysts are, what is still uncertain, what actions are available, and what may happen if action is delayed. MITRE ATT&CK provides a common knowledge base of adversary tactics and techniques grounded in real-world observations, while STIX and TAXII provide standardized mechanisms for representing and exchanging CTI. These resources improve analytical structure and interoperability, but they do not by themselves guarantee that intelligence is understandable to an executive or operational stakeholder.

Ainslie, Thompson, Maynard, and Ahmad argue that CTI should be understood in relation to organizational security decision-making rather than simply as a collection of threat feeds or technical indicators. Their review identifies the need to better connect intelligence practices with decisions made at different organizational levels. This challenge is intensified by automation and artificial intelligence. Modern security platforms can prioritize or classify large numbers of alerts, yet opaque automated reasoning can create difficulties when users need to understand why an alert has been prioritized or which evidence supports a recommendation. Research on explainable artificial intelligence similarly emphasizes that model outputs become more useful when users can inspect intelligible reasons for predictions rather than relying exclusively on black-box scores. Ribeiro, Singh, and Guestrin introduced LIME to provide locally interpretable explanations, while Lundberg and Lee developed SHAP to assign interpretable feature contributions to model predictions.

Explainability is nevertheless broader than displaying model feature importance. A business leader deciding whether to stop production is unlikely to benefit from a technically correct SHAP plot unless the explanation also clarifies operational implications. Explainable threat intelligence therefore requires translation between technical evidence and decision context. Recent work by Rastogi and colleagues on explainable AI in threat intelligence found that cybersecurity professionals expressed interest in elements such as attack attribution, confidence information, and explanations that improve



alert interpretability and triage. Their findings reinforce the principle that effective explanations should not simply expose internal model mechanics but should improve the user's ability to understand and act upon a security assessment. NIST's broader trustworthy-AI guidance likewise treats explainability and interpretability, transparency, reliability, security, and accountability as interconnected characteristics rather than isolated technical features.

The present study therefore proposes Explainable Threat Intelligence (ETI) as threat intelligence deliberately structured for comprehensibility, traceability, uncertainty awareness, and actionability among decision makers who do not require detailed cybersecurity expertise to discharge their responsibilities. Because no real participant dataset was supplied, the study does not fabricate empirical survey findings. Instead, it integrates established CTI, cybersecurity-governance, and explainability literature with a reproducible simulation involving 270 hypothetical decision-maker observations. The analytical comparison examines three presentation formats—raw technical CTI, conventional executive-summary CTI, and ETI—to illustrate how future empirical studies could evaluate decision accuracy, decision confidence, and decision time.

## 2. Objectives of the Study

### 2.1 To Examine the Communication Gap Between Technical CTI and Organizational Decision-Making

The first objective is to examine why technically accurate threat intelligence may still be ineffective when it does not provide nontechnical decision makers with sufficient contextual interpretation. Particular attention is given to business relevance, uncertainty, evidence, expected consequences, and recommended actions.

### 2.2 To Evaluate the Decision-Support Potential of Explainable Threat Intelligence

The second objective is to assess, through simulation, whether an explainable CTI format could produce stronger decision outcomes than raw technical intelligence or a conventional short executive summary. The principal analytical outcomes are decision accuracy, decision confidence, and decision time.

### 2.3 To Develop a Practical Explainability Framework for CTI

The third objective is to identify design principles that can help cybersecurity teams communicate threat intelligence to executives and other nontechnical stakeholders without oversimplifying evidence. The study emphasizes layered explanations, traceability, confidence communication, organizational impact, and decision-oriented recommendations.

## 3. Review of Literature

### 3.1 From Threat Data to Decision-Oriented Intelligence

NIST's *Guide to Cyber Threat Information Sharing* established an important distinction between raw threat information and information that becomes useful through correlation, enrichment, interpretation, and contextualization. The guidance identifies indicators, tactics, techniques and procedures, security alerts, intelligence reports, and recommended security configurations as forms of threat information and notes that combining information from multiple sources can reduce ambiguity and improve actionability.

This conceptual distinction is central to explainable threat intelligence because a decision maker rarely needs every raw observable; instead, the decision maker needs sufficient evidence to understand the nature, credibility, and organizational relevance of the threat.

