

Global Journal of Engineering and Technology Review

Volume: 02 Issue: 09 September 2026

Accelerating National Defense with Large Language Models: A Conceptual Framework for the Real-Time Processing of Shared Cyber Threat Indicators

Chukwunenye Amadi

Independent Researcher, USA

Abolaji Adebayo

Computacenter, USA

Ayokunle Olamide Ijagbemi

Independent Researcher, USA

Published Date: September 05, 2026**Article DOI : 10.65150/EP-gjetr/V2E9/2026-04**

ABSTRACT: National defense increasingly depends on the capacity to detect, interpret, and act upon vast streams of shared cyber threat intelligence in real time. Yet defense and security organizations are overwhelmed by unstructured, redundant, and fragmented threat data drawn from heterogeneous global sources, a condition that undermines timely and accurate analysis. This paper develops a conceptual framework for applying Large Language Models (LLMs) to the real-time processing of shared threat indicators within national defense intelligence systems. Grounded in Information Processing Theory, Sociotechnical Systems Theory, and Signal Detection Theory, the framework positions LLMs as advanced cognitive processors that ingest, contextualize, and prioritize language-based threat data while preserving human judgment and accountability. This paper reports no new experimental results. It advances a conceptual argument, supported by an illustrative and explicitly non-experimental comparison, that transformer-based semantic processing is better suited in principle than conventional keyword-matching pipelines to the interpretive demands of heterogeneous and redundantly reported threat intelligence. The comparison is expository in purpose and is offered to generate testable hypotheses rather than to establish validated performance claims. The paper contributes an integrated, governance-aware architecture for intelligent defense systems and articulates the ethical, organizational, and policy conditions under which such systems can be responsibly deployed. The framework yields the testable propositions that LLM-driven semantic processing may reduce analyst cognitive burden, shorten response cycles, and strengthen national cybersecurity posture. Each of these propositions remains to be established through instrumented empirical evaluation, which the paper identifies as the immediate priority for future work.

KEYWORDS: Large Language Models, national defense, cyber threat intelligence, information overload, situational awareness, responsible AI.

I. INTRODUCTION

Modern national defense is inseparable from the security of the digital infrastructure on which military, governmental, and critical civilian functions now depend. Adversaries operating across cyberspace exploit shared vulnerabilities at machine speed, and the defensive community has responded by institutionalizing the collective exchange of cyber threat intelligence among allied nations, sector-specific coordinating bodies, and commercial security vendors (Skopik et al., 2016). This collaborative posture rests on a straightforward premise: an indicator of compromise observed by one participant can inoculate many others if it is disseminated quickly and understood correctly. The practice of technical threat intelligence sharing has consequently matured from ad hoc email exchanges into a structured ecosystem of feeds, standards, and platforms (Tounsi & Rais, 2018). Yet the very success of this ecosystem has generated a second-order problem. The aggregate volume of shared indicators now arrives faster than human defenders can triage, contextualize, and act upon it, so that the intelligence intended to accelerate defense can paradoxically decelerate it. This paper argues that the emerging generation of large language models offers a principled means of restoring the tempo advantage to defenders. It develops a conceptual framework in which these models serve as real-time interpreters of shared indicators, bridging the gap between the machine speed of attack and the cognitive limits of the analysts charged with response.

The difficulty is fundamentally one of cognitive economics rather than of data availability. Decades of research on human information processing establish that working memory and attention impose narrow bounds on the quantity of discrete items an individual can hold and manipulate simultaneously (Miller, 1956). The transfer of incoming stimuli into durable, actionable knowledge is mediated by structured memory stages whose throughput cannot be arbitrarily expanded (Atkinson & Shiffrin, 1968). When the influx of threat indicators exceeds these bounds, analysts resort to heuristics, defer low-salience items, and accumulate backlogs, with the result that genuinely urgent signals are delayed or discarded alongside noise. Empirical studies of operational threat-intelligence programs corroborate this picture, documenting analyst fatigue, alert saturation, and uneven exploitation of shared feeds (Sauerwein et al., 2017). The organizational response has typically been to hire additional staff

License: This is an open access article under the CC BY 4.0 license: <https://creativecommons.org/licenses/by/4.0/>

or to layer additional filtering rules, but neither strategy scales with an adversary population that automates its own tooling. What is required is a mechanism capable of absorbing the interpretive burden at the same velocity and volume as the indicators themselves, while preserving the semantic richness that rigid rule-based filters discard. The remainder of this paper contends that language models, properly framed within established theories of cognition and organization, constitute precisely such a mechanism.

The technical basis for this contention is the rapid advance of transformer-based language models over the past several years. The introduction of the attention mechanism enabled models to represent long-range dependencies in sequential data without the bottlenecks of recurrence (Vaswani et al., 2017), and subsequent scaling of these architectures across parameters and data yielded systems capable of performing a broad range of language tasks from natural instructions and a handful of examples (Brown et al., 2020). These systems, now widely characterized as foundation models, exhibit a generality that distinguishes them from the narrow, single-purpose classifiers that dominated earlier applications of machine learning to security (Bommasani et al., 2021). A single pretrained model can summarize a prose incident report, extract structured fields from a malware advisory, normalize inconsistent terminology, and answer analyst queries about relationships among observed artifacts, all without task-specific engineering. This flexibility is decisive for threat intelligence, whose source material is heterogeneous, partly unstructured, and continuously evolving. The framework advanced here treats the language model not as an oracle that replaces the analyst but as a high-throughput interpretive layer that renders the flood of shared indicators tractable to human judgment. Establishing that role rigorously requires connecting the capabilities of these models to the specific pathologies of contemporary intelligence sharing and to the theoretical constructs that explain why those pathologies arise.

The application of computational methods to cyber defense is not itself new, and the present framework builds deliberately on that lineage. Machine-learning techniques have been studied extensively for intrusion detection, where surveys catalog the trade-offs among classifier families, feature representations, and the persistent challenge of skewed and shifting data distributions (Buczak & Guven, 2016). Deep neural architectures subsequently extended these efforts, offering learned representations that reduced dependence on hand-crafted features across malware analysis, traffic classification, and anomaly detection (Berman et al., 2019). Parallel work has sought to situate such techniques within a coherent discipline of cybersecurity data science, emphasizing data quality, model interpretability, and the operational context of deployment (Sarker et al., 2020). More recently, researchers have turned specifically to the mining of cyber threat intelligence for proactive defense, surveying methods that transform raw reports and feeds into actionable knowledge (Sun et al., 2023). The distinguishing contribution of the language-model paradigm relative to this body of work lies in its capacity to operate directly on natural-language and semi-structured intelligence, to generalize across tasks without bespoke retraining, and to expose its reasoning in a form that human analysts can interrogate. This paper positions its framework as a natural successor to prior applied efforts rather than a departure from them, extending their trajectory toward real-time interpretation.

The contribution of this work is conceptual and integrative rather than empirical. It advances a framework in which large language models mediate the real-time processing of shared cyber threat indicators, and it grounds that framework in both the technical properties of the models and the human and organizational theories that explain why unassisted processing fails at scale. The paper is organized as follows. Section II characterizes the contemporary threat-intelligence environment, describing the sources and standards of shared indicators, the operational consequences of information overload, and the structural problems of fragmentation, duplication, and semantic drift. Section III reviews the foundations and capabilities of large language models, tracing the path from the transformer architecture and large-scale pretraining through contextual representation to alignment and reasoning. Section IV develops the theoretical foundations of the framework, drawing on information processing theory, sociotechnical systems theory, signal detection theory, and situation awareness. Section V introduces the related-work survey that situates the proposal within applied cyber defense and cross-sector intelligent systems. Together these sections establish the conceptual scaffolding on which the second half of the paper constructs its detailed architectural and evaluative proposals, and they clarify the assumptions under which language-model mediation can be expected to succeed and the boundaries within which its outputs remain subject to human authority.

Relationship to Prior Work and Statement of Original Contribution

In the interest of full transparency, the authors disclose that this manuscript builds upon an earlier working paper by the lead author, circulated as Amadi (2025), which addressed the same operational problem of semantic correlation and de-duplication of shared cyber threat indicators. That earlier paper is cited throughout where its framing is carried forward, and the present manuscript should be read as a successor to it rather than as an independent treatment. The two works share a problem statement, namely that the volume and redundancy of shared threat intelligence outpace human interpretive capacity, and they share the general proposition that language-model representations offer a route to meaning-level correlation that lexical matching cannot supply. Portions of the abstract and the framing of the information-overload problem overlap in substance with that earlier text, and the present abstract has been revised to remove wording carried over from it.

The two works differ in kind, and the distinction matters for how the present contribution should be assessed. The earlier paper was presented as the design and evaluation of an operational system and reported comparative performance in favour of the language-model approach over keyword matching. The present manuscript makes no such performance claim. It is a conceptual and design-science contribution whose evaluation is explicitly illustrative, and one of its purposes is to withdraw and replace the earlier claim of demonstrated empirical superiority, which the evidence available to the authors does not support at the standard the claim implied. Readers should treat the comparative statements in this paper as expository reasoning about mechanism, not as measured results.

Against that background, the substantive additions the present manuscript makes are the following. First, it supplies a theoretical grounding that the earlier work lacked, deriving the interpretive bottleneck from information processing theory and bounded rationality, the division of labour between analyst and model from sociotechnical systems theory, and the triage decision itself from signal detection theory and the levels of situation awareness, so that the architecture is answerable to explanatory constructs rather than to intuition. Second, it introduces a retrieval-grounded reasoning layer that conditions every synthesized assessment on citable passages of validated intelligence, which is absent from the earlier design and which converts explicability from an aspiration into an architectural property. Third, it develops a governance and ethics analysis that translates established principles of responsible artificial intelligence into functional specifications enforced by the architecture, together with an account of the institutional practices on which responsible deployment depends. Fourth, it states its limitations and threats to validity explicitly, marking the

boundary between what has been argued and what remains to be demonstrated. Fifth, it situates the proposal within the applied literature on cyber defence and adjacent intelligent systems, clarifying what is inherited and what is new. The contribution of this work is therefore the disciplined reconstruction of a previously asserted capability as a falsifiable, governance-aware design, not the demonstration of that capability.

II. THE CONTEMPORARY THREAT-INTELLIGENCE ENVIRONMENT

A. Sources and Types of Shared Threat Indicators

Shared cyber threat indicators originate from a diverse constellation of producers, including government computer emergency response teams, information sharing and analysis centers organized by economic sector, commercial threat-intelligence vendors, open-source research communities, and the internal telemetry of individual enterprises. To render observations from these heterogeneous sources mutually intelligible, the community has converged on structured representations, most prominently the Structured Threat Information Expression standard, which provides a common vocabulary for encoding indicators, adversary tactics, campaigns, and their interrelationships (Barnum, 2014). Complementing this representation is a dedicated transport protocol that governs the automated exchange of such structured intelligence among trusted parties, enabling machine-to-machine dissemination at scale (Connolly et al., 2014). Beyond formal standards, collaborative platforms have emerged to support the collection, correlation, and redistribution of indicators within and across trust communities, offering practical infrastructure for members to contribute and consume shared observations (Wagner et al., 2016). The types of indicators exchanged span a wide spectrum, from low-level atomic artifacts such as file hashes, network addresses, and domain names to higher-order behavioral descriptors that characterize adversary methods. This layering reflects a recognition that the defensive value of an observation depends heavily on the level of abstraction at which it is expressed, and that atomic indicators alone provide brittle and short-lived protection against a determined and adaptive adversary.

The interpretive value of shared indicators is amplified when they are situated within models of adversary behavior. The intrusion kill chain frames an intrusion as a sequence of stages, from reconnaissance through exploitation to action on objectives, and thereby allows individual indicators to be understood as evidence of an attacker's progression rather than as isolated facts (Hutchins et al., 2011). A complementary and now widely adopted knowledge base catalogs the concrete tactics and techniques that adversaries employ across these stages, providing a shared reference against which observed behavior can be classified and communicated (Strom et al., 2018). These behavioral frameworks matter because they shift the emphasis of intelligence from ephemeral artifacts, which adversaries can trivially rotate, toward durable descriptions of method that retain defensive utility over longer horizons. A comprehensive survey of technical threat intelligence situates these standards, platforms, and behavioral models within a broader taxonomy and highlights the persistent tension between the machine-readable precision required for automation and the contextual nuance required for human decision-making (Tounsi & Rais, 2018). This tension is central to the present work. Structured standards make indicators exchangeable, and behavioral models make them meaningful, but neither, on its own, resolves the burden of interpreting the resulting volume in the time available. It is precisely at this interpretive juncture that a language-model layer promises to contribute.

B. Information Overload and Its Operational Consequences

The proliferation of sources and standards that makes collective defense possible also produces an aggregate stream whose magnitude exceeds any plausible bound on human throughput. A large enterprise subscribing to multiple feeds may receive hundreds of thousands of discrete indicators daily, each nominally requiring assessment for relevance, credibility, and priority. The constraint is not storage or transport, both of which scale cheaply, but interpretation, which remains bound by the cognitive limits documented in the psychological literature (Miller, 1956). Studies of operational threat-intelligence programs report that analysts frequently cannot consume the feeds to which their organizations subscribe, that substantial fractions of received indicators are never examined, and that the perceived quality and actionability of shared data vary widely across sources (Sauerwein et al., 2017). The consequence is a defensive posture that is nominally well-informed but practically saturated, in which the marginal indicator adds to the backlog rather than to protection. This condition of overload is qualitatively different from a simple shortage of staff, because adding analysts increases coordination cost and does not address the underlying mismatch between the velocity of automated production and the pace of human sense-making. The problem therefore invites a solution that operates on the interpretive bottleneck directly, compressing and prioritizing the stream before it reaches the analyst rather than merely enlarging the workforce.

The operational consequences of this overload extend beyond individual analyst fatigue to the effectiveness of collective defense as a system. Research on cyber threat intelligence sharing identifies a recurring gap between the availability of shared data and its actual exploitation, noting that organizational, technical, and trust-related barriers frequently prevent received intelligence from influencing defensive action (Wagner et al., 2019). When indicators arrive faster than they can be triaged, the latency between disclosure and defensive response widens, and this latency is precisely the window an adversary exploits to propagate across the many organizations that received but did not act upon a warning. Overload thus undermines the fundamental economic argument for sharing, which presumes that timely dissemination converts a single victim's misfortune into the wider community's immunity. Moreover, saturation degrades judgment quality in subtler ways: under time pressure and high volume, analysts systematically favor easily assessed atomic indicators over the more consequential but interpretively demanding behavioral intelligence, inverting the priority that the behavioral frameworks were designed to encourage. The net effect is that the community accumulates ever more intelligence while extracting diminishing defensive value from each increment. A processing layer capable of sustaining interpretive throughput at the volume and velocity of production would restore the latent value of shared indicators and reconnect dissemination to defensive action, which is the outcome the proposed framework seeks.

C. Fragmentation, Duplication, and Semantic Drift

Beyond sheer volume, the shared-intelligence ecosystem suffers from structural heterogeneity that further complicates interpretation. Because indicators originate from many independent producers, the same underlying threat is frequently reported multiple times in slightly different forms, with divergent naming conventions, inconsistent field population, and varying levels of confidence and context. This fragmentation produces substantial duplication, so that a defender confronting what appears to be a large set of distinct warnings may in fact be facing a much smaller number of underlying events described redundantly across sources. Research on collective cyber defense through information sharing emphasizes that the value of a sharing community depends not merely on the quantity of contributions but on the community's ability to correlate and de-

conflict them into a coherent operational picture (Skopik et al., 2016). The reconciliation task is complicated by the prevalence of unstructured and semi-structured source material, since a large proportion of actionable intelligence is first published as prose advisories, blog posts, and reports rather than in fully structured form. Reviews of methods for mining threat intelligence from unstructured text document the difficulty of reliably extracting entities and relationships from such material, and the corresponding need for techniques that tolerate linguistic variability while preserving semantic fidelity (Rahman et al., 2020). Fragmentation, in short, imposes an integration cost that grows with the number of participating sources.

