

Artificial Intelligence, Imagery Intelligence, and Early Warning Systems: A Strategic Intelligence Model for Addressing Contemporary Security Threats

Kharis Kurnia^{1*}, Apsas Saputra²

^{1,2} Universitas Bangka Belitung

Corresponding Author's e-mail : kharis-kurnia@ubb.ac.id



e-ISSN: 2964-0962

SEIKAT: Jurnal Ilmu Sosial, Politik dan Hukum

<https://ejournal.45mataram.ac.id/index.php/seikat>

Vol. 5, No. 4, Agustus 2026

Page: 2311 - 2322

Available at:

<https://ejournal.45mataram.ac.id/index.php/seikat/article/view/3556>

DOI:

<https://doi.org/10.55681/seikat.v5i4.3556>

Article History:

Received: 20-07-2026

Revised: 15-08-2026

Accepted: 15-08-2026

Abstract : *This study develops a conceptual model integrating artificial intelligence (AI), imagery intelligence (IMINT), and an early warning system (EWS) within a strategic intelligence framework to address contemporary security threats. Adopting a qualitative conceptual approach based on a systematic literature review (SLR), the article synthesizes recent scholarship on AI, intelligence collection, and security governance in order to explain how the three components can be conceptually integrated into a single analytical flow. The findings indicate that AI improves image-processing speed, expands change-detection capacity, and strengthens the accuracy of object identification, thereby enabling EWS to operate faster and more predictively than conventional, manually driven analytical processes. However, AI deployment also introduces risks of algorithmic bias, false positives, limited explainability, and potential image manipulation, making human oversight indispensable rather than optional. The article proposes an AI-IMINT-EWS model consisting of four interrelated layers, acquisition, machine analysis, human verification, and early warning, that positions AI as an analytical accelerator, human analysts as strategic verifiers, and EWS as the integrative node for early warning. This model offers a conceptual contribution to the development of modern intelligence practice that is fast, accountable, and adaptive, while also opening an agenda for future empirical validation.*

Keywords : *Artificial Intelligence; Imagery Intelligence; Early Warning System; Strategic Intelligence; Contemporary Security.*

Abstrak : Penelitian ini mengembangkan model konseptual yang mengintegrasikan kecerdasan buatan (AI), intelijen citra (IMINT), dan sistem peringatan dini (EWS) dalam kerangka intelijen strategis untuk menghadapi ancaman keamanan kontemporer. Dengan menggunakan pendekatan konseptual kualitatif berbasis systematic literature review (SLR), artikel ini mensintesis kajian-kajian terbaru mengenai AI, pengumpulan data intelijen, dan tata kelola keamanan guna menjelaskan bagaimana ketiga komponen tersebut dapat diintegrasikan secara konseptual ke dalam satu alur analitik. Temuan menunjukkan bahwa AI meningkatkan kecepatan pemrosesan citra, memperluas kapasitas deteksi perubahan, dan memperkuat akurasi identifikasi objek, sehingga memungkinkan EWS bekerja lebih cepat dan lebih prediktif dibandingkan proses analisis konvensional yang dilakukan secara manual. Namun, penerapan AI juga memunculkan risiko bias algoritmik, positif palsu (false positive), keterbatasan keterjelasan (explainability), dan potensi manipulasi citra, sehingga pengawasan oleh manusia menjadi suatu keharusan. Artikel ini mengajukan model AI-IMINT-EWS yang terdiri atas empat lapisan yang saling berkaitan, akuisisi, analisis mesin, verifikasi manusia, dan peringatan dini, yang memposisikan AI sebagai akselerator analitis, analis manusia sebagai verifikasi strategis, dan EWS sebagai simpul integratif bagi peringatan dini. Model ini menawarkan kontribusi konseptual bagi pengembangan praktik intelijen

modern yang cepat, akuntabel, dan adaptif, sekaligus membuka agenda validasi empiris di masa mendatang.

Kata Kunci : Kecerdasan Buatan; Intelijen Citra; Sistem Peringatan Dini; Intelijen Strategis; Keamanan Kontemporer.

INTRODUCTION

The development of visual data over the past two decades has fundamentally transformed the modern intelligence landscape. High-resolution satellite imagery, drone footage, optical sensors, and multi-source data are now generated in volumes that can no longer be adequately processed through manual reading alone. Under these conditions, intelligence work must not only be accurate but must also operate within extremely short time frames in order to capture threat signals before they escalate into larger crises. Consequently, the ability to interpret imagery quickly, consistently, and contextually has become one of the core competencies in contemporary security (Maslej et al., 2024). This shift is not merely a matter of technological convenience; it reflects a structural change in the nature of threats themselves, which increasingly unfold through incremental, hard-to-observe movements rather than sudden, single, easily detectable events.

The development of artificial intelligence (AI) offers considerable potential to address this challenge. AI enables the automation of object classification, change detection, pattern recognition, and cross-source data integration at a scale that far exceeds human capacity. Maslej et al. (2024) note that AI continues to advance rapidly in terms of performance, adoption, and social impact, while simultaneously raising serious challenges related to responsible AI, model transparency, and risk evaluation (Stanford Institute for Human-Centered Artificial Intelligence, 2024). In the intelligence context, this is particularly important, since the systems used to interpret threats must not only be fast but also explainable, auditable, and accountable. A system that produces rapid outputs without an adequate explanatory basis may accelerate analysis while simultaneously degrading the quality of the strategic judgment that follows from it.

Research on intelligence collection disciplines shows that the field continues to evolve and demands broader conceptual synthesis in order to respond to the complexity of modern threats (Henrico & Putter, 2025). Imagery intelligence (IMINT), as one of the principal forms of intelligence collection, occupies a strategic position because it provides visual representations of activities, changes, and patterns that are not always captured by other intelligence sources. However, the continually growing scale of IMINT data has made traditional analytical processes increasingly vulnerable to delay, analytical fatigue, and perceptual bias. This is precisely where AI offers added value not as a replacement for analysts, but as a means of augmenting human analytical capacity (Henrico & Putter, 2025). The augmentation logic is important to underline from the outset, because much of the public discourse surrounding AI in security tends to oscillate between two extremes: uncritical enthusiasm about full automation on one hand, and reflexive rejection of AI as inherently untrustworthy on the other. Neither extreme adequately captures how AI is actually positioned within a functioning intelligence cycle.

This article addresses a specific gap in that discourse. Although the literature on AI in security continues to grow, much of it remains at a technical level, focusing on computer vision, machine learning, or AI model security. Relatively few studies conceptually connect AI, IMINT, and early warning systems (EWS) within a strategic intelligence framework. Yet contemporary threats whether in the form of conflict escalation, military mobilization, infrastructure sabotage, or shifts in the balance of power, require systems capable of detecting change at an early stage and translating it into actionable strategic warning ((El Morr et al., 2024; Kimaita & Irungu, 2024).

Without a conceptual bridge linking the technical capability of AI, the collection discipline of IMINT, and the anticipatory function of EWS, organizations risk adopting AI tools in a fragmented manner, gaining processing speed without a corresponding improvement in strategic foresight.

The novelty of this article lies