OASIS subsequently standardized machine-readable CTI through STIX and TAXII. STIX provides a structured language for expressing cyber threat and observable information, while TAXII provides an application-layer protocol for exchanging threat information. These standards improve interoperability and automated exchange among security technologies and organizations. Their principal strength is machine-to-machine consistency, but ETI addresses a complementary requirement: human-to-decision consistency. An intelligence record may be syntactically interoperable while still being difficult for an executive to interpret.

MITRE ATT&CK contributes another essential layer by organizing adversary behavior into tactics and techniques based on real-world observations. ATT&CK gives analysts a common language for describing adversary behavior and can support threat modeling, defensive prioritization, detection analysis, and threat-informed defense. For nontechnical audiences, however, an ATT&CK identifier becomes meaningful only when it is translated into questions such as which business process could be disrupted, which assets are exposed, how the technique contributes to an attack pathway, and what preventive or response decision is required.

Ainslie and colleagues' 2023 review directly connected CTI with security decision-making and argued for greater attention to the practical use of threat intelligence within organizations. Saeed and colleagues' systematic literature review similarly emphasizes CTI's contribution to organizational cybersecurity resilience and the growing use of analytics for extracting value from large and heterogeneous threat-information sources. Together, these studies indicate that CTI effectiveness cannot be evaluated only by the volume of collected indicators or feeds. Its value also depends on whether intelligence changes priorities, defensive actions, resource allocation, or organizational preparedness.

### **3.2 Explainability, Trust, and Cybersecurity Interpretation**

Explainable artificial intelligence literature provides a useful theoretical foundation for ETI. Ribeiro, Singh, and Guestrin's LIME framework demonstrated that users can receive interpretable local explanations of otherwise complex classifier outputs. The authors specifically related explanation to appropriate trust and demonstrated its usefulness in tasks where human users need to determine whether a prediction should be relied upon. Lundberg and Lee later proposed SHAP, a unified framework that assigns feature-importance values to individual predictions and provides a theoretically grounded mechanism for interpreting complex models.

Cybersecurity introduces additional explainability requirements because security decisions often involve uncertainty, adversarial behavior, incomplete information, rapidly changing evidence, and potentially high operational consequences. Charmet and colleagues' literature survey on explainable AI for cybersecurity reviewed the role of explanation across security applications and highlighted the broader tension between increasingly complex detection systems and the need for interpretable outputs. Mendes and Rios likewise observed that black-box AI models can create adoption and comprehension problems in cybersecurity and reviewed techniques used to make model decisions more interpretable.



Rastogi, Dhanuka, Saxena, Mairal, and Nguyen brought this problem specifically into the threat-intelligence environment. Their industry-oriented research examined how security analysts process AI-based alerts and reported strong interest in explanations involving attribution, confidence scores, and feature contributions. Although their principal audience includes security professionals rather than nontechnical executives, the implications extend further: different users require explanations that correspond to their decision responsibilities.

For this reason, the present paper distinguishes model explainability from decision explainability. Model explainability answers questions such as why a classifier assigned a malicious score. Decision explainability answers a broader set of questions: why the threat matters to this organization, which evidence supports the assessment, what uncertainty remains, what consequences are plausible, what options are available, and why one response is recommended over another. An effective ETI system may use LIME, SHAP, ATT&CK mappings, causal reasoning, analyst judgment, or other technical methods internally, but its executive output must be constructed around organizational decisions.

### 3.3 Cybersecurity Governance and the Need for Audience-Adaptive Intelligence

NIST Cybersecurity Framework 2.0 strengthens the governance dimension of cybersecurity by introducing GOVERN alongside IDENTIFY, PROTECT, DETECT, RESPOND, and RECOVER. The framework states that its high-level outcomes are intended to be understandable by executives, managers, and practitioners regardless of cybersecurity expertise and can be used to understand, prioritize, and communicate cybersecurity risks. This creates a strong institutional basis for audience-adaptive threat intelligence.