A more insidious problem than duplication is semantic drift, the gradual divergence in the meaning attached to shared terms as they propagate across communities, tools, and time. Adversary group names, malware families, and technique labels are frequently assigned independently by different vendors, so that a single campaign may accumulate several incompatible designations, while a single label may be applied inconsistently to distinct phenomena. As indicators are re-shared, enriched, and re-encoded, contextual qualifications are often stripped away, leaving atomic values detached from the conditions under which they were meaningful. Surveys of methods for mining cyber threat intelligence to support proactive defense highlight normalization and disambiguation as persistent obstacles to automated exploitation, precisely because the surface form of an indicator underdetermines its operational significance (Sun et al., 2023). Semantic drift is corrosive because it defeats naive de-duplication and correlation, which depend on stable identity, and because it silently degrades the confidence that downstream consumers can place in aggregated intelligence. Addressing it requires more than string matching; it requires an interpretive capacity that can recognize when differently expressed observations refer to the same underlying entity and when identically labeled observations do not. This capacity for meaning-preserving normalization across linguistic and structural variation is a natural strength of contextual language models, and it motivates the technical review that follows.

III. LARGE LANGUAGE MODELS: FOUNDATIONS AND CAPABILITIES

A. The Transformer Architecture and Pretraining

The capabilities invoked by the proposed framework rest on a specific architectural innovation. The transformer dispenses with the sequential recurrence that previously dominated neural sequence modeling and instead relies entirely on self-attention, a mechanism that computes, for each element of a sequence, a weighted combination of all other elements according to learned measures of relevance (Vaswani et al., 2017). This design confers two decisive advantages for language processing. First, it captures dependencies between distant tokens directly, without the signal degradation that afflicts recurrent models over long spans, which is essential for interpreting the extended and reference-laden prose typical of threat advisories. Second, because attention over a sequence can be computed in parallel rather than step by step, the architecture is highly amenable to acceleration on modern hardware, enabling training on corpora far larger than earlier methods could accommodate. This parallelism proved to be the enabling condition for scale. Empirical investigation of the relationship between model size, dataset size, and computational budget revealed smooth and predictable improvements in language-modeling performance as these quantities increase together, a regularity that guided the deliberate construction of progressively larger systems (Kaplan et al., 2020). The transformer thus provided not only a superior mechanism for modeling language but a substrate whose returns to scale could be exploited systematically, setting the trajectory that produced today's foundation models.

Pretraining is the second pillar of the paradigm. Rather than training a model from scratch for each task, contemporary systems are first exposed to vast quantities of unlabeled text through self-supervised objectives that require no human annotation. One influential line of work demonstrated that a transformer trained purely to predict the next token across a large and diverse corpus acquires, as a byproduct, the ability to perform many language tasks without explicit supervision, revealing that broad linguistic and even factual competence can emerge from a single generative objective (Radford et al., 2019). This self-supervised pretraining is attractive for the threat-intelligence domain because labeled security data are scarce, expensive, and quickly outdated, whereas unlabeled text describing systems, vulnerabilities, and adversary behavior is abundant. A model pretrained on such material internalizes the vocabulary and discourse patterns of the domain and can subsequently be specialized with comparatively little task-specific data. The separation of a costly, general pretraining phase from a cheap, targeted adaptation phase is what makes the economics of the approach favorable for defenders, who cannot afford to engineer a bespoke system for every interpretive task the flood of indicators presents. The architecture and the pretraining regime together yield a single, reusable interpretive asset, and the subsequent subsections examine how that asset represents meaning and how its behavior is shaped toward reliable use.

B. From Representation to Contextual Understanding

The interpretive power relevant to threat intelligence derives from how these models represent meaning. Whereas earlier approaches assigned each word a fixed vector regardless of context, bidirectional transformer models learn representations in which the meaning assigned to a token depends jointly on the tokens that precede and follow it, produced by a pretraining objective that masks portions of the input and requires the model to reconstruct them (Devlin et al., 2019). Such contextual representation is precisely what the disambiguation of security intelligence demands, since the significance of an address, a filename, or a technique label is determined by its surrounding context rather than by its surface form alone. Refinements to this pretraining recipe demonstrated that careful attention to training data volume, batch construction, and optimization substantially improves the quality of the learned representations without any change to the underlying architecture, underscoring how sensitive downstream competence is to the pretraining regime (Liu et al., 2019). For the framework proposed here, the salient consequence is that a contextual encoder can recognize when two differently worded descriptions denote the same underlying entity and when superficially similar expressions diverge in meaning, directly addressing the semantic drift and fragmentation identified earlier. The representation is not a lookup table of terms but a learned space in which relationships of similarity and difference reflect usage, which is the property that meaning-preserving normalization requires.

A further generalization reframed disparate language tasks within a single unified format, casting classification, extraction, summarization, and question answering alike as transformations from an input text to an output text (Raffel et al., 2020). This text-to-text formulation is particularly well suited to the interpretation of shared indicators, whose processing entails exactly such transformations: converting a prose advisory into

structured fields, condensing a lengthy report into an analyst-ready summary, or answering a targeted question about the relationships among observed artifacts. Because a single model handles all of these under one interface, an operational pipeline need not integrate and maintain a proliferation of specialized components. The capacity to perform many such tasks from natural-language instruction and only a few demonstrations was demonstrated at scale, showing that sufficiently large models generalize to novel tasks specified purely in context, without gradient updates (Brown et al., 2020). This in-context adaptability is valuable in a domain where the interpretive tasks themselves shift as adversary behavior and reporting conventions evolve, since new processing needs can be expressed as instructions rather than requiring retraining. Together, contextual representation and unified, instruction-driven task formulation supply the interpretive flexibility that the heterogeneous and drifting stream of shared intelligence demands, while consolidating that capability within a single maintainable model.

C. Alignment, Reasoning, and Reliability

Raw generative capability is necessary but not sufficient for a system entrusted with defensive interpretation; the model's behavior must be aligned with the intentions of its operators. A pretrained model optimized solely to predict text does not reliably follow instructions or respect the priorities of a human user, and may produce fluent but unhelpful or inappropriate output. Alignment techniques address this gap by fine-tuning the model on demonstrations of desired behavior and then further optimizing it against a reward model learned from human preferences, yielding systems that follow instructions more faithfully and that better reflect the operator's intent (Ouyang et al., 2022). For threat-intelligence applications, alignment is not a cosmetic refinement but a functional requirement, because the value of an interpretive layer depends on its adherence to the analyst's actual priorities, its willingness to express uncertainty, and its restraint in the face of ambiguous evidence. The broader literature on foundation models emphasizes that such capability and alignment are accompanied by risks, including the propagation of biases latent in training data, brittleness under distribution shift, and the potential for confident but incorrect assertions (Bommasani et al., 2021). A responsible framework must therefore treat alignment and its attendant hazards as first-class design concerns rather than afterthoughts, a stance the second half of this paper adopts explicitly in its architectural proposals and evaluation criteria.

Reliability in interpretation also depends on the model's capacity for structured reasoning, not merely pattern completion. Prompting a sufficiently capable model to articulate intermediate reasoning steps before committing to an answer has been shown to substantially improve performance on tasks that require multi-step inference, a technique that renders the model's problem-solving partly explicit (Wei et al., 2022). This property carries direct significance for threat analysis, where correctly assessing an indicator frequently requires chaining several inferences, relating an observed artifact to a technique, the technique to a stage of an intrusion, and that stage to a plausible adversary objective. Explicit intermediate reasoning offers a further benefit beyond accuracy: it produces a trace that a human analyst can inspect and contest, which is indispensable in a defensive setting where accountability and the calibration of trust are paramount. Yet the same literature cautions that articulated reasoning can be persuasive without being correct, so that a plausible chain of steps may lend unwarranted confidence to a flawed conclusion. The framework advanced here consequently positions the language model as a reasoning aid whose outputs remain subject to human verification, rather than as an autonomous decision-maker. Establishing where that boundary should fall requires theoretical grounding in how humans and organizations process information and detect signals, which the following section supplies.

IV. THEORETICAL FOUNDATIONS

A. Information Processing Theory

The claim that language models can relieve a defensive bottleneck presupposes a theory of where that bottleneck lies, and information processing theory supplies it. The classical model of human memory describes the flow of information through a sequence of stores, from a brief sensory register through a capacity-limited short-term store to a more durable long-term store, with transfer between stages governed by control processes such as rehearsal and attention (Atkinson & Shiffrin, 1968). The critical constraint for present purposes is the limited capacity of the short-term store, whose bound on the number of discrete chunks that can be held and manipulated simultaneously establishes a hard ceiling on the rate at which incoming information can be consciously processed (Miller, 1956). Threat-intelligence triage is an information-processing task of exactly this kind, in which each indicator competes for a share of a fixed and modest cognitive budget. When the arrival rate of indicators exceeds the analyst's processing rate, items necessarily queue, and the backlog grows without bound regardless of analyst diligence. The theory thus reframes information overload not as a failure of effort or organization but as a structural consequence of a mismatch between input rate and processing capacity, and it identifies the transfer of interpretive load away from the capacity-limited human channel as the appropriate remedy.

Information processing theory also clarifies the mechanism by which an automated interpretive layer can help, which is by performing the chunking and pre-attentive filtering that would otherwise consume scarce human capacity. If a language model condenses many redundant indicators into a single coherent summary, normalizes divergent terminology to a common referent, and surfaces only the items whose novelty or severity warrants human attention, it effectively enlarges the throughput of the human-machine system without violating the individual analyst's fixed cognitive bound. The theory of communication provides a complementary framing, treating the problem as one of transmitting information reliably over a channel of limited capacity in the presence of noise (Shannon, 1948). Under this framing, duplication and semantic drift function as noise that consumes channel capacity without conveying additional defensive value, and the interpretive layer acts to increase the effective signal-to-noise ratio delivered to the analyst. The consequence of these constraints for decision quality is captured by the principle of bounded rationality, which holds that decision-makers facing limited information and limited processing capacity do not optimize but satisfice, selecting the first alternative that meets an acceptable threshold (Simon, 1955). An interpretive layer that reduces the informational burden of each decision therefore does not merely accelerate triage; it enlarges the space of alternatives an analyst can realistically consider, improving the substantive quality of defensive choices.

B. Sociotechnical Systems Theory

A framework that introduces automated interpretation into an operational defensive setting must account for more than individual cognition, because threat analysis is embedded in organizations with their own structures, roles, and norms. Sociotechnical systems theory holds that the technical and social subsystems of a work system are jointly optimized only when they are designed in concert, and that improvements to the

technical subsystem which disregard the social subsystem tend to fail or to produce unintended dysfunction (Trist & Bamforth, 1951). The foundational studies underlying this theory observed that introducing a technically superior method of work degraded overall performance when it disrupted the established social organization of the workers, and that restoring effectiveness required designing the technical and human elements together. Applied to threat intelligence, the lesson is that inserting a language model into an analysis workflow is not a purely technical act; it reconfigures the division of labor, the locus of judgment, and the accountability structure of the defensive team. A model that assumes tasks previously performed by analysts alters their roles, their skill requirements, and their relationship to the intelligence they consume. Ignoring these effects risks producing a technically capable system that operators distrust, circumvent, or misuse, thereby forfeiting the very throughput gains the technology was adopted to secure.

The sociotechnical perspective therefore imposes concrete design constraints on any language-model-mediated processing framework. It requires that the allocation of tasks between model and analyst be treated as a deliberate joint-design decision rather than a residual outcome of what the technology happens to automate, and that the human retain meaningful authority over consequential judgments so that responsibility remains clearly located. It further requires that the system's outputs be intelligible to the analysts who must act on them, since a technical subsystem whose reasoning is opaque cannot be effectively integrated with the social subsystem that bears accountability for defensive decisions. This principle connects directly to the earlier discussion of explicit reasoning and alignment, which are the technical properties that make joint optimization achievable. The theory also cautions against the assumption that efficiency gains transfer automatically to the organization: the redistribution of interpretive labor may concentrate residual, high-difficulty cases on human analysts, changing the character and cognitive demand of their remaining work in ways that must be anticipated and supported. A well-designed framework consequently attends not only to the model's accuracy but to the roles, trust relationships, and feedback pathways through which analysts and models jointly produce defensive outcomes, treating the human-machine ensemble, rather than the model in isolation, as the unit of design and evaluation.

C. Signal Detection Theory and Situation Awareness

The triage of threat indicators is at its core a detection problem under uncertainty, and signal detection theory provides the formal apparatus for reasoning about it. The theory distinguishes an observer's sensitivity, the capacity to discriminate genuine signals from noise, from the decision criterion, the threshold of evidence at which the observer chooses to respond, and it shows that these two quantities can be manipulated independently (Green & Swets, 1966). This distinction is directly applicable to the interpretation of shared indicators, where each triage decision trades off the risk of missing a true threat against the cost of responding to a false one. Under overload, analysts implicitly shift their decision criterion toward conservatism simply to manage volume, accepting more missed detections in exchange for a manageable workload, a trade-off that the theory makes explicit and measurable. An interpretive layer can improve outcomes along either dimension: by enriching and disambiguating indicators it can raise the effective sensitivity of the human-machine system, and by absorbing volume it can relieve the pressure that forces an unfavorable criterion shift. Framing the contribution of a language model in these terms yields principled evaluation metrics, since the value of the system can be assessed not merely by aggregate accuracy but by its effect on the sensitivity and criterion of the overall detection process.

Detection, however, is only the entry point to the higher-order cognitive achievement that defensive decision-making ultimately requires, namely situation awareness. Situation awareness is characterized as comprising the perception of relevant elements in the environment, the comprehension of their meaning in relation to one another, and the projection of their future status (Endsley, 1995). Threat-intelligence analysis maps onto these three levels directly: perceiving the individual indicators, comprehending how they combine to indicate an adversary campaign, and projecting the likely next actions the defender must forestall. The pathologies described earlier attack situation awareness at each level, as overload impairs perception, fragmentation and semantic drift impede comprehension, and the resulting backlog forestalls projection. A language-model interpretive layer can be understood as an instrument for restoring situation awareness across all three levels, filtering and normalizing to support reliable perception, correlating and relating indicators to support comprehension, and, through structured reasoning over adversary models, contributing to projection. This framing unifies the theoretical strands of the section: information processing theory explains the capacity constraint, sociotechnical theory governs the human-machine division of labor, signal detection theory formalizes the triage decision, and situation awareness specifies the cognitive end state the whole system exists to produce. Together they establish the criteria against which the framework's later design must be judged.

V. RELATED WORK: APPLIED CYBER DEFENCE AND CROSS-SECTOR INTELLIGENT SYSTEMS

Having established the operational problem, the technical means, and the theoretical foundations, this section turns to the body of applied research within which the proposed framework must be situated. The purpose of the survey that follows is to trace how intelligent computational methods have been brought to bear on cyber defense and on analogous high-tempo, high-stakes domains in other sectors, and thereby to clarify both the precedents on which the framework draws and the gaps it seeks to fill. The relevant literature spans several partly overlapping strands, including the application of machine learning and deep learning to intrusion detection and malware analysis, the emerging discipline of cybersecurity data science, the mining of structured and unstructured threat intelligence for proactive defense, and the deployment of language and decision-support technologies in adjacent fields that share the defining characteristics of time pressure, uncertainty, and information abundance. Reviewing these strands together serves to demonstrate that the language-model-mediated processing of shared indicators is a natural extension of established trajectories rather than an isolated proposal, while also exposing the specific respects in which prior approaches fall short of the real-time, meaning-preserving interpretation that collective defense now demands. The detailed treatment of individual contributions follows.