A nontechnical decision maker does not need to become a threat hunter to participate responsibly in cyber-risk governance. The intelligence product should instead support the decisions for which that person is accountable. A chief financial officer may need exposure estimates and expected financial consequences. A chief operating officer may need service-disruption scenarios and recovery options. A public administrator may need continuity implications and public communication requirements. A board member may need strategic risk, risk appetite, and assurance information. The underlying technical evidence can remain the same while the explanation layer changes according to decision context.

This audience-adaptive approach must nevertheless avoid oversimplification. An explanation that merely labels risk as “high” can create false certainty if underlying confidence is weak. Similarly, visually attractive dashboards may encourage action without making evidence quality or uncertainty visible. NIST's AI Risk Management Framework treats explainability together with reliability, transparency, accountability, and context of use, which supports the position that explanation quality should be evaluated not by simplicity alone but by whether it enables informed and appropriately calibrated human judgment.

The literature therefore reveals a specific research gap. CTI research has examined collection, sharing, standards, automation, organizational resilience, and security decision-making, while XAI research has examined model transparency and analyst trust. Less attention has been directed toward a unified communication framework that converts technical CTI into evidence-preserving explanations specifically optimized for nontechnical organizational decision makers. The present study addresses that gap through a conceptual ETI model and simulation-based evaluation.

## 4. Research Methodology

### 4.1 Research Design

The study adopts a simulation-based comparative decision-support design. No human participants were recruited and no actual organizational security incidents or confidential threat-intelligence records were used. A synthetic dataset representing 270 hypothetical nontechnical decision makers was created solely to demonstrate how alternative CTI presentation formats could be evaluated experimentally.

- The hypothetical observations were divided equally among three presentation conditions:
- Raw Technical CTI: Threat intelligence dominated by technical indicators, vulnerability references, ATT&CK identifiers, event telemetry, and specialist terminology.
- Executive Summary CTI: Technical information condensed into a short narrative containing threat severity and a high-level recommendation.
- Explainable CTI: A structured intelligence product containing threat description, organizational relevance, affected assets or processes, supporting evidence, confidence level, uncertainty, likely consequences, prioritized actions, and access to underlying technical evidence.
- Each hypothetical decision maker was represented as making five scenario-based decisions concerning incident escalation, operational continuity, resource allocation, risk acceptance, or implementation of countermeasures.
- 

### 4.2 Proposed Measurement Instrument

The following five-item questionnaire is proposed for a future empirical validation study. Responses may be collected using a five-point Likert scale ranging from 1 = Strongly Disagree to 5 = Strongly Agree.

**Table 1: Proposed Questionnaire for Evaluating Explainable Threat Intelligence**

| Item | Proposed Statement                                                                                                       | Construct                              |
|------|--------------------------------------------------------------------------------------------------------------------------|----------------------------------------|
| Q1   | The threat-intelligence report clearly explains why the identified cyber threat matters to my organization.              | Organizational relevance               |
| Q2   | I can understand the evidence and reasoning supporting the threat assessment without specialist cybersecurity knowledge. | Explanation comprehensibility          |
| Q3   | The report clearly communicates the level of confidence and uncertainty associated with the assessment.                  | Uncertainty transparency               |
| Q4   | The recommended response options are sufficiently clear for me to make an organizational decision.                       | Decision actionability                 |
| Q5   | I am confident that I can justify my decision using the information provided in the threat-intelligence report.          | Decision confidence and accountability |

**Source:** Proposed by the author for future empirical testing. The questionnaire was not administered in the present study.



A full empirical study should expand these constructs into validated multi-item scales and test reliability, convergent validity, discriminant validity, and measurement invariance. Decision accuracy should ideally be evaluated independently against expert-defined scenario criteria rather than through self-report alone.

### 4.3 Simulation Procedure and Statistical Analysis

A fixed random seed was used to make the synthetic analysis reproducible. Ninety observations were generated for each intelligence-presentation condition. Decision accuracy represented the proportion of five hypothetical scenario decisions answered consistently with a predetermined expert decision criterion. Decision confidence was represented on a five-point scale. Decision time was represented in minutes.

The simulated population parameters intentionally represented progressively stronger decision support across the three conditions. These parameters constitute analytical assumptions, not claims about real decision makers. Random variation was introduced so that outcomes differed across simulated individuals.