The framework advanced in this paper is positioned against a broad and rapidly growing body of applied scholarship. Three strands are especially pertinent: research on cyber defence and threat intelligence, which defines the operational problem; research on artificial intelligence and secure digital systems, which supplies the analytical machinery; and a wider literature that applies data-driven modeling, risk governance, and intelligent decision support across other critical-infrastructure and enterprise sectors, which demonstrates the breadth and transferability of the methods on which real-time semantic processing draws. The remainder of this section surveys these strands in turn.

A mature stream of applied research addresses the detection, prioritization, and containment of evolving threats in enterprise and national settings. This work develops maturity models, threat-actor analysis, incident-response and security-orchestration frameworks, and threat-informed defence engineering that together define the operational requirements the present study seeks to accelerate: rapid triage of high-volume indicators, measurable control effectiveness, and coordinated response across organizational boundaries (Adebayo et al., 2023b; Akeju et al., 2018; Dosunmu & Ogundele, 2021, 2022, 2023, 2024a, 2024b, 2024d, 2025a, 2025c; Mbonu et al., 2019b, 2021c; Odejobi et al., 2025; Ogbole et al., 2025; Smart, 2020; Smart & Jooda, 2023a, 2023b, 2025, 2026).

Complementary studies examine the analytical capabilities that artificial intelligence and machine learning bring to security and decision support, including predictive analytics, human-in-the-loop annotation, generative models, edge intelligence, and interpretable learning. Collectively they establish both the promise of automated inference over large, noisy datasets and the necessity of retaining human oversight where consequences are severe (Adeyelu, 2020; Adeyelu & Dagodzo, 2024b; Annan, 2024; Borah et al., 2022; Dagodzo et al., 2022; Eboh & Aliliele, 2025b; Elebe et al., 2023; Essandoh et al., 2025; Komi, 2023; Komi & Adeniji, 2023; Komi & Ganiu, 2023; Ladapo et al., 2022a, 2024b, 2026; Mayo et al., 2021; Mbonu et al., 2021b; Okoruwa et al., 2025; Tonoyan et al., 2021a).

A substantial governance literature specifies the control, oversight, and accountability conditions under which intelligent systems must operate, spanning regulatory compliance, data-protection and privacy engineering, identity and access management, continuous auditing, and audit-automation. These contributions are directly relevant to fielding language models in defence, where transparency, provenance, and role-based control are prerequisites rather than afterthoughts (Adeyelu & Dagodzo, 2024a; Aliliele et al., 2023b, 2024a, 2025a, 2025b; Annan, 2022, 2025a, 2025b; Arumosoye & Obriki, 2020; Badmus et al., 2019b, 2022a, 2024, 2025; Bello et al., 2024; Dosunmu & Ogundele, 2025b; Edivri et al., 2026; Ekechi et al., 2026; Eyetsemitan et al., 2020, 2021, 2025; Fadayomi et al., 2021; Jooda & Smart, 2024b; Kumuyi et al., 2023, 2024a; Mbonu et al., 2018, 2019a, 2020a, 2020c, 2022a, 2022b, 2022c; Obogo et al., 2021c, 2022; Obriki & Arumosoye, 2024a; Ogbole et al., 2021b; Okonkwo et al., 2019; Okoruwa et al., 2020; Smart & Jooda, 2024b).

Work on cloud engineering, secure software delivery, and enterprise integration informs the engineering foundations required to deploy language-based analytics reliably at scale, addressing configuration governance, integration patterns, secure testing, and the operational discipline needed to keep data pipelines dependable under load (Aliliele et al., 2023a).

Beyond the security domain, research on supply-chain optimization, procurement, and logistics demonstrates closely analogous data-driven approaches to coordinating, de-duplicating, and governing information across distributed, multi-stakeholder operations. The problems of reconciling heterogeneous records, eliminating redundant entries, and maintaining a single authoritative view mirror those confronted in threat-intelligence consolidation (Dosunmu & Ogundele, 2019, 2024c; Okoruwa et al., 2024; Oteri & Edivri, 2026).

Studies of industrial and occupational safety contribute mature models of hazard detection, near-miss signal analysis, and operational risk governance. Their central concern, distinguishing weak leading signals from background noise before an incident occurs, directly parallels the signal-versus-noise framing that motivates semantic filtering and prioritization in this study (Obogo et al., 2020a, 2020b).

Contributions addressing geospatial intelligence, UAV and remote-sensing systems, and aviation-safety oversight demonstrate how heterogeneous sensor and spatial data are fused into actionable operational intelligence and how regulatory oversight is exercised over autonomous platforms, a fusion-and-oversight challenge closely related to correlating multi-source threat indicators (Adeyelu, 2023; Dagodzo & Ahiaeké Patrick, 2022; Smart & Jooda, 2024a).

A body of work on financial analytics, fraud detection, and business-process optimization shows how analytical models detect anomalies, reduce redundancy, and streamline decision workflows in high-stakes, data-intensive settings, providing methodological parallels for the anomaly-sensitive, redundancy-averse processing pursued in defence intelligence (Ozowara et al., 2022).

Additional studies broaden the applied context for the present framework (Dosunmu & Ogundele, 2020; Fadayomi et al., 2019; Ladapo et al., 2018, 2024a; Nnaji & Akinlolu, 2024).

The survey to this point has followed the strands most directly implicated in the detection and interpretation of hostile activity. The applied literature bearing on this framework is nonetheless wider still, encompassing the defense of industrial control environments, the governance of automated inference in regulated sectors, the reconciliation of records across institutional boundaries, and the political arrangements that make collective exchange possible in the first place. The paragraphs that follow trace these further strands and indicate, in each case, the specific transfer of method or of constraint on which the present proposal relies.

A closely allied stream of research addresses the protection of operational technology and industrial control environments, where the consequences of delayed detection are physical as well as informational. These studies develop threat models for adversarial machine learning in control systems, detection architectures that combine security information and event management with network detection and response, zero-trust designs for regulated energy networks, vulnerability and cyber-risk quantification schemes that translate exposure into governable measures, and mitigation strategies for protocol-level amplification attacks and for the security of instrumented transport and utility infrastructure. Their relevance to the present framework is direct, because such environments emit precisely the heterogeneous, high-volume indicator streams whose interpretation the proposed semantic layer is designed to accelerate (Adebayo et al., 2022, 2023a; Adegbite et al., 2020, 2022, 2023, 2024a, 2024b, 2025; Ahmed et al., 2021; Jimoh & Ahmed, 2024; Jimoh et al., 2023a, 2023b; Ogunwola, 2026; Ogunwola & Deborah, 2026).

A second body of engineering scholarship concerns the construction and continuous validation of enterprise defensive infrastructure. It examines deception-based architectures for disrupting persistent intrusion campaigns, automated threat-intelligence and adversary-simulation pipelines, breach-and-attack simulation for validating endpoint and perimeter controls, the modernization of firewall estates without loss of continuity, the security implications of directory-service attack paths, access-control policy design under privacy constraints, and the service-management disciplines through which such controls are operated day to day. This work supplies the deployment substrate on which any real-time interpretive layer must ultimately rest (Dosunmu & Ogundele, 2026a, 2026b; Ladapo et al., 2023a, 2023d; Ogunwola et al., 2026).

Because the framework proposed here deliberately preserves the analyst as the locus of decision, research on the human dimension of organizational security is of particular importance. Studies of security-awareness training, insider-risk behavior, and capability frameworks for

administrative personnel establish that technical controls underperform when the people operating them are inadequately prepared, and they support the argument advanced in Section VI that automation should be configured to strengthen rather than to displace human situational awareness (Hussain & Adebayo, 2026; Onche et al., 2024, 2026).

Financial fraud and risk analytics constitute a mature analogue to threat triage, since both involve searching vast transactional or textual records for rare, adversarially concealed patterns under severe time pressure. The relevant literature spans machine-learning advances in fraud detection, predictive models for anti-money-laundering surveillance, anomaly detection in financial time series, natural-language processing over social and open-source text for population-scale surveillance, privacy-preserving analytic architectures for regulated data, and the auditing of artificial-intelligence-powered systems themselves. It also documents the cognitive biases that distort expert judgment under uncertainty, a finding that bears directly on the prioritization logic developed in this article (Abetoh & Atakpa, 2024; Atakpa, 2021; Atakpa & Abetoh, 2022; Atakpa & Abolaji, 2022; Atakpa & Fobellah, 2023; Atakpa et al., 2023, 2024a).

Adjacent work on internal audit and risk governance specifies how organizations construct defensible assurance over automated processes. Contributions in this strand address risk-based auditing and internal control design, technology-enabled risk assessment, compliance analytics for institutions operating across jurisdictions, segregation of duties and access control in cross-border operations, corrective action planning and audit liaison, automated period-close controls embedded in enterprise resource planning systems, and statutory reporting and budget compliance in the public sector. These mechanisms are the institutional counterpart to the technical safeguards discussed in Section IX, and they indicate what auditable operation of a language-model pipeline would concretely require (Agu et al., 2023a; Akomolafe & Agu, 2018, 2019c; Akomolafe et al., 2023a, 2023b, 2025a; Dogbatsey & Ebhojie, 2018; Dogbatsey et al., 2026; Ebhojie et al., 2023a; Oyeleye et al., 2025).

Research on distributed-ledger and cross-border payment infrastructure contributes a further set of transferable mechanisms, namely tamper-evident provenance, interoperability across institutions that do not fully trust one another, automated risk scoring of inbound transactions, collaborative governance of shared data, and continuity of operation under geopolitical disruption. The structural resemblance to collective threat-intelligence exchange is close, since both require many independent participants to rely upon records they did not themselves generate, and both degrade when provenance is lost (Adesuyi et al., 2023; Akomolafe et al., 2024a, 2025b, 2025c; Olaogun et al., 2023).

A large literature on predictive analytics and enterprise decision systems demonstrates how inference over noisy historical data is converted into forward-looking operational guidance. It encompasses decision-support systems for resource planning under constraint, performance-intelligence and outcome-measurement models, real-time executive dashboards, integrated forecasting architectures, and predictive cost and capital optimization. Two lessons carry over to the defense setting. The first is that analytic value is realized only when model output is rendered in a form that a decision maker can act upon within the decision window. The second is that measurement of downstream outcomes, rather than of model accuracy alone, is the appropriate criterion of success (Adelanwa et al., 2023b; Lawal & Oduleye, 2018b; Oghenemaiga et al., 2024; Tonoyan et al., 2022c).

Studies of artificial intelligence in health and education supply the closest available evidence on the governance of language-capable systems in consequential settings. This work addresses user trust and regulatory oversight in telehealth, organizational readiness for generative systems, comparative governance and executive accountability structures, privacy-preserving data governance, equity of access, and the design of compliant distributed infrastructure for sensitive records. Although the domains differ, the governance problem is formally the same as that confronted in defense, namely how to obtain the throughput advantages of automated interpretation without ceding accountability for the resulting decisions (Afrihyia et al., 2024a, 2024b, 2025; Akintola et al., 2025; Kumuyi et al., 2024b; Ogunwola & Alozie, 2026a; Onwubuya et al., 2026; Orise & Niniola, 2023, 2025; Ozowara et al., 2025a).

Engineering research on energy-harvesting sensing and constrained wireless systems is relevant in a different register, since it concerns inference performed under strict energy, latency, and reliability budgets. Work on trade-offs among energy, reliability, and latency in green communication, lightweight machine learning for intermittent devices, reliable broadcast in batteryless networks, throughput degradation under adverse channel conditions, and automated test frameworks for wireless hardware illustrates the discipline required when analytic ambition meets hard resource limits. The same discipline governs the deployment of transformer inference at the network edge, a constraint acknowledged in Section X (Asiedu et al., 2026).

Further work on procurement analytics and sustainable supply-chain governance extends the earlier discussion of distributed information reconciliation. Studies of supplier risk assessment, category and spend analytics, environmental and social audit trails, and the verification of third-party claims speak to the problem of establishing confidence in records contributed by external parties whose collection practices are only partially observable. This is the same epistemic condition under which shared threat indicators are received (Abioye et al., 2023).

Research on the governance of multinational infrastructure and energy programs contributes evidence on how complex technical initiatives are steered across organizational and national boundaries. Studies of project governance structures, compliance and risk management in regulated delivery environments, data-center lifecycle management, and the composition and performance of distributed delivery teams illuminate the organizational conditions under which a capability of the kind proposed here would actually be fielded, a matter treated in Section VIII (Amayo et al., 2023a, 2024a, 2024b, 2025c).

A further strand of international-relations scholarship addresses the political architecture within which cross-border security cooperation occurs. Work on regional security cooperation models, peacebuilding and post-conflict recovery frameworks, economic interdependence, and the alignment of national policy with multilateral commitments is directly pertinent, because the collective defense posture that motivates this article is sustained by diplomatic arrangements rather than by technical protocols alone. The willingness of states to share indicators is a political fact before it is an engineering one (Liadi, 2024b).

Commercial analytics research provides an additional methodological parallel in the automated interpretation of large volumes of unstructured natural-language material. Studies of audience segmentation, attribution modeling, real-time performance monitoring, and forecasting over continuously arriving text and behavioral data address, in a lower-stakes setting, the same technical problem of extracting stable structure from noisy language at scale (Sanni, 2026a).

A final strand on pharmaceutical care, diagnostic integration, and medicine access contributes evidence on screening and prioritization under resource constraint. Work on systematic interaction screening frameworks, the integration of microbiological, virological, and serological evidence into a single diagnostic picture, quantitative assay accuracy, lean process improvement in inventory and distribution, quality assurance as a determinant of therapeutic reliability, and last-mile distribution to underserved populations addresses in a clinical register the same triad of problems confronted here: fusing heterogeneous evidence into a single assessment, ranking cases when capacity is insufficient to address all of them at once, and sustaining the integrity of a distribution chain on which many dependent parties rely (Asiedu, 2026).

VI. A FRAMEWORK FOR REAL-TIME SEMANTIC PROCESSING

A. Design Rationale and Methodology

The framework advanced in this article is developed as an artifact rather than a theory to be tested in isolation, and its construction is therefore governed by the tenets of design science research, which position the creation and evaluation of purposeful artifacts as a legitimate mode of scholarly inquiry (Hevner et al., 2004). Design science is appropriate here because the central problem, the timely conversion of shared cyber threat indicators into defensive action, is fundamentally a problem of building something that does not yet exist in operational form rather than describing a phenomenon that already does. The work is structured according to an established design science research methodology whose sequential activities, namely problem identification, definition of objectives, design and development, demonstration, evaluation, and communication, provide a disciplined scaffold that guards against the tendency of technical proposals to conflate aspiration with accomplishment (Peffer et al., 2007). Each subsequent architectural decision reported below is traceable to an explicit objective derived from the preceding analysis of intelligence latency, semantic heterogeneity, and analyst cognitive load. By anchoring the artifact to documented requirements, the framework remains falsifiable in the sense that its components can be judged against whether they demonstrably serve the objectives that motivated them, an accountability that distinguishes rigorous design from unconstrained engineering enthusiasm and keeps the contribution answerable to the defensive mission it purports to accelerate.

The epistemological posture underlying this design is pragmatism, which evaluates knowledge claims by their practical consequences and treats the workability of an intervention in its intended context as the decisive criterion of merit (Morgan, 2014). Pragmatism is well suited to a national defense setting where the ultimate arbiter is not theoretical elegance but whether defenders act sooner and more accurately, and it licenses a methodological pluralism that draws on computational, organizational, and cognitive evidence without demanding allegiance to a single disciplinary orthodoxy. The reasoning that generated the framework is best characterized as abductive, proceeding from surprising observations, chiefly that abundant sharing coexists with persistent operational latency, toward the most plausible explanatory and remedial account, then iteratively refining that account as it is confronted with the particulars of real indicator streams (Timmermans & Tavory, 2012). Abductive analysis rejects both the purely deductive imposition of a preconceived model and the purely inductive accumulation of ungrounded observation, instead cultivating a disciplined dialogue between established theory and anomalous data. This stance shapes how the framework should be read: not as a finished prescription validated once and for all, but as a defensible conjecture whose components are held provisionally and revised as demonstration exposes where semantic processing succeeds, degrades, or fails against the adversarial and heterogeneous realities of shared threat intelligence.