Descriptive statistics were calculated for decision accuracy, confidence, and decision time. A one-way analysis of variance was then used to demonstrate how researchers could test whether differences among presentation formats were statistically distinguishable. Because the data are synthetic, statistical significance indicates consistency within the generated simulation and must not be interpreted as evidence about real organizations.

Ethical approval was not required for this purely synthetic methodological demonstration. A future human-participant study should obtain appropriate institutional ethics approval or exemption, provide informed consent, avoid collecting unnecessary employment identifiers, protect potentially sensitive cybersecurity judgments, and ensure that participation cannot adversely affect employment evaluation.

## 5. Analysis

### 5.1 Decision Accuracy Across Intelligence Formats

The most direct simulated effect concerns decision accuracy. Participants represented in the raw technical CTI condition achieved a mean accuracy of 60.4%, whereas the executive-summary condition achieved 76.9%. The explainable CTI condition produced the highest simulated accuracy at 87.3%.

The difference between raw technical CTI and explainable CTI was approximately 26.9 percentage points. The result demonstrates the central analytical proposition of the paper: technically richer information does not necessarily generate better organizational decisions when the decision maker cannot efficiently interpret its significance.

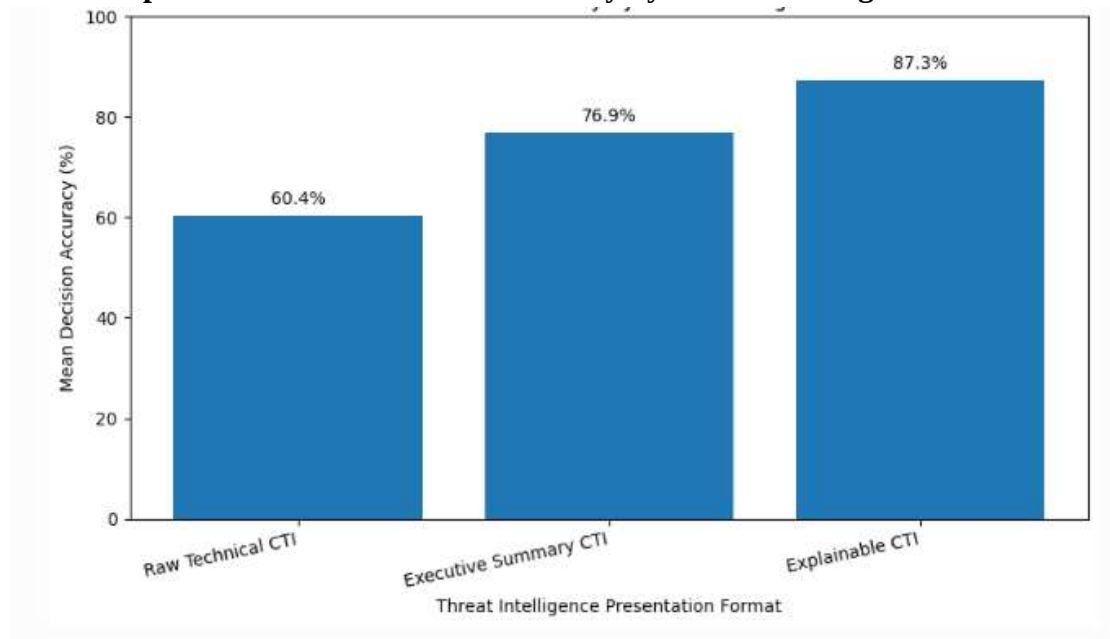
**Table 2: Simulated Decision Outcomes by Threat-Intelligence Presentation Format**

| CTI Presentation Format | Simulated n | Mean Decision Accuracy | Mean Confidence / 5 | Mean Decision Time |
|-------------------------|-------------|------------------------|---------------------|--------------------|
| Raw Technical CTI       | 90          | 60.4%                  | 3.12                | 12.96 min          |
| Executive Summary CTI   | 90          | 76.9%                  | 3.68                | 9.73 min           |
| Explainable CTI         | 90          | 87.3%                  | 4.16                | 7.22 min           |

**Source:** Author-generated synthetic dataset; total simulated n = 270. Values are illustrative and must not be represented as findings from real participants.