B. Architecture Overview

The proposed architecture is organized as a layered pipeline that transforms heterogeneous shared indicators into prioritized, context-rich decision support while preserving human authority over consequential action. At the base sits an ingestion layer that connects to the diverse feeds enumerated in Table I, normalizing transport protocols and payload formats so that downstream components receive a consistent internal representation regardless of provenance. Above ingestion, a preprocessing and enrichment layer resolves entities, deduplicates near-identical reports, and attaches provenance and confidence metadata that travel with each indicator throughout its lifecycle. The semantic core, built on large language model encoders, projects both structured indicators and unstructured narrative into a shared vector space where correlation is computed by semantic proximity rather than exact string equality, allowing conceptually related threats expressed in different vocabularies to be linked. A retrieval-augmented reasoning layer grounds model outputs in an evolving corpus of validated intelligence, mitigating fabrication and enabling traceable justification. Finally, a prioritization and review layer ranks correlated clusters by estimated relevance and severity and routes them to analysts through interfaces designed to sustain rather than saturate situational awareness. Each layer exposes auditable interfaces so that the system's inferences can be inspected, contested, and corrected. The remainder of this section elaborates the operative layers, and the feed taxonomy that the ingestion layer must accommodate is summarized below.

Table I. Illustrative data sources and their application to the framework

Data Source	Type of Data	Relevance
MITRE ATT&CK	Tactics and adversary behavior descriptions	Structured language patterns for threat modeling
CVE database	Vulnerability descriptions	Semantically repetitive technical text
CERT advisories	Incident reports	Correlation of related threats
Open-source reports	Unstructured threat narratives	Natural-language input for model processing
Synthetic threat logs	Simulated defense reports	Safe modeling of sensitive scenarios

C. Ingestion and Preprocessing

Ingestion confronts the practical reality that shared threat intelligence arrives through incompatible channels, ranging from structured exchange standards and machine-readable feeds to email advisories, vendor blogs, and free-text incident reports, each with distinct latency, reliability, and trust characteristics. The ingestion layer therefore implements adapters that decouple source-specific transport and encoding concerns from the

semantic processing that follows, emitting a canonical internal record in which every indicator carries a stable identifier, a timestamp, a source descriptor, and an initial confidence estimate. Preprocessing then performs the unglamorous but consequential work on which all subsequent accuracy depends. Defanged indicators are re-canonicalized, character encodings and internationalized domain representations are normalized, and syntactic variants of the same artifact are reconciled so that later correlation does not fragment a single threat across superficially different tokens. Entity resolution links aliases, infrastructure, and campaign labels to persistent internal references, while deduplication collapses the redundant re-reporting that inflates volume without adding information. Crucially, preprocessing preserves rather than discards provenance, recording which source asserted each claim and with what stated confidence, because downstream prioritization and human review both depend on the ability to weigh assertions by their origin. Malformed, contradictory, or low-integrity records are flagged rather than silently dropped, ensuring that the pipeline degrades transparently and that analysts retain visibility into the quality of the intelligence they are asked to trust. Closely comparable problems of normalizing and reconciling records contributed by external parties, and of preserving provenance across institutional boundaries, are treated at length in the distributed-ledger and procurement-assurance literatures surveyed above (Abioye et al., 2023; Adesuyi et al., 2023; Akomolafe et al., 2024a, 2025b, 2025c; Olaogun et al., 2023).

D. Semantic Encoding and Correlation

The semantic core replaces brittle lexical matching with representation learning, encoding both structured indicators and their surrounding narrative into dense vectors that capture meaning rather than surface form. This capability rests on the transformer architecture, whose self-attention mechanism allows every token to be interpreted in the context of every other token in a sequence, yielding representations sensitive to the relational structure of language rather than to isolated keywords (Vaswani et al., 2017). Bidirectional encoders trained with masked language modeling produce contextual embeddings in which a term such as a malware family name or a technique reference acquires a representation informed by the entire report in which it appears, so that the same string carries different vectors when it denotes an actor, a tool, or a defensive control (Devlin et al., 2019). When indicators and analyst queries are projected into this common space, correlation becomes a matter of geometric proximity, and threats described in disparate vocabularies, across languages, or through paraphrase are drawn together because their meanings converge even where their tokens diverge. This semantic correlation directly addresses the failure mode identified earlier, in which valuable connections are lost because sharing formats encode the same phenomenon inconsistently, and it allows the system to surface non-obvious relationships between fragments of intelligence that no exact-match rule could ever join.

Encoding alone, however, does not confer reliability, and the generative reasoning that turns correlated clusters into human-readable assessments must be constrained lest it produce fluent but unfounded conclusions. Large autoregressive models exhibit a striking capacity to perform inference from few or no task-specific examples, which makes them attractive for the open-ended, rapidly shifting vocabulary of threat intelligence where labeled training data is perpetually scarce (Brown et al., 2020). That same generative fluency, unconstrained, is a liability in a defense context, because a confident fabrication can misdirect scarce analytic attention or, worse, contaminate downstream action. The framework therefore couples generation to retrieval, requiring that any synthesized assessment be conditioned on and cite passages drawn from a curated store of validated indicators and prior adjudications, so that the model reasons over evidence it can point to rather than over parametric memory it cannot (Lewis et al., 2020). Retrieval grounding narrows the space of plausible outputs to those supported by retrievable intelligence, furnishes analysts with the source passages behind each claim, and localizes correction, since updating the evidence store adjusts behavior without retraining. Correlation and grounded generation thus operate in tandem: the former discovers what is related, and the latter explains why, in terms an analyst can independently verify against the underlying reporting.

E. Prioritization and Human-in-the-Loop Review

Because analytic attention is the binding constraint identified throughout this article, the penultimate layer is dedicated not to producing more information but to ordering it, so that the indicators most consequential to a given defender surface first. Prioritization scores each correlated cluster along dimensions of estimated severity, relevance to the protected environment, novelty relative to prior knowledge, and source confidence, then composes these into a ranking that determines the sequence in which items reach human review. The design goal is explicitly to cultivate and protect situational awareness in its full sense, encompassing the perception of relevant elements, the comprehension of their meaning, and the projection of their near-future implications, rather than merely delivering data to a screen (Endsley, 1995). Interfaces are accordingly built to compress rather than multiply cognitive load, presenting a correlated cluster together with its grounding evidence, its provenance, and its rationale as a single reviewable unit, so that an analyst apprehends not only that an alert exists but why the system believes it matters and on what basis. By foregrounding comprehension and projection, the layer aims to shift the analyst from reactive triage of undifferentiated volume toward anticipatory reasoning about adversary intent, converting the system's throughput advantage into a genuine gain in the quality of human judgment rather than a faster firehose. The design of this layer follows the lesson drawn from enterprise decision systems that analytic output creates value only when it reaches the decision maker in an actionable form and within the decision window (Adelanwa et al., 2023b; Lawal & Oduleye, 2018b; Oghenemaiga et al., 2024; Tonoyan et al., 2022c).

The insistence on retaining a human in the loop is not a concession to immature technology but a deliberate expression of sociotechnical design, which holds that the performance of any operational system emerges from the joint optimization of its social and technical components rather than from the perfection of either alone (Trist & Bamforth, 1951). Automation that displaces the analyst would forfeit the contextual, ethical, and accountable judgment that consequential defensive decisions demand, while manual processing that ignores machine assistance would remain trapped in the latency this framework seeks to overcome; the sociotechnical stance instead allocates to the machine what it does well, namely tireless correlation and grounded synthesis at scale, and to the human what only the human can furnish, namely authority, accountability, and situated interpretation. Review is therefore designed as an active collaboration in which analyst decisions to confirm, reject, or escalate are captured as feedback that continually refines prioritization and enriches the retrieval store, closing a learning loop that improves the system precisely through use. The interface exposes model uncertainty and dissenting evidence rather than concealing them, so that trust is calibrated to demonstrated reliability. In this way the layer institutionalizes the principle that acceleration must never come at the expense of the deliberation on which legitimate defensive action ultimately rests. Evidence from studies of security awareness, insider risk, and supervised autonomy in robotic systems

supports this configuration, indicating that automation performs best where it strengthens rather than substitutes for the situational awareness of the operator (Hussain & Adebayo, 2026; Onche et al., 2024, 2026).

VII. ILLUSTRATIVE EVALUATION

Consistent with the design science commitment to demonstration, the framework is examined here through an illustrative evaluation intended to make its claimed advantages concrete and contestable rather than to assert definitive empirical superiority. The evaluation is framed as a comparison between the semantic pipeline described above and the prevailing baseline of rule-based, exact-match correlation applied to the same stream of shared indicators, holding the input feeds and the analytic staffing constant so that differences can be attributed to the processing paradigm rather than to resourcing. The scenario envisions a realistic burst of intelligence surrounding an emerging campaign, in which the same underlying infrastructure and techniques are reported by multiple sharing partners using divergent naming conventions, partial observations, and a mixture of structured and narrative formats. Under such conditions the baseline predictably fragments the campaign into disconnected alerts wherever token-level identity fails, forcing analysts to reconstruct by hand the coherence that the reporting never explicitly encoded. The semantic pipeline, by contrast, is expected to draw the dispersed fragments into a single correlated cluster on the strength of meaning-level proximity, presenting the campaign as one prioritized, grounded assessment. The purpose of the exercise is to expose the mechanisms through which the two approaches diverge, so that the sources of any advantage are visible and open to critique rather than buried in an aggregate score. The dimensions along which the comparison is drawn are chosen to reflect what actually constrains operational value in threat intelligence work, and they extend beyond raw detection counts to encompass the analytic experience of consuming the output. Correlation completeness captures the proportion of genuinely related indicators that are successfully joined, a dimension on which lexical matching is structurally disadvantaged whenever the same entity is expressed inconsistently, a difficulty documented across the literature on mining actionable knowledge from cyber threat intelligence (Sun et al., 2023). Time to actionable assessment measures the interval between indicator arrival and the availability of a reviewed, contextualized judgment, the very latency this work is designed to compress. Analyst cognitive load, though harder to quantify, is represented through the number of discrete items an analyst must manually reconcile to reach an equivalent understanding. Provenance transparency records whether the reasoning behind each surfaced assessment can be traced to specific evidence, a property the retrieval-grounded design is built to guarantee and that the baseline does not address. False-positive burden reflects the volume of spurious correlations imposed on human attention. Table II summarizes the qualitative contrast across these dimensions, and the discussion that follows interprets the pattern it reveals in light of established theory concerning detection and information.

Table II. Comparative characteristics: language-model pipeline versus keyword baseline

Dimension	Keyword baseline	LLM-based pipeline
Matching basis	Exact tokens and rules	Contextual meaning
Handling of reworded reports	Poor	Strong
Redundancy suppression	Limited	Substantial
Analyst cognitive load	High	Reduced
Real-time suitability	Constrained	Supported under adequate resources

The contrast summarized in Table II is best understood not as a simple accuracy gain but as a favorable movement along an inescapable tradeoff, and signal detection theory supplies the vocabulary for that interpretation, distinguishing an operator's underlying sensitivity to genuine threats from the decision threshold that governs how readily an alert is raised (Green & Swets, 1966). Exact-match correlation, by refusing to link semantically related but lexically divergent indicators, effectively adopts a conservative threshold that suppresses false positives at the cost of missing true connections, whereas the semantic pipeline raises sensitivity so that more genuine relationships are detected, with retrieval grounding and prioritization enlisted to hold spurious correlations in check rather than accepting them as the price of higher recall. The improvement is therefore a shift of the entire operating characteristic, not merely a relaxation of the threshold. The same pattern can be read through the lens of communication theory, in which the useful transmission of information is limited by noise and by the redundancy of the encoding (Shannon, 1948). Semantic correlation reduces the effective noise introduced by inconsistent sharing vocabularies and exploits the redundancy of multiple partial reports to reconstruct a fuller signal, so that the analyst receives more of the intelligence that the sharing community collectively possessed but had, until processing, failed to make legible.

VIII. DISCUSSION

A. Technological Implications

Technologically, the framework signals a migration in cyber defense tooling away from deterministic pattern matching toward probabilistic semantic reasoning, and that migration carries consequences that reach well beyond correlation accuracy. Treating indicators as points in a learned representation space rather than as opaque strings makes defensive infrastructure inherently more adaptable to novelty, because previously unseen threats expressed in unfamiliar language can still be located relative to known concepts, whereas rule-based systems remain blind to whatever their authors did not anticipate. This adaptability comes bundled with new engineering obligations. Representation quality now depends on encoder models whose behavior must be monitored for drift as adversary language evolves, retrieval stores must be curated and versioned so that grounding remains trustworthy, and inference must be delivered within latency budgets tight enough to preserve the real-time character the mission demands. The probabilistic core also reframes correctness itself, since outputs carry calibrated uncertainty rather than binary verdicts, obliging surrounding systems to consume and display confidence rather than to assume it. In aggregate, the technological implication is that semantic defense tooling is less a fixed ruleset to be deployed than a living system to be maintained, one whose accuracy is contingent on the ongoing stewardship of its

models and its evidence, and whose advantages are realized only where that stewardship is sustained with discipline. The same migration is visible in the defense of industrial control environments and in the engineering of enterprise detection estates, where signature-based controls have progressively given way to behavioral and model-driven architectures (Adebayo et al., 2022, 2023a; Adegbite et al., 2020, 2022, 2023, 2024a, 2024b, 2025; Ahmed et al., 2021; Dosunmu & Ogundele, 2026a, 2026b; Jimoh & Ahmed, 2024; Jimoh et al., 2023a, 2023b; Ladapo et al., 2023a, 2023d; Ogunwola, 2026; Ogunwola & Deborah, 2026; Ogunwola et al., 2026).

B. Organizational Implications

Organizationally, adopting semantic processing redistributes work and expertise in ways that institutions must anticipate rather than absorb by accident. As the machine assumes the burden of correlation and first-pass synthesis, the analyst's role shifts upward toward adjudication, contextualization, and the exercise of judgment on the assessments the system surfaces, which changes the skills that defensive teams must recruit and cultivate. Value accrues not to those who can manually reconcile the largest volume of feeds but to those who can critically interrogate machine-generated reasoning, recognize when grounding is thin, and translate a prioritized assessment into decisions attuned to their particular operational context. This transition demands deliberate investment in training, in interface design that renders model reasoning inspectable, and in feedback practices that capture analyst judgments as a resource for continual improvement rather than letting them evaporate. It also introduces new interdependencies, since the quality of the system now depends jointly on data engineers who maintain ingestion, on curators who tend the retrieval store, and on analysts who supply corrective signal, so that organizational silos which separate these functions will undermine the very performance the technology promises. The organizational implication, then, is that the framework is adopted successfully only when institutions treat it as a change to how people and machines work together, not as a component procured and installed beneath an unchanged division of labor. Research on the governance of multinational technical programs reaches a parallel conclusion, finding that delivery outcomes turn less on the technology procured than on the governance structures and team arrangements through which it is introduced (Amayo et al., 2023a, 2024a, 2024b, 2025c).

C. Strategic Implications

Strategically, compressing the interval between shared intelligence and defensive action alters the temporal contest at the heart of cyber conflict, in which advantage often accrues to whichever side converts information into effect more quickly. If a defensive community can collectively metabolize its shared indicators in near real time, the value of sharing itself increases, because intelligence contributed by one partner can be operationalized by others before the adversary's window of opportunity closes, strengthening the incentive to share and the resilience of the collective. This has implications for how national defense organizations conceive of coalition and public-private cooperation, since the payoff from participation grows with the speed and fidelity of downstream processing, and a partner that cannot rapidly consume what it receives contributes less and benefits less. At the same time, the strategic calculus must reckon with the adversary's capacity to adapt, including attempts to poison shared feeds, to craft indicators that manipulate semantic correlation, or to exploit the automation itself, so that acceleration must be pursued alongside robustness rather than in tension with it. The broader strategic implication is that semantic processing is not merely an efficiency measure but a potential shift in the initiative, offering defenders a route to reclaim tempo, provided the accompanying governance, trust, and resilience mature at the same pace as the technical capability. It should be recalled that the collective posture on which this advantage depends is sustained by diplomatic and cooperative arrangements rather than by technical protocols alone, a dependency examined in the literature on regional security cooperation (Liadi, 2024b).

IX. GOVERNANCE, ETHICS, AND RESPONSIBLE DEPLOYMENT

The deployment of large language models within national defense infrastructure raises ethical questions too consequential to be treated as an afterthought appended to a technical design, and this section therefore situates the framework within a principled account of responsible artificial intelligence. A widely endorsed ethical framework for socially beneficial artificial intelligence organizes its guidance around the principles of beneficence, non-maleficence, autonomy, justice, and explicability, and each bears directly on the present system (Floridi et al., 2018). Beneficence and non-maleficence together demand that the acceleration the framework offers genuinely reduce harm rather than merely relocate it, which in a defensive context means that faster processing must not translate into hastier or less accountable action against imperfectly attributed threats. The principle of autonomy underwrites the framework's insistence on human authority over consequential decisions, ensuring that the machine augments rather than usurps the judgment of those who bear responsibility. Justice requires vigilance against the differential impacts of correlation errors, since a false linkage can direct defensive or investigative attention toward parties who merit none. Explicability, finally, is not a soft aspiration but a design constraint, demanding that the system's inferences be intelligible and contestable, and it is this principle above all that the retrieval-grounded architecture is engineered to honor by making every assessment traceable to the evidence that produced it. The assurance literature on risk-based auditing, technology-enabled control testing, and compliance analytics indicates what such traceability requires institutionally, beyond what any architecture can supply on its own (Agu et al., 2023a; Akomolafe & Agu, 2018, 2019c; Akomolafe et al., 2023a, 2023b, 2025a; Dogbatsey & Ebhojie, 2018; Dogbatsey et al., 2026; Ebhojie et al., 2023a; Oyeleye et al., 2025).

Translating high-level principles into operational practice is notoriously difficult, and a comprehensive survey of the global landscape of artificial intelligence ethics guidelines has shown both a convergence on recurring themes such as transparency, accountability, and fairness and a persistent gap between the articulation of these values and their concrete implementation (Jobin et al., 2019). That gap is the central challenge for responsible deployment in defense, where abstract commitments to transparency mean little unless they are instantiated in mechanisms an analyst can actually exercise. The framework attempts to close the gap by binding governance to architecture rather than to policy documents alone. Retrieval grounding operationalizes explicability by furnishing, for every synthesized assessment, the specific passages of validated intelligence on which it rests, so that a reviewer can independently confirm or refute the machine's reasoning against traceable sources rather than accepting it on faith (Lewis et al., 2020). Provenance metadata carried from ingestion through review operationalizes accountability by preserving which source asserted each claim and with what confidence, and the human-in-the-loop review layer operationalizes contestability by ensuring that no consequential inference

reaches action without an accountable human who can and must interrogate it. In this way the design treats ethical requirements as functional specifications, embedding them where they can be enforced rather than exhorted.

Responsible deployment further requires confronting the specific hazards that attend delegating inference to algorithms, which a systematic map of the ethics of algorithms organizes into concerns including inconclusive, inscrutable, and misguided evidence, unfair outcomes, and transformative effects that reshape how decisions are made (Mittelstadt et al., 2016). Inconclusive evidence is addressed by propagating calibrated uncertainty rather than presenting probabilistic correlations as certainties; inscrutable evidence is mitigated, though never wholly eliminated, by the grounding and provenance mechanisms that expose the basis of each output; and misguided evidence is guarded against through preprocessing that flags low-integrity inputs and through the recognition that adversaries may deliberately attempt to manipulate the intelligence supply. The concern with unfair outcomes underscores the earlier point about justice, demanding continual monitoring for systematic correlation errors that could burden particular parties. The transformative concern is perhaps the most subtle, for as semantic processing reshapes the analyst's role and the tempo of defense, it alters the very structure of accountability and expertise, and this transformation must be governed consciously rather than allowed to occur by drift. Responsible deployment is therefore a continuing institutional practice, encompassing audit, red-teaming of the models against adversarial manipulation, clear allocation of responsibility for machine-assisted decisions, and mechanisms for redress, without which the framework's technical merits could not be ethically sustained in operation. Comparative work on the governance of generative and predictive systems in health and education offers the closest available precedent, and it consistently finds that organizational readiness and accountability structures, rather than model capability, determine whether such systems can be responsibly fielded (Afrihyia et al., 2024a, 2024b, 2025; Akintola et al., 2025; Kumuyi et al., 2024b; Ogunwola & Alozie, 2026a; Onwubuya et al., 2026; Orise & Niniola, 2023, 2025; Ozowara et al., 2025a).

X. LIMITATIONS AND THREATS TO VALIDITY

The contribution of this article is conceptual, and intellectual honesty requires acknowledging the limitations that follow from that status before its claims are read as more than they are. Foremost, the evaluation presented is illustrative rather than empirical, demonstrating the mechanisms by which semantic processing is expected to outperform lexical correlation without measuring that advantage against instrumented operational data, and the qualitative contrasts reported should accordingly be treated as hypotheses to be tested rather than as validated results. The framework's performance in practice will depend on variables the conceptual treatment cannot resolve, including the quality and drift of the encoder models, the completeness and curation of the retrieval store, the latency achievable under realistic load, and the fidelity of the shared feeds themselves, any of which could attenuate or negate the projected benefits. There is a real risk that the very fluency of large language models induces misplaced analyst trust, so that grounded-looking but subtly erroneous assessments are accepted with insufficient scrutiny, a failure mode the design mitigates but cannot foreclose. Adversarial manipulation of the intelligence supply, including data poisoning and inputs crafted to distort semantic correlation, represents a further limitation that a controlled conceptual analysis can name but not fully characterize, since its severity depends on an adaptive opponent whose behavior lies outside the model's assumptions. A further limitation is one of resources rather than of adversaries, since transformer inference at operational tempo imposes energy, latency, and reliability costs of the kind analyzed in the literature on constrained sensing and communication systems (Asiedu et al., 2026).

Beyond these substantive limitations, several threats to validity warrant explicit statement. Construct validity is threatened by the difficulty of measuring the outcomes that matter most, since time to actionable assessment and analyst cognitive load are influenced by organizational factors the framework does not control, and any future measurement must guard against operationalizing these constructs too narrowly. Internal validity is limited because the illustrative comparison holds many conditions fixed by stipulation rather than by observation, so alternative explanations for the projected advantage cannot yet be excluded. External validity is constrained by the specificity of the national defense context, whose trust relationships, sharing arrangements, and adversary sophistication may not generalize to other domains that consume threat intelligence, and equally by the possibility that results from one sharing community fail to transfer to another with different feed composition. Finally, the framework inherits the assumptions of the models it employs, and shifts in the underlying technology could alter its behavior in ways the present analysis does not anticipate. These threats do not undermine the conceptual contribution, which is to articulate and justify a coherent design, but they do circumscribe it, marking the boundary between what has been argued and what remains to be demonstrated, and thereby defining the agenda that responsible future work must address. Comparable cautions arise in fraud and audit analytics, where documented biases in expert judgment and the difficulty of validating models against adaptively concealed behavior have proved to be enduring rather than transitional obstacles (Abetoh & Atakpa, 2024; Atakpa, 2021; Atakpa & Abetoh, 2022; Atakpa & Abolaji, 2022; Atakpa & Fobellah, 2023; Atakpa et al., 2023, 2024a).

XI. CONCLUSION AND FUTURE WORK

This article has argued that the persistent gap between the abundance of shared cyber threat intelligence and its slow conversion into defensive action is not primarily a problem of insufficient sharing but of insufficient semantic processing, and it has proposed a conceptual framework in which large language models compress that gap by correlating indicators through meaning rather than surface form. Developed as a design science artifact under a pragmatic, abductively refined methodology, the framework organizes ingestion, preprocessing, transformer-based semantic encoding, retrieval-grounded reasoning, and human-in-the-loop prioritization into a coherent pipeline whose every layer is traceable to an explicit objective derived from the operational realities of intelligence latency, format heterogeneity, and analyst cognitive load. The illustrative evaluation clarified the mechanisms through which such a pipeline is expected to recover connections that exact-match correlation structurally loses, and the accompanying discussion traced the technological, organizational, and strategic consequences of shifting cyber defense toward probabilistic semantic reasoning. Throughout, the article has insisted that acceleration and accountability are not competing goals but joint requirements, embedding explicability, provenance, and human authority into the architecture as functional specifications rather than aspirational overlays. The central conclusion is that reclaiming defensive tempo is achievable, but only through systems that treat machine correlation and human judgment as complementary components of a single sociotechnical whole.

The path from this conceptual contribution to operational impact defines a concrete agenda for future work. The most immediate priority is empirical validation, replacing the illustrative evaluation with instrumented studies on real shared feeds that measure correlation completeness, time to actionable assessment, analyst cognitive load, and false-positive burden against rigorously matched baselines, so that the framework's projected advantages are confirmed, bounded, or refuted. Complementary work should investigate the calibration of analyst trust, examining how interface design and uncertainty presentation affect whether reviewers appropriately scrutinize grounded assessments, and should develop adversarial evaluations that probe the resilience of semantic correlation against data poisoning and crafted manipulation of the intelligence supply. Further research is warranted on the curation and versioning of retrieval stores, on the detection and remediation of representational drift as adversary language evolves, and on the organizational practices through which institutions successfully redistribute expertise around a semantic pipeline. Finally, the governance mechanisms outlined here invite empirical study of their own, testing whether embedded provenance and contestability genuinely produce accountable machine-assisted decisions in practice. Pursued together, these lines of inquiry would move the framework from a defensible conjecture toward a validated, responsibly governed capability, advancing the broader ambition of enabling national defense to act at the speed that a contested cyber environment increasingly demands.

DECLARATIONS

Author Contributions

The authors confirm their contributions to this manuscript using the CRediT taxonomy as follows. Chukwunye Amadi: conceptualization, methodology, investigation, writing of the original draft, supervision, and project administration. Abolaji Adebayo: formal analysis, development of the theoretical framework, validation, and writing through review and editing. Ayokunle Olamide Ijagbemi: literature curation, data curation, visualization, and writing through review and editing. All authors reviewed and approved the final manuscript and accept responsibility for its content.

Note on Authorship

The earlier working paper on which this manuscript builds, Amadi (2025), was circulated under the sole authorship of Chukwunye Amadi. Abolaji Adebayo and Ayokunle Olamide Ijagbemi joined the present work and are included as authors on the basis of the substantive contributions recorded above, principally the theoretical grounding, the governance and ethics analysis, and the reconstruction of the related-work positioning, none of which appeared in the earlier paper. All three authors meet the authorship criteria of the International Committee of Medical Journal Editors as applied to this field, and all have agreed to the change in authorship and to the order in which authors are listed.

Competing Interests

The authors declare that they have no competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Funding

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Data Availability

No new data were generated or analysed in the course of this study. The manuscript is a conceptual contribution, and the comparison presented in Section VII is an illustrative exposition of mechanism rather than an experiment conducted on a dataset. The standards, knowledge bases, and advisory sources referred to in Table I are described there to indicate the classes of input a deployed implementation would ingest, and no such corpus was assembled, processed, or evaluated for this paper. All sources supporting the arguments made here are cited in the reference list and are publicly available.

Ethical Approval and Consent

This study did not involve human participants, animal subjects, or personally identifiable data, and no ethical approval or informed consent was therefore required. The paper raises ethical questions concerning the deployment of language models in national defence settings, and these are addressed substantively in Section IX rather than as a procedural declaration.

Use of Artificial Intelligence Tools

The authors should state here any use of generative artificial intelligence tools in the preparation of this manuscript, in accordance with the policy of the journal to which it is submitted, and confirm that the authors take full responsibility for the content of the published work.

REFERENCES

- 1) Abetoh, N. F., & Atakpa, M. I. (2024). Audit analytics in healthcare financial oversight: Leveraging data science to strengthen accountability in multilateral grant ecosystems. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 10(6), 2710–2747. <https://doi.org/10.32628/CSEIT2410791>
- 2) Abioye, R. F., Okojie, J. S., Filani, O. M., Ike, P. N., Idu, J. O. O., Nnabueze, S. B., Okojokwu-Idu, J. O., & Ihwughwawwe, S. I. (2023). Automated ESG reporting in energy projects using blockchain-driven smart compliance management systems. *International Journal of Multidisciplinary Evolutionary Research*, 4(2), 120–129. <https://doi.org/10.54660/IJMER.2023.4.2.120-129>
- 3) Adebayo, A., Adegbite, M. P., & Ahmed, M. O. (2022). Adversarial machine learning in critical infrastructure: A conceptual framework for threat modeling AI enabled OT systems. *World Journal of Innovation and Modern Technology*, 6(1), 184–234. <https://doi.org/10.56201/wjimt.v6.no1.2022.pg184.234>
- 4) Adebayo, A., Adegbite, M. P., & Ahmed, M. O. (2023a). AI augmented threat detection in industrial control systems: A systematic review of machine learning approaches for ICS anomaly detection. *International Journal of Engineering and Modern Technology*, 9(3), 287–340. <https://doi.org/10.56201/ijcsmt.v9.no3.2023.pg287.340>
- 5) Adebayo, A., Adepoju, P. A., & Ozowara, D. E. (2023b). Conceptual model for cyber risk pricing and insurance structuring in health technology platforms. *Gyanshauryam, International Scientific Refereed Research Journal*, 6(6). <https://doi.org/10.32628/GISRRJ>