The simulated one-way analysis of variance for decision accuracy produced  $F = 44.75$ ,  $p < .001$ . Because the observations are synthetic, this value serves only to demonstrate the analytical procedure that could be applied in a controlled experiment.

**Graph 1: Simulated Decision Accuracy by Threat-Intelligence Format**



**Interpretation:** Decision accuracy increases progressively as the threat-intelligence presentation moves from specialist technical reporting toward decision-oriented explainability.



**Key Finding:** The simulation suggests that the greatest decision value is generated not merely by shortening technical information but by preserving important evidence while explicitly explaining relevance, uncertainty, consequences, and available actions.

## 5.2 Decision Confidence and Decision Time

Decision confidence also increased across the three simulated conditions. Raw technical CTI produced an average confidence score of 3.12, conventional executive summaries produced 3.68, and explainable CTI produced 4.16. The simulated ANOVA yielded  $F = 84.80$ ,  $p < .001$ .

This pattern does not mean that higher confidence is automatically desirable. Excessive confidence can itself become a governance risk. The objective of ETI should therefore be calibrated confidence: decision makers should feel more confident when evidence is strong and appropriately cautious when evidence is incomplete. This is why explicit uncertainty communication is included as a core ETI element.

Decision time decreased from approximately 12.96 minutes in the technical-report condition to 9.73 minutes in the executive-summary condition and 7.22 minutes in the explainable condition. The simulated ANOVA for decision time produced  $F = 189.47$ ,  $p < .001$ .

The reduction in time reflects a plausible advantage of structured explanations. Decision makers do not need to search through specialist information to determine which details are relevant. Instead, the intelligence product directs attention toward the elements that affect the decision while retaining access to the technical evidence for validation or escalation.

## 5.3 Interpretation of the Combined Results

The three outcomes should be interpreted together. A report that produces faster decisions but lower accuracy would not constitute effective decision support. Similarly, a report that increases confidence without improving understanding could encourage automation bias or unjustified certainty. In the present simulation, explainable CTI simultaneously increased decision accuracy and confidence while reducing decision time.

This pattern is conceptually consistent with broader explainability research. Ribeiro and colleagues demonstrated that explanations can support judgments about whether predictions should be trusted, while Lundberg and Lee emphasized the importance of understanding why complex models produce particular predictions. Within threat intelligence, Rastogi and colleagues similarly identify confidence information and interpretable reasoning as important requirements for supporting cyber analysts.

For nontechnical decision makers, however, explanations need to go beyond technical model interpretation. The central unit of explanation should be the decision, not the algorithm. A high-quality ETI product should enable the reader to answer six questions: What is the threat? Why does it matter here? What evidence supports the assessment? How certain are we? What could happen? What should we decide now?

## 6. Discussion

### 6.1 Explainability as a Translation Layer Between Security Operations and Governance

The findings support the conceptualization of explainable threat intelligence as a translation layer between specialist cybersecurity analysis and organizational governance. NIST's definition of CTI already emphasizes transformation, analysis, interpretation, and enrichment for decision-making, indicating that intelligence value arises from context rather than raw data alone. ETI extends this principle by making the intended decision maker an explicit design variable. Intelligence for a malware analyst and intelligence for a chief executive may originate from the same evidence but should not be presented through the same informational structure.

This does not imply that technical detail should be removed. Oversimplification can conceal uncertainty and create decisions that appear rational while lacking evidential support. A layered architecture is therefore preferable. The primary view should communicate the decision-relevant interpretation, while deeper layers should allow the reader, security analyst, auditor, or investigator to inspect supporting indicators, ATT&CK mappings, vulnerability details, source provenance, analytical assumptions, and alternative explanations. MITRE ATT&CK and standardized CTI formats such as STIX can provide the technical structure beneath this decision layer.

The governance value of this design is particularly important because NIST CSF 2.0 places cybersecurity risk within organizational governance and explicitly seeks to communicate outcomes to audiences ranging from practitioners to executives. ETI can therefore serve as an interface between DETECT and RESPOND activities conducted by security teams and GOVERN decisions involving risk ownership, resource allocation, escalation, and accountability.