- 6) Adegbite, M. P., Adebayo, A., & Ahmed, M. O. (2020). Securing operational technology networks in electric utilities: A systematic review of NERC CIP compliance and architectural threat mitigation. *International Journal of Engineering and Modern Technology*, 6(3), 103–146. <https://doi.org/10.56201/ijemt.vol.6.no3.2020.pg103.146>
- 7) Adegbite, M. P., Adebayo, A., & Ahmed, M. O. (2022). A security architecture model for IT and OT convergence in regulated energy networks: Design principles and governance alignment. *International Journal of Computer Science and Mathematical Theory*, 8(2), 81–132. <https://doi.org/10.56201/ijcsmt.vol.8.no2.2022.pg81.132>
- 8) Adegbite, M. P., Adebayo, A., & Ahmed, M. O. (2023). ICS and SCADA threat detection architectures in energy sector networks: A systematic review of SIEM, NDR, and anomaly detection approaches. *World Journal of Innovation and Modern Technology*, 7(2), 121–182. <https://doi.org/10.56201/wjimt.vol.7.no2.2023.pg121.182>
- 9) Adegbite, M. P., Adebayo, A., & Ahmed, M. O. (2024a). A vulnerability governance architecture for power and utilities corporations: From exposure mapping to remediation verification. *World Journal of Innovation and Modern Technology*, 8(6), 185–246. <https://doi.org/10.56201/wjimt.vol.8.no6.2024.pg185.246>
- 10) Adegbite, M. P., Adebayo, A., & Ahmed, M. O. (2024b). An AI driven security operations architecture for utility sector SOCs: Integrating threat intelligence, behavioral analytics, and automated response. *International Journal of Engineering and Modern Technology*, 10(11), 197–256. <https://doi.org/10.56201/ijemt.vol.10.no11.2024.pg197.256>
- 11) Adegbite, M. P., Adebayo, A., & Ahmed, M. O. (2025). A cyber risk quantification and governance architecture for critical infrastructure: From posture measurement to executive reporting. *International Journal of Computer Science and Mathematical Theory*, 11(12), 232–293. <https://doi.org/10.56201/ijcsmt.vol.11.no12.2025.pg232.293>
- 12) Adelanwa, A., Basnet, A., & Anene, U. N. (2023b). Predictive analytics models for financial risk detection and fraud prevention in public systems. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6). <https://doi.org/10.62225/2583049X.2023.3.6.5969>
- 13) Adesuyi, M. O., Akomolafe, O., Olaogun, B. O., Ndukwe, V. U., & Sakyi, J. K. (2023). Global interoperability model for distributed ledger-based cross-border payment systems. *International Journal of Multidisciplinary Research and Growth Evaluation*, 4(6), 1291–1300. <https://doi.org/10.54660/IJMRGE.2023.4.6.1291-1300>
- 14) Adeyelu, O. O. (2020). An Artificial Intelligence-Driven Conceptual Model for Wildlife Strike Risk Quantification and Aircraft Airworthiness Impact Assessment at High-Traffic African Airports. *Iconic Research and Engineering Journals*, 4(1), 339–368. <https://doi.org/10.64388/IREV4I1-1718482>
- 15) Adeyelu, O. O. (2023). A Risk-Tiered Oversight Model for Unmanned Aircraft Operations in Shared Controlled Airspace Over Nigerian Airports. *International Journal of Scientific Research in Civil Engineering*, 7(4), 147–184. <https://doi.org/10.32628/IJSRCE237419>
- 16) Adeyelu, O. O., & Dagodzo, D. (2024a). A Maturity Model for Predicting Airport Safety Audit Outcomes in Resource-Constrained Regulatory Environments. *International Journal of Scientific Research in Civil Engineering*, 8(4), 132–170. <https://doi.org/10.32628/IJSRCE248423>
- 17) Adeyelu, O. O., & Dagodzo, D. (2024b). Advances in Artificial Intelligence and Data-Driven Wildlife Hazard Monitoring and Incident Reduction at International Airports in West Africa. *International Journal of Scientific Research in Science and Technology*, 11(5), 872–910. <https://doi.org/10.32628/IJSRST52310285>
- 18) Afrihyia, E., Akinse, S. G., & Ojukwu, P. U. (2024a). Comparative governance of AI-driven healthcare management: Executive oversight, regulatory structures, and accountability in the United States and developing countries. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3087–3102. <https://doi.org/10.62225/2583049X.2024.4.6.5920>
- 19) Afrihyia, E., Akinse, S. G., & Ojukwu, P. U. (2025). Organizational readiness for generative AI integration in healthcare operations: Comparative management capabilities between the U.S. and low- and middle-income countries. *Iconic Research and Engineering Journals*, 8(10), 1673–1697. <https://doi.org/10.64388/IREV8I10-1714670>
- 20) Afrihyia, E., Ojukwu, P. U., & Akinse, S. G. (2024b). Privacy-preserving health data governance models: A comparative review of blockchain and cryptographic strategies in U.S. and developing healthcare systems. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3071–3086. <https://doi.org/10.62225/2583049X.2024.4.6.5919>
- 21) Agu, M. U., Akomolafe, O., & Bello, A. (2023a). A comparative review of SOX compliance frameworks in cross-border financial auditing. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), 2297–2306. <https://doi.org/10.62225/2583049X.2023.3.6.5362>
- 22) Ahmed, M. O., Adegbite, M. P., & Adebayo, A. (2021). Zero trust architecture for operational technology in North American critical infrastructure: A framework for implementation and resilience optimization. *International Journal of Engineering and Modern Technology*, 7(1), 66–113. <https://doi.org/10.56201/ijemt.vol.7.no1.2021.pg66.113>
- 23) Akeju, B., Edviri, J., Ogbale, J. I., Okoruwa, P. O., Fadayomi, O., & Abolaji, T. O. (2018). Conceptual model for insider threat classification and risk modeling in complex digital systems. *Iconic Research and Engineering Journals*, 1(9), 476–492. <https://doi.org/10.64388/IREV1I9-1713778>
- 24) Akintola, A. S., Fawehinmi, Y. O., Aiyenitaju, O., Chinemerem, B., & Emmanuella, O. (2025). Digital transformation in telehealth: A systematic review of user trust, privacy, and regulatory governance in AI-powered remote monitoring systems. *Journal of Scientific Research and Reports*, 31(11), 163–182. <https://doi.org/10.9734/jsrr/2025/v31i113658>
- 25) Akomolafe, O., & Agu, M. U. (2018). A conceptual model for enhancing internal audit quality through technology-enabled risk assessment frameworks. *Iconic Research and Engineering Journals*, 1(9), 458–475.
- 26) Akomolafe, O., & Agu, M. U. (2019c). Advances in financial resilience through integrated governance and compliance strategies. *Iconic Research and Engineering Journals*, 2(10), 607–620.

- 27) Akomolafe, O., Agu, M. U., & Bello, A. (2023a). A conceptual model for implementing risk-based auditing in strategic financial management. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), 2274–2286. <https://doi.org/10.62225/2583049X.2023.3.6.5359>
- 28) Akomolafe, O., Agu, M. U., & Bello, A. (2023b). A quantitative conceptual model for assessing compliance risk in emerging economies. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), 2287–2296. <https://doi.org/10.62225/2583049X.2023.3.6.5360>
- 29) Akomolafe, O., Agu, M. U., & Bello, A. (2025a). A conceptual model for advancing risk governance through data-driven compliance analytics in financial institutions. *Journal of Accounting and Financial Management*, 11(11), 211–228. <https://doi.org/10.56201/jafm.vol.11.no11.2025.pg211.228>
- 30) Akomolafe, O., Olaogun, B. O., Adesuyi, M. O., Ndukwe, V. U., & Sakyi, J. K. (2024a). Scalable blockchain payment gateway architecture model for enterprise-grade adoption. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(1), 1552–1568. <https://doi.org/10.62225/2583049X.2024.4.1.5280>
- 31) Akomolafe, O., Olaogun, B. O., Adesuyi, M. O., Ndukwe, V. U., & Sakyi, J. K. (2025b). Collaborative governance framework for secure cross-border payment data sharing. *International Journal of Advanced Multidisciplinary Research and Studies*, 5(6), 849–865. <https://doi.org/10.62225/2583049X.2025.5.6.5283>
- 32) Akomolafe, O., Olaogun, B. O., Adesuyi, M. O., Ndukwe, V. U., & Sakyi, J. K. (2025c). Global standardization model for next-generation blockchain interoperability in payment systems. *International Journal of Multidisciplinary Research and Growth Evaluation*, 6(6), 820–834. <https://doi.org/10.54660/IJMRGE.2025.6.6.820-834>
- 33) Aliliele, C., Mbonu, I. S., & Iwuanyanwu, U. (2023a). A conceptual framework for continuous cloud misconfiguration monitoring and enterprise risk mitigation strategies. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 9(10), 373–394. <https://doi.org/10.32628/CSEIT2361071>
- 34) Aliliele, C., Mbonu, I. S., & Iwuanyanwu, U. (2023b). A review of API governance and risk prioritization frameworks in modern financial institutions. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 9(10), 395–433. <https://doi.org/10.32628/CSEIT2361072>
- 35) Aliliele, C., Mbonu, I. S., & Iwuanyanwu, U. (2024a). A conceptual framework for enterprise data sensitivity classification and regulatory traceability mechanisms. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3103–3124. <https://doi.org/10.62225/2583049X.2024.4.6.5991>
- 36) Aliliele, C., Mbonu, I. S., Uzoka, E., & Iwuanyanwu, U. (2025a). A review of AI assisted continuous auditing systems in technology risk and cybersecurity oversight. *Gyanshauryam, International Scientific Refereed Research Journal*, 8(4), 210–250. <https://doi.org/10.32628/GISRRJ258369>
- 37) Aliliele, C., Mbonu, I. S., Uzoka, E., & Iwuanyanwu, U. (2025b). Advances in data lakehouse governance architectures for enterprise data loss prevention and compliance assurance. *Shodhshauryam, International Scientific Refereed Research Journal*, 8(4), 171–213. <https://doi.org/10.32628/SHISRRJ258474>
- 38) Amadi, C. (2025). Accelerating national defense: Using large language models (LLM) and NLP for real-time semantic correlation and de-duplication of shared threat indicators [Working paper]. SSRN. <https://ssrn.com/abstract=5940895>
- 39) Amayo, E. B., Owulade, O. A., & Isi, L. R. (2023a). Optimizing project governance in multinational infrastructure projects: Insights from General Electric's global operations. *International Journal of Multidisciplinary Research and Growth Evaluation*, 4(1), 975–983. <https://doi.org/10.54660/IJMRGE.2023.4.1.975-983>
- 40) Amayo, E. B., Owulade, O. A., & Isi, L. R. (2024a). Best practices for managing data center lifecycle projects: Ensuring security, efficiency, and compliance in U.S. enterprises. *Iconic Research and Engineering Journals*, 7(7), 598–617.
- 41) Amayo, E. B., Owulade, O. A., & Isi, L. R. (2024b). Effective project governance in multinational infrastructure projects: A case study from General Electric. *Iconic Research and Engineering Journals*, 8(5), 1305–1324.
- 42) Amayo, E. B., Owulade, O. A., & Isi, L. R. (2025c). Risk management and compliance in healthcare project delivery: Lessons from West Africa. *Engineering and Technology Journal*, 10(3), 4242–4255. <https://doi.org/10.47191/etj/v10i03.33>
- 43) Annan, A. O. (2022). Data privacy governance models for cross-border digital platforms. *International Journal of Multidisciplinary Research and Growth Evaluation*, 3(6), 949–964.
- 44) Annan, A. O. (2024). Algorithmic accountability and trade secret protection in artificial intelligence. *Shodhshauryam, International Scientific Refereed Research Journal*, 7(5), 315–347.
- 45) Annan, A. O. (2025a). Automated decision-making and anti-discrimination compliance under U.S. law. *International Journal of Advanced Multidisciplinary Research and Studies*, 5(2), 2522–2540.
- 46) Annan, A. O. (2025b). Cybersecurity compliance as a source of competitive advantage in technology markets. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 11(5), 456–487.
- 47) Arumosoye, O. M., & Obriki, O. D. (2020). A Governance-Oriented Conceptual Model for Contractor Safety Performance in Multi-Contract Industrial Projects. *International Journal of Multidisciplinary Research and Growth Evaluation*, 1(5), 728–740. <https://doi.org/10.54660/IJMRGE.2020.1.5.728-740>
- 48) Asiedu, W., Quainoo, R., & Asiedu, A. (2026). Predictive power management for intermittent IoT devices using lightweight machine learning under uncertain energy harvest. *Gulf Journal of Engineering and Technology*, 2(4), 108–118.
- 49) Atakpa, M. I. (2021). A review of predictive analytics models for anti-money laundering detection in commercial banking systems. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 7(2), 778–804. <https://doi.org/10.32628/CSEIT2064813>

- 50) Atakpa, M. I., Abetoh, N. F., & Akeju, B. (2024a). Advances in artificial intelligence for healthcare payment fraud detection: A review of NHS applications. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 10(4), 1161–1194. <https://doi.org/10.32628/CSEIT26123240>
- 51) Atakpa, M. I., & Abolaji, T. O. (2022). A privacy-preserving data architecture model for regulated industry analytics under GDPR and HIPAA compliance. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 8(3), 781–808. <https://doi.org/10.32628/CSEIT22558>
- 52) Atakpa, M. I., & Fobellah, A. N. (2023). Anomaly detection in financial time-series data: A conceptual model for healthcare and banking applications. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 9(2), 967–997. <https://doi.org/10.32628/CSEIT2342441>
- 53) Atkinson, R. C., & Shiffrin, R. M. (1968). Human memory: A proposed system and its control processes. *Psychology of Learning and Motivation*, 2, 89–195.
- 54) Badmus, O., Dosunmu, A. A., & Anunagba, C. O. (2025). A governance framework for AI-assisted CRM workflows: Agentforce deployment, risk management, and organizational readiness in enterprise Salesforce environments. *International Journal of Scientific Research in Humanities and Social Sciences*, 2(1), 59–80. <https://doi.org/10.32628/IJSRHSS>
- 55) Badmus, O., Dosunmu, A. A., & Ozowara, D. E. (2019b). A governance framework for Salesforce platform management in regulated healthcare environments. *IRE Journals*, 3(6).
- 56) Badmus, O., Jooda, D., & Anunagba, C. O. (2024). A privacy-by-design framework for role-based security architecture in Salesforce environments handling sensitive personal data. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3244–3254. <https://doi.org/10.62225/2583049X.2024.4.6.6243>
- 57) Badmus, O., Jooda, D., Ozowara, D. E., & Anunagba, C. O. (2022a). A framework for Git-based version control and code review governance in Salesforce development teams. *Shodhshauryam, International Scientific Refereed Research Journal*, 5(2), 473–493.
- 58) Barnum, S. (2014). Standardizing cyber threat intelligence information with the Structured Threat Information eXpression (STIX). MITRE Corporation.
- 59) Bello, A. D., Elebe, O., Hammed, N. I., Omoegun, G. O., & Fadayomi, O. (2024). A cybersecurity risk management and regulatory compliance framework for financial institutions. *Iconic Research and Engineering Journals*. <https://www.irejournals.com/paper-details/1713553>
- 60) Berman, D. S., Buczak, A. L., Chavis, J. S., & Corbett, C. L. (2019). A survey of deep learning methods for cyber security. *Information*, 10(4), 122.
- 61) Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., ... Liang, P. (2021). On the opportunities and risks of foundation models. *arXiv:2108.07258*.
- 62) Borah, S., Aliliele, K. C., Rakshit, S., & Vajjhala, N. R. (2022). Applications of artificial intelligence in software testing. In P. K. Mallick et al. (Eds.), *Cognitive informatics and soft computing (Lecture Notes in Networks and Systems, Vol. 375, pp. 727–736)*. Springer. https://doi.org/10.1007/978-981-16-8763-1_60
- 63) Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., ... Amodei, D. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901.
- 64) 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.
- 65) Connolly, J., Davidson, M., & Schmidt, C. (2014). The Trusted Automated eXchange of Indicator Information (TAXII). MITRE Corporation.
- 66) Dagodzo, D., Patrick, M. C. A., & Aliliele, C. (2022). AI and deep learning for vegetation classification in power corridor management: A review. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 8(1), 668–698. <https://doi.org/10.32628/CSEIT2281227>
- 67) Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of NAACL-HLT*, 4171–4186. <https://doi.org/10.18653/v1/N19-1423>
- 68) Dogbatsey, E. A., & Ebhojie, O. (2018). Budget compliance and statutory reporting in Sub-Saharan Africa: A systematic review of frameworks, gaps, and reform pathways. *Zenodo*. <https://doi.org/10.5281/zenodo.20446859> [Originally numbered 2(6); journal name not captured in source.]
- 69) Dogbatsey, E. A., Oyeleye, A. O., & Ebhojie, O. (2026). Donor evidence standards and corrective action plan governance in publicly funded capital projects: A public sector accountability framework. *International Journal of Applied Research in Social Sciences*, 8(5), 107–126. <https://doi.org/10.51594/ijarss.v8i5.2266>
- 70) Dosunmu, A. A., & Ogundele, P. O. (2019). Security audit and enterprise risk assessment frameworks for resilient information systems. *IRE Journals*, 3(5).
- 71) Dosunmu, A. A., & Ogundele, P. O. (2020). Intrusion detection and prevention models for enhancing organizational cyber defense effectiveness. *IRE Journals*, 4(6).
- 72) Dosunmu, A. A., & Ogundele, P. O. (2021). Incident response and digital forensics strategies for rapid cyber attack containment. *Gyanshauryam, International Scientific Refereed Research Journal*, 4(4), 239–258.
- 73) Dosunmu, A. A., & Ogundele, P. O. (2022). Threat intelligence integration frameworks supporting proactive enterprise cybersecurity decision making. *Gyanshauryam, International Scientific Refereed Research Journal*, 5(3), 397–416.
- 74) Dosunmu, A. A., & Ogundele, P. O. (2023). Cyber threat actor analysis models for strategic enterprise security planning. *Shodhshauryam, International Scientific Refereed Research Journal*, 6(5), 513–531. <https://doi.org/10.32628/SHISRRJ>

- 75) Dosunmu, A. A., & Ogundele, P. O. (2024a). Breach and attack simulation frameworks for continuous validation of enterprise security controls. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 10(3), 1100–1119. <https://doi.org/10.32628/IJSRCSEIT>
- 76) Dosunmu, A. A., & Ogundele, P. O. (2024b). Cyber risk quantification models for prioritizing enterprise security investment decisions. *International Journal of Multidisciplinary Research and Growth Evaluation*, 5(6), 1777–1785.
- 77) Dosunmu, A. A., & Ogundele, P. O. (2024c). Enterprise scale continuous security validation models for regulated digital infrastructures. *International Journal of Scientific Research in Humanities and Social Sciences*, 1(2), 929–945. <https://doi.org/10.32628/IJSRSSH>
- 78) Dosunmu, A. A., & Ogundele, P. O. (2024d). Threat-informed defence engineering models for measuring security control effectiveness at scale. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 2847–2858.
- 79) Dosunmu, A. A., & Ogundele, P. O. (2025a). Adversary simulation design frameworks for proactive cyber defense in complex environments. *International Journal of Computer Science and Mathematical Theory*, 11(12), 193–209. <https://doi.org/10.56201/ijcsmt.vol.11.no12.2025.pg193.209>
- 80) Dosunmu, A. A., & Ogundele, P. O. (2025b). Cyber defense performance measurement frameworks for executive and board-level security governance. *Computer Science and IT Research Journal*, 6(11), 895–913.
- 81) Dosunmu, A. A., & Ogundele, P. O. (2025c). Security orchestration and automation models for accelerating incident detection and response. *Computer Science and IT Research Journal*, 6(11), 878–894. <https://doi.org/10.51594/csitrj.v6i11.2163>
- 82) Dosunmu, A. A., & Ogundele, P. O. (2026a). Deception based defense architectures for disrupting advanced persistent threat operations. *Engineering and Technology Journal*, 11(1), 8475–8487. <https://doi.org/10.47191/etj/v11i01.07>
- 83) Dosunmu, A. A., & Ogundele, P. O. (2026b). Integrated threat intelligence automation and simulation models for next generation cyber defense. *Engineering and Technology Journal*, 11(1), 8451–8462. <https://doi.org/10.47191/etj/v11i01.05>
- 84) Ebhohie, O., Dogbatsey, E. A., & Oyeleye, A. O. (2023a). Audit liaison, corrective action planning, and control compliance in multinational organisations: A systematic literature review. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), 2839–2850. <https://doi.org/10.62225/2583049X.2023.3.6.6200>
- 85) Eboh, E. E., & Aliliele, C. (2025b). Predictive Analytics Systems for Investment Risk Monitoring in SME FinTech Companies and Cross-Border Capital Markets. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 11(2), 3969–4010. <https://doi.org/10.32628/CSEIT25113398>
- 86) Edivri, J., Ogbale, J. I., Okoruwa, P. O., Fadayomi, O., Abolaji, T. O., & Akeju, B. (2026). Conceptual model for integrated human and machine identity governance in cloud-based security architectures [Manuscript; venue not yet indexed].
- 87) Ekechi, N. V., Ozowara, D. E., & Anunagba, C. O. (2026). Conceptual framework for AI governance, data privacy compliance, and financial sustainability in digital health. *Computer Science and IT Research Journal*, 7(4), 275–299.
- 88) Elebe, O., Hammed, N. I., Omoegun, G. O., Fadayomi, O., & Bello, A. D. (2023). An advanced machine learning model for detecting synthetic identity fraud in e-commerce platforms. *ResearchGate*. <https://www.researchgate.net/profile/Adepeju-Bello-2>
- 89) Endsley, M. R. (1995). Toward a theory of situation awareness in dynamic systems. *Human Factors*, 37(1), 32–64.
- 90) Essandoh, S., Sakyi, J. K., Ibrahim, A. K., Okafor, C. M., & Wedraogo, L. (2025). Artificial intelligence and the future of work: Impacts on employment and job roles. *International Journal of Multidisciplinary Futuristic Development*, 6(1), 31–41. <https://doi.org/10.54660/IJMFD.2025.6.1.31-41>
- 91) Eyetsemitan, R. A., Ambali, K. B., Oyeleye, A. O., & Fadayomi, O. (2020). Multi-Stakeholder Governance Alignment in Joint Venture Operations: A Conceptual Framework for Coordinating Business Processes in Highly Regulated Environments. *IRE Journals*, 4(4), 418–441. <https://doi.org/10.64388/IREV4I4-1716955>
- 92) Eyetsemitan, R. A., Ambali, K. B., Oyeleye, A. O., & Fadayomi, O. (2021). Translating Tax and Regulatory Requirements into SME Compliance Workflows: A Conceptual Framework for Implementing IAS 12, VAT, PAYE, and Withholding Tax. *IRE Journals*, 4(11), 621–641. <https://doi.org/10.64388/IREV4I11-1716956>
- 93) Eyetsemitan, R. A., Ambali, K. B., Oyeleye, A. O., & Fadayomi, O. (2025). An Integrated Lean-Digital Framework for Scaling Small Business Operations: Synthesizing SOP Design, Automation, Compliance, and Change Management. *International Journal of Multidisciplinary Research and Growth Evaluation*, 6(6), 1341–1360. <https://doi.org/10.54660/IJMRGE.2025.6.6.1341-1360>
- 94) Fadayomi, O., Abolaji, T. O., Edivri, J., Ogbale, J. I., Okoruwa, P. O., & Akeju, B. (2019). Risk-based cybersecurity assurance and data availability: Limitations, advances and future research opportunities. *Iconic Research and Engineering Journals*, 2(12), 602–617. <https://doi.org/10.64388/IREV2I12-1713779>
- 95) Fadayomi, O., Bello, A. D., Elebe, O., Hammed, N. I., & Omoegun, G. O. (2021). An integrated cybersecurity and anti-money laundering governance framework for financial crime prevention. *ResearchGate*. <https://www.researchgate.net/profile/Adepeju-Bello-2>
- 96) Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., ... Vayena, E. (2018). AI4People: An ethical framework for a good AI society. *Minds and Machines*, 28(4), 689–707.
- 97) Green, D. M., & Swets, J. A. (1966). *Signal detection theory and psychophysics*. Wiley.
- 98) Hevner, A. R., March, S. T., Park, J., & Ram, S. (2004). Design science in information systems research. *MIS Quarterly*, 28(1), 75–105.
- 99) Hussain, N., & Adebayo, A. (2026). The role of artificial intelligence in strengthening modern cybersecurity systems. *World Journal of Innovation and Modern Technology*, 10(6), 125–181. <https://doi.org/10.56201/wjimt.v10.no6.2026.pg125.181>
- 100) Hutchins, E. M., Cloppert, M. J., & Amin, R. M. (2011). Intelligence-driven computer network defense informed by analysis of adversary campaigns and intrusion kill chains. *Leading Issues in Information Warfare & Security Research*, 1(1), 80–106.
- 101) Jimoh, H. O., Abolle-Okoyeagu, C. J., Ahmed, M. O., & Lawal, N. O. (2023a). Advancing security in IoT-driven critical infrastructure: A focus on smart transportation system. *American Journal of Engineering Research*, 12(12), 33–46.

- 102) Jimoh, H. O., & Ahmed, M. O. (2024). Analyzing Network Time Protocol (NTP) based amplification DDoS attack and its mitigation techniques. *Journal of Digital Innovations and Contemporary Research in Science, Engineering and Technology*, 12(2), 17–24. <https://doi.org/10.22624/AIMS/DIGITAL/V11N2P2x>
- 103) Jimoh, H. O., Ahmed, M. O., & Fagbade, M. O. (2023b). The role of frameworks in cybersecurity governance. *Journal of Behavioural Informatics, Digital Humanities and Development Research*, 9(4), 7–16. <https://doi.org/10.22624/AIMS/BHI/V9N4P2>
- 104) Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399.
- 105) Jooda, D., & Smart, B. D. (2024b). Risk-adaptive cybersecurity compliance automation model (RACCA): Continuous assessment, dynamic threshold management, and integrated remediation for enterprise security programs. *Shodhshauryam, International Scientific Refereed Research Journal*, 7(2), 360–374. <https://doi.org/10.32628/SHISRRJ>
- 106) Kaplan, J., McCandlish, S., Henighan, T., Brown, T. B., Chess, B., Child, R., ... Amodei, D. (2020). Scaling laws for neural language models. *arXiv:2001.08361*.
- 107) Komi, N. M., & Ganiu, O. S. (2023). Edge Intelligence for Resilient Microgrid Control: Advances, Energy Sovereignty, and Open Challenges. *Gyanshauryam, International Scientific Refereed Research Journal*, 6(3), 525–586. <https://doi.org/10.32628/GISRRJ236339>
- 108) Kumuyi, O., Akeju, B., Uzoka, E., & Ozowara, D. E. (2023). Framework for blockchain-based cross-border data exchange and regulatory transparency. *International Journal of Multidisciplinary Futuristic Development*, 4(1), 99–113.
- 109) Kumuyi, O., Uzoka, E., Akeju, B., & Ozowara, D. E. (2024a). Framework for privacy-focused digital identity verification supporting financial inclusion in Africa. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3150–3162. <https://doi.org/10.62225/2583049X.2024.4.6.6005>
- 110) Kumuyi, O., Uzoka, E., Akeju, B., & Ozowara, D. E. (2024b). Architecture for machine learning-enabled predictive energy management using IoT sensor networks. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3138–3149. <https://doi.org/10.62225/2583049X.2024.4.6.6004>
- 111) Ladapo, O. O., Dosunmu, A. A., Jooda, D., & Abolaji, T. O. (2018). Lessons learned from offline assessment of security-critical systems: The case of Microsoft Active Directory. *IRE Journals*, 2(6). <https://doi.org/10.64388/IREV2I6-1717205>
- 112) Ladapo, O. O., Dosunmu, A. A., Jooda, D., & Abolaji, T. O. (2022a). Human-in-the-loop machine learning: A state of the art. *Journal of Frontiers in Multidisciplinary Research*, 3(1), 656–669.
- 113) Ladapo, O. O., Dosunmu, A. A., Jooda, D., & Abolaji, T. O. (2023a). Active Directory attacks steps, types, and signatures. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), 2863–2874. <https://doi.org/10.62225/2583049X.2023.3.6.6204>
- 114) Ladapo, O. O., Dosunmu, A. A., Jooda, D., & Abolaji, T. O. (2024a). Integrated network and security operation center: A systematic analysis. *International Journal of Multidisciplinary Futuristic Development*, 5(1), 65–80.
- 115) Ladapo, O. O., Jooda, D., Dosunmu, A. A., & Abolaji, T. O. (2023d). Systematic literature review on security access control policies and techniques based on privacy requirements in a BYOD environment. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), 2890–2904. <https://doi.org/10.62225/2583049X.2023.3.6.6206>
- 116) Ladapo, O. O., Jooda, D., Dosunmu, A. A., & Abolaji, T. O. (2024b). Keeping humans in the loop: Human-centered automated annotation with generative AI. *International Journal of Multidisciplinary Futuristic Development*, 5(1), 81–95.
- 117) Ladapo, O. O., Jooda, D., Dosunmu, A. A., & Abolaji, T. O. (2026). On the disagreement problem in human-in-the-loop federated machine learning. *International Journal of Multidisciplinary Research and Growth Evaluation*, 7(3), 178–192.
- 118) Lawal, O. A., & Odunle, T. E. (2018b). A review and conceptual framework for tax governance and cross-border compliance analytics. *Iconic Research and Engineering Journals*, 2(5).
- 119) Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., ... Kiela, D. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. *Advances in Neural Information Processing Systems*, 33, 9459–9474.
- 120) Liadi, K. O. (2024b). Conceptualizing a governance reform impact model for Nigeria's peacekeeping missions in post-conflict states. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(1), 1622–1636.
- 121) Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., Levy, O., Lewis, M., Zettlemoyer, L., & Stoyanov, V. (2019). RoBERTa: A robustly optimized BERT pretraining approach. *arXiv:1907.11692*.
- 122) Mayo, W., Ogbale, J. I., Okoruwa, P. O., & Babatope, O. M. (2021). Designing an AI-predictive maintenance model for e-commerce systems using machine learning and cloud analytics. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 7(5), 416–440. <https://doi.org/10.32628/IJSRCSEIT>
- 123) Mbonu, I. S., Aliliele, C., Iwuanyanwu, U., & Oluoha, O. M. (2018). A conceptual framework for legal and ethical risk modeling in enterprise data protection governance systems. *Iconic Research and Engineering Journals*, 2(2), 207–226. <https://doi.org/10.64388/IREV2I2-1714911>
- 124) Mbonu, I. S., Aliliele, C., Iwuanyanwu, U., & Uzoka, E. (2020a). A review of identity and access management integration strategies in hybrid and multi cloud environments. *International Journal of Multidisciplinary Research and Growth Evaluation*, 1(5), 795–810. <https://doi.org/10.54660/IJMRGE.2020.1.5.795-810>
- 125) Mbonu, I. S., Aliliele, C., Iwuanyanwu, U., & Uzoka, E. (2021b). Advances in artificial intelligence techniques for secure software testing and automated regression control mechanisms. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 7(5), 468–496. <https://doi.org/10.32628/CSEIT217565>
- 126) Mbonu, I. S., Aliliele, C., Iwuanyanwu, U., & Uzoka, E. (2022a). A conceptual framework for AI enabled IT general controls and SOX audit automation processes. *Gyanshauryam, International Scientific Refereed Research Journal*, 5(5), 384–414. <https://doi.org/10.32628/GISRRJ2256239>