### 6.2 Explanation Quality Should Be Measured Through Decisions, Not Presentation Aesthetics

The simulated comparison also highlights the importance of evaluating explanation systems against actual decision outcomes. Cybersecurity dashboards are frequently evaluated by technical performance, information coverage, interface usability, or user satisfaction. These metrics are valuable, but they do not directly establish whether an explanation improves decisions.

A more rigorous ETI evaluation should measure whether users select appropriate actions, how quickly they reach decisions, whether confidence corresponds to evidence quality, whether important uncertainty is recognized, and whether decision makers can later justify their reasoning. These outcome-oriented measures align with the broader logic of explainable AI, where explanations are intended to help humans understand and appropriately rely upon algorithmic outputs rather than merely expose technical complexity.

The distinction between an executive summary and ETI is particularly important. Summarization primarily reduces information volume. Explainability establishes relationships among evidence, threat interpretation, organizational impact, uncertainty, and recommended action. A short report can therefore remain unexplainable if it merely states that a threat is "critical." Conversely, a moderately detailed report can remain accessible if its logic is clearly organized.

A practical ETI statement might therefore communicate that a ransomware actor has recently used a particular access technique, that the organization exposes an asset consistent with the observed attack pathway, that evidence is assessed with moderate confidence because one telemetry source is

missing, that disruption of a named service is the most significant business consequence, and that two response options are available with different operational costs. Such an explanation allows the decision maker to understand both the recommendation and the basis on which it rests.

### 6.3 Human Oversight, Uncertainty, and Responsible Automation

Explainable threat intelligence becomes more important as AI and automated analytics are integrated into threat detection and prioritization. NIST's AI Risk Management Framework treats explainability and interpretability as part of a wider set of trustworthiness considerations that also includes transparency, accountability, validity, reliability, security, and resilience. This means that an explanation should not be treated as decorative justification added after an automated system has already reached a conclusion.

An ETI system should clearly distinguish observed evidence from machine inference and human analytical judgment. It should identify source provenance where disclosure is permissible, provide confidence or uncertainty information, indicate alternative plausible interpretations when relevant, and preserve an audit trail linking recommendations to underlying evidence. This approach can reduce the risk that decision makers interpret algorithmic outputs as unquestionable facts.

Human oversight should also be matched to consequence. Low-impact repetitive actions may justify greater automation, whereas decisions involving major service shutdowns, significant financial expenditure, public communication, regulatory reporting, employee consequences, or other substantial organizational effects should generally preserve meaningful human review. Explainability supports this review because it gives the responsible decision maker a basis for questioning, accepting, modifying, or rejecting a recommendation.

The most sustainable role for ETI is therefore neither to replace cyber analysts nor to require executives to become cyber analysts. Its function is to preserve the analytical depth of specialist intelligence while presenting the minimum sufficient explanation required for accountable decision-making. In this sense, explainability becomes an organizational coordination mechanism as much as a technical capability.

### 7. Conclusion

Cyber threat intelligence is valuable only when it informs an appropriate decision. Modern CTI infrastructures can collect, exchange, correlate, and analyze large volumes of information, while standards such as STIX and TAXII and knowledge bases such as MITRE ATT&CK improve the technical organization of threat data. However, the growing participation of executives and other nontechnical stakeholders in cyber-risk governance creates an additional requirement: intelligence must be understandable without being misleadingly simplified.

The present study develops Explainable Threat Intelligence as a decision-centered communication approach that connects technical evidence with organizational relevance, confidence, uncertainty, consequences, and response options. In the synthetic experiment, explainable CTI produced a simulated decision accuracy of 87.3%, compared with 76.9% for conventional executive summaries and 60.4% for raw technical CTI. Explainable CTI also generated higher simulated confidence and shorter decision time. These findings are not empirical claims about real decision makers; they illustrate testable hypotheses and provide a framework for subsequent controlled research.