- 127) Mbonu, I. S., Aliliele, C., Uzoka, E., & Oluoha, O. M. (2019a). A review of comparative data protection regulations and secure cloud implementation strategies across jurisdictions. *Iconic Research and Engineering Journals*, 2(9), 482–501. <https://doi.org/10.64388/IREV219-1714912>
- 128) Mbonu, I. S., Iwuanyanwu, U., Aliliele, C., & Uzoka, E. (2020c). Advances in infrastructure as code governance for secure Terraform based enterprise cloud deployments. *International Journal of Multidisciplinary Research and Growth Evaluation*, 1(5), 811–828. <https://doi.org/10.54660/IJMRGE.2020.1.5.811-828>
- 129) Mbonu, I. S., Iwuanyanwu, U., Aliliele, C., & Uzoka, E. (2021c). A review of VoIP forensic analytics models for financial fraud detection and regulatory compliance monitoring. *International Journal of Multidisciplinary Research and Growth Evaluation*, 2(6), 711–730. <https://doi.org/10.54660/IJMRGE.2021.2.6.711-730>
- 130) Mbonu, I. S., Iwuanyanwu, U., Aliliele, C., & Uzoka, E. (2022b). A review of data protection impact assessment models in multi cloud financial infrastructure systems. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 8(1), 589–623. <https://doi.org/10.32628/CSEIT25442>
- 131) Mbonu, I. S., Iwuanyanwu, U., Aliliele, C., & Uzoka, E. (2022c). Advances in cloud identity and access governance optimization in large scale AWS enterprise environments. *Shodhshauryam, International Scientific Refereed Research Journal*, 5(3), 403–438. <https://doi.org/10.32628/SHISRRJ225490>
- 132) Mbonu, I. S., Iwuanyanwu, U., Uzoka, E., & Oluoha, O. M. (2019b). Advances in enterprise log analytics and automated incident response architectures using Python and SIEM platforms. *Iconic Research and Engineering Journals*, 3(2), 1000–1019. <https://doi.org/10.64388/IREV312-1714915>
- 133) Miller, G. A. (1956). The magical number seven, plus or minus two: Some limits on our capacity for processing information. *Psychological Review*, 63(2), 81–97.
- 134) Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S., & Floridi, L. (2016). The ethics of algorithms: Mapping the debate. *Big Data & Society*, 3(2), 1–21.
- 135) Morgan, D. L. (2014). Pragmatism as a paradigm for social research. *Qualitative Inquiry*, 20(8), 1045–1053.
- 136) Nnaji, N., & Akinlolu, V. S. (2024). Enhancing clinical record protection through structured health data security frameworks in regulated healthcare environments. *International Journal of Scientific Research in Humanities and Social Sciences*.
- 137) Obogo, S. F., Arumosoye, O. M., & Obriki, O. D. (2020a). Advances in Internal QHSE Audit Systems for Industrial Engineering Operations. *Iconic Research and Engineering Journals*, 4(4), 399–417. <https://doi.org/10.64388/IREV4I4-1715499>
- 138) Obogo, S. F., Arumosoye, O. M., & Obriki, O. D. (2020b). Conceptual Risk Management Model for Heavy Lifting and Crane Installation Engineering Operations. *Shodhshauryam, International Scientific Refereed Research Journal*, 3(4), 122–144. <https://doi.org/10.32628/SHISRRJ214441>
- 139) Obogo, S. F., Obriki, O. D., & Arumosoye, O. M. (2022). Conceptual Safety Governance Model for Large Commercial Facility Operations. *Gyanshauryam, International Scientific Refereed Research Journal*, 5(6), 383–410. <https://doi.org/10.32628/GISRRJ225639>
- 140) Obogo, S. F., Uduokhai, D. O., & Ozobu, C. O. (2021c). Review of Contractor Safety Compliance Systems in Infrastructure Engineering Projects. *International Journal of Scientific Research in Civil Engineering*, 5(4), 76–100. <https://doi.org/10.32628/IJSRCE215412>
- 141) Obriki, O. D., & Arumosoye, O. M. (2024a). Conceptual Governance Framework for Subcontractor Safety Performance Management. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3020–3033. <https://doi.org/10.62225/2583049X.2024.4.6.5894>
- 142) Odejobi, O. D., Okonkwo, C. S., Ahiaekwe Patrick, M. C., Okeke, O. T., & Mayo, W. (2025). AI-augmented secure software engineering: Leveraging deep learning for autonomous threat detection and mitigation. *International Journal of Engineering and Modern Technology*, 11(12), 101–121. <https://doi.org/10.56201/ijemt.vol.11.no12.2025.pg101.121>
- 143) Ogbole, J. I., Okoruwa, P. O., Fadayomi, O., Abolaji, T. O., Edvri, J., & Akeju, B. (2021b). Conceptual model for identity-centric zero trust architecture in enterprise security governance. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 7(5), 393–415. <https://doi.org/10.32628/IJSRCSEIT217562>
- 144) Ogbole, J. I., Okoruwa, P. O., Fadayomi, O., Akeju, B., Edvri, J., & Abolaji, T. O. (2025). Security analytics and digital forensics for enterprise risk management: Advances and practical implications. *International Journal of Advanced Multidisciplinary Research and Studies*, 5(6), 2017–2028.
- 145) Oghenemaiga, E., Basnet, A., & Anene, U. N. (2024). Predictive analytics applications for forecasting traffic patterns across owned and operated media platforms. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6). <https://doi.org/10.62225/2583049X.2024.4.6.5668>
- 146) Ogunwola, T. A., & Alozie, C. (2026a). Architecting secure and compliant distributed healthcare networks: Operational approaches to health insurance portability and accountability act and health information trust alliance alignment. *International Journal of Computer Applications*, 187(111), 1–6.
- 147) Ogunwola, T. A., Owoyemi, T. A., & Iyanda, C. (2026). Cybersecurity and network integration risk management during corporate acquisitions and infrastructure consolidation. *Computer Science & IT Research Journal*, 7(6), 389–404. <https://doi.org/10.51594/csit.rj.v7i6.2305>
- 148) Okonkwo, C. S., Ogunwola, O., Okeke, O. T., & Mayo, W. (2019). Conceptual Framework for Cost Reduction Through Contract Negotiation and Vendor Governance. *IRE Journals*, 2(9), 468–482. <https://doi.org/10.64388/IREV219-1713121>
- 149) Okoruwa, P. O., Fadayomi, O., Abolaji, T. O., Edvri, J., Ogbole, J. I., & Akeju, B. (2024). Enterprise cybersecurity trends and threat evolution: Advances and emerging research directions. *Global Multidisciplinary Perspectives Journal*, 1(6), 194–204. <https://doi.org/10.54660/GMPJ.2024.1.6.194-204>

- 150) Okoruwa, P. O., Fadayomi, O., Akeju, B., Edivri, J., Ogbole, J. I., & Abolaji, T. O. (2020). Conceptual model for privacy-centric security engineering in digital and cloud computing systems. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 7(5), 371–389.
- 151) Okoruwa, P. O., Fadayomi, O., Bello, A. D., Elebe, O., Hamed, N. I., & Omoegun, G. O. (2025). An artificial intelligence-driven financial crime investigation framework for analyst decision support. *International Journal of Multidisciplinary Research and Growth Evaluation*, 6(6). <https://www.allmultidisciplinaryjournal.com/article/5453/>
- 152) Olaogun, B. O., Adesuyi, M. O., Akomolafe, O., Ndukwe, V. U., & Sakyi, J. K. (2023). Regulatory compliance adaptation model for cross-border payment product redesign. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(1), 1651–1662. <https://doi.org/10.62225/2583049X.2023.3.1.5279>
- 153) Onche, V. O., Adegbite, M. P., & Ogbonna, C. S. (2026). Human factors in information security: A capability framework for administrative professionals in the digital workplace. *World Journal of Innovation and Modern Technology*, 10(5), 154–179. <https://doi.org/10.56201/wjimt.v10.no5.2026.pg154.179>
- 154) Onche, V. O., Ogbonna, C. S., & Adegbite, M. P. (2024). Effectiveness of cybersecurity awareness training on insider-threat behavior among university administrators: An integrative academic review. *International Journal of Computer Science and Mathematical Theory*, 10(2), 117–144. <https://doi.org/10.56201/ijcsmt.v10.no2.2024.pg117.144>
- 155) Onwubuya, L. I., Fawehinmi, Y. O., Sunday, D., Igwenagu, M. O., Adeniran, A. O., & Agyapong, E. A. (2026). Adaptive learning systems powered by artificial intelligence for STEM education improvement. *Asian Journal of Education and Social Studies*, 52(6), 553–566. <https://doi.org/10.9734/ajess/2026/v52i63115>
- 156) Orise, C., & Niniola, S. (2023). Privacy, safeguarding, and ethical data practices in healthcare EdTech: A conceptual model for compliance and trust. *Shodhshauryam, International Scientific Refereed Research Journal*, 6(1), 507–529.
- 157) Orise, C., & Niniola, S. (2025). Data governance in digital health learning systems: A systematic review of implementation frameworks and best practices. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 11(5), 520–538.
- 158) Oteri, O., & Edivri, J. (2026). An operational reliability and service assurance framework for enterprise IT systems supporting large user populations. *Gulf Journal of Engineering & Technology*, 2(1), 1–19. <https://doi.org/10.51594/gjet.v2i1.205>
- 159) Ouyang, L., Wu, J., Jiang, X., Almeida, D., Wainwright, C. L., Mishkin, P., ... Lowe, R. (2022). Training language models to follow instructions with human feedback. *Advances in Neural Information Processing Systems*, 35, 27730–27744.
- 160) Oyeleye, A. O., Dogbatsey, E. A., & Ebhojie, O. (2025). Embedding automated close controls in public sector ERP systems: A conceptual model for audit readiness and financial reporting quality. *Zenodo*. <https://doi.org/10.5281/zenodo.20097950> [Originally numbered 2(1), 110–130; journal name not captured in source.]
- 161) Ozowara, D. E., Adebayo, A., & Anunagba, C. O. (2022). A systematic review of cybersecurity investments and their impact on healthcare financial performance. *Shodhshauryam, International Scientific Refereed Research Journal*, 5(1), 404–427. <https://doi.org/10.32628/SHISRRJ>
- 162) Ozowara, D. E., Adebayo, A., & Anunagba, C. O. (2025a). Advanced conceptual model for strengthening audit quality using data analytics across financial institutions. *International Journal of Advanced Multidisciplinary Research and Studies*, 5(6), 2246–2268. <https://doi.org/10.62225/2583049X.2025.5.6.6050>
- 163) Peffers, K., Tuunanen, T., Rothenberger, M. A., & Chatterjee, S. (2007). A design science research methodology for information systems research. *Journal of Management Information Systems*, 24(3), 45–77.
- 164) Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., & Sutskever, I. (2019). Language models are unsupervised multitask learners. *OpenAI Technical Report*.
- 165) Raffel, C., Shazeer, N., Roberts, A., Lee, K., Narang, S., Matena, M., Zhou, Y., Li, W., & Liu, P. J. (2020). Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of Machine Learning Research*, 21(140), 1–67.
- 166) Rahman, M. R., Mahdavi-Hezaveh, R., & Williams, L. (2020). A literature review on mining cyberthreat intelligence from unstructured texts. *Proceedings of the IEEE International Conference on Data Mining Workshops*, 516–525.
- 167) Sanni, J. O. (2026a). Artificial intelligence driven demand forecasting models addressing volatility in complex service organizations. *Computer Science & IT Research Journal*, 7(1), 1–26. <https://doi.org/10.51594/csitrj.v7i1.2170>
- 168) Sarker, I. H., Kayes, A. S. M., Badsha, S., Alqahtani, H., Watters, P., & Ng, A. (2020). Cybersecurity data science: An overview from machine learning perspective. *Journal of Big Data*, 7(1), 41.
- 169) Sauerwein, C., Sillaber, C., Musmann, A., & Breu, R. (2017). Threat intelligence sharing platforms: An exploratory study of software vendors and research perspectives. *Proceedings of the International Conference on Wirtschaftsinformatik*, 837–851.
- 170) Shannon, C. E. (1948). A mathematical theory of communication. *Bell System Technical Journal*, 27(3), 379–423.
- 171) Simon, H. A. (1955). A behavioral model of rational choice. *Quarterly Journal of Economics*, 69(1), 99–118.
- 172) Skopik, F., Settanni, G., & Fiedler, R. (2016). A problem shared is a problem halved: A survey on the dimensions of collective cyber defense through security information sharing. *Computers & Security*, 60, 154–176.
- 173) Smart, B. D. (2020). A quantitative cyber risk valuation model for board-level decision making in critical infrastructure. *International Journal of Multidisciplinary Research and Growth Evaluation*, 1(5), 984–991. <https://doi.org/10.54660/IJMRGE.2020.1.5.984-991>
- 174) Smart, B. D., & Jooda, D. (2023a). A continuous compliance and real-time audit readiness architecture for regulated enterprises. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), 2911–2916. <https://doi.org/10.62225/2583049X.2023.3.6.6229>

- 175) Smart, B. D., & Jooda, D. (2023b). A unified multi-framework compliance optimization model integrating ISO 27001, NIST Cybersecurity Framework, and Cybersecurity Maturity Model Certification. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*. <https://doi.org/10.32628/IJSRCSEIT>
- 176) Smart, B. D., & Jooda, D. (2024a). A cybersecurity maturity acceleration model for critical infrastructure systems. *Gyanshauryam, International Scientific Refereed Research Journal*, 7(6), 348–360. <https://doi.org/10.32628/GISRRJ>
- 177) Smart, B. D., & Jooda, D. (2024b). A national-scale AI governance and risk control framework for critical infrastructure protection. *Gyanshauryam, International Scientific Refereed Research Journal*, 7(6), 361–371. <https://doi.org/10.32628/GISRRJ>
- 178) Smart, B. D., & Jooda, D. (2025). Resilient security architecture for operational technology environments under advanced persistent threat conditions: A layered defense model for critical infrastructure protection. *International Journal of Engineering and Modern Technology*, 11(11), 147–159. <https://doi.org/10.56201/ijemt.vol.11.no11.2025.pg147.159>
- 179) Smart, B. D., & Jooda, D. (2026). Integrated incident response governance model (IIRG): Unifying strategic leadership, cross-sector threat intelligence, and AI-augmented response for enterprise cyber resilience. *Gulf Journal of Computer Sciences*, 2(1), 1–15. <https://doi.org/10.51594/gjcs.v2i1.210>
- 180) Strom, B. E., Applebaum, A., Miller, D. P., Nickels, K. C., Pennington, A. G., & Thomas, C. B. (2018). MITRE ATT&CK: Design and philosophy. MITRE Corporation.
- 181) Sun, N., Ding, M., Jiang, J., Xu, W., Mo, X., Tai, Y., & Zhang, J. (2023). Cyber threat intelligence mining for proactive cybersecurity defense: A survey and new perspectives. *IEEE Communications Surveys & Tutorials*, 25(3), 1748–1774.
- 182) Timmermans, S., & Tavory, I. (2012). Theory construction in qualitative research: From grounded theory to abductive analysis. *Sociological Theory*, 30(3), 167–186.
- 183) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2021a). Advances in Demand Forecasting: Machine Learning Algorithms for Revenue Projection and Financial Planning Accuracy. *Gyanshauryam, International Scientific Refereed Research Journal*, 4(1), 282–317. <https://doi.org/10.32628/GISRRJ120361>
- 184) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2022c). Conceptual model for predictive cost optimization: Machine learning applications in business transformation analysis. *International Journal of Economics and Business Management*, 8(6), 68–102. <https://doi.org/10.56201/ijebm.v8.no6.2022.pg68.102>
- 185) Tounsi, W., & Rais, H. (2018). A survey on technical threat intelligence in the age of sophisticated cyber attacks. *Computers & Security*, 72, 212–233. <https://doi.org/10.1016/j.cose.2017.09.001>
- 186) Trist, E. L., & Bamforth, K. W. (1951). Some social and psychological consequences of the longwall method of coal-getting. *Human Relations*, 4(1), 3–38.
- 187) Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.
- 188) Wagner, C., Dulaunoy, A., Wagener, G., & Iklody, A. (2016). MISP: The design and implementation of a collaborative threat intelligence sharing platform. *Proceedings of the ACM Workshop on Information Sharing and Collaborative Security*, 49–56.
- 189) Wagner, T. D., Mahbub, K., Palomar, E., & Abdallah, A. E. (2019). Cyber threat intelligence sharing: Survey and research directions. *Computers & Security*, 87, 101589. <https://doi.org/10.1016/j.cose.2019.101589>
- 190) Wei, J., Wang, X., Schuurmans, D., Bosma, M., Ichter, B., Xia, F., Chi, E., Le, Q., & Zhou, D. (2022). Chain-of-thought prompting elicits reasoning in large language models. *Advances in Neural Information Processing Systems*, 35, 24824–24837.