The principal practical implication is that organizations should not evaluate their threat-intelligence capability solely by the number of feeds collected, indicators processed, alerts generated, or ATT&CK techniques mapped. The more consequential question is whether the intelligence enables the responsible person to determine what should be done. An effective ETI product should therefore explain the threat, affected organizational context, evidential basis, analytical confidence, important uncertainty, plausible consequences, response alternatives, and rationale for the recommended action. Technical detail should remain available through layered evidence rather than being removed from the intelligence chain.

Future empirical research should test the proposed framework with executives, public administrators, managers, compliance leaders, and other nontechnical decision makers using realistic cybersecurity scenarios. Studies should compare alternative explanation structures, examine confidence calibration, evaluate cognitive workload, investigate differences across professional roles, and determine when additional explanation begins to create information overload. Researchers should also examine whether AI-generated explanations improve or distort decision quality and whether users can distinguish factual evidence from probabilistic inference. Such work would move cyber threat intelligence beyond information delivery toward transparent, accountable, and genuinely decision-oriented cybersecurity governance.

## References

1. Moffat, N., & Crichton, M. (2015). Threat intelligence: Concepts and challenges. *Proceedings of the 2015 IEEE International Conference on Cyber Security and Cloud Computing*.
2. Tounsi, W., & Rais, H. (2018). A survey on technical threat intelligence in the age of sophisticated cyber attacks. *Computers & Security*, 72, 212–233.
3. Choo, K.-K. R. (2011). The cyber threat landscape: Challenges and future research directions. *Computers & Security*, 30(8), 719–731.
4. Sommer, R., & Paxson, V. (2010). Outside the closed world: On using machine learning for network intrusion detection. *2010 IEEE Symposium on Security and Privacy*, 305–316.
5. Buczak, A. L., & Guven, E. (2016). A survey of data mining and machine learning methods for cyber security intrusion detection. *IEEE Communications Surveys & Tutorials*, 18(2), 1153–1176.
6. Kotsiantis, S. B. (2007). Supervised machine learning: A review of classification techniques. *Informatica*, 31, 249–268.
7. Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). “Why should I trust you?” Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135–1144.
8. Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems*, 30.
9. Guidotti, R., Monreale, A., Ruggieri, S., Turini, F., Giannotti, F., & Pedreschi, D. (2018). A survey of methods for explaining black box models. *ACM Computing Surveys*, 51(5), Article 93.
10. Miller, T. (2019). Explanation in artificial intelligence: Insights from the social sciences. *Artificial Intelligence*, 267, 1–38.
11. Doshi-Velez, F., & Kim, B. (2017). Towards a rigorous science of interpretable machine learning. *arXiv preprint arXiv:1702.08608*.



12. Lipton, Z. C. (2018). The mythos of model interpretability. *Communications of the ACM*, 61(10), 36–43.
13. Adadi, A., & Berrada, M. (2018). Peeking inside the black-box: A survey on explainable artificial intelligence (XAI). *IEEE Access*, 6, 52138–52160.
14. Arrieta, A. B., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., García, S., Gil-López, S., Molina, D., Benjamins, R., Chatila, R., & Herrera, F. (2020). Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82–115.
15. Endsley, M. R. (1995). Toward a theory of situation awareness in dynamic systems. *Human Factors*, 37(1), 32–64.
16. Endsley, M. R. (2000). Theoretical underpinnings of situation awareness: A critical review. In *Situation Awareness Analysis and Measurement*. Lawrence Erlbaum Associates.
17. Parasuraman, R., Sheridan, T. B., & Wickens, C. D. (2000). A model for types and levels of human interaction with automation. *IEEE Transactions on Systems, Man, and Cybernetics—Part A: Systems and Humans*, 30(3), 286–297.
18. Dzindolet, M. T., Peterson, S. A., Pomranky, R. A., Pierce, L. G., & Beck, H. P. (2003). The role of trust in automation reliance. *International Journal of Human-Computer Studies*, 58(6), 697–718.
19. Sillaber, C., Sauerwein, C., Mussmann, A., & Breu, R. (2018). Data quality challenges and future research directions in threat intelligence sharing. *Journal of Cyber Security Technology*, 2(3–4), 167–200.
20. Wagner, T. D., Mahbub, K., Palomar, E., & Abdallah, A. E. (2019). Cyber threat intelligence sharing: Survey and research directions. *Computers & Security*, 87, 101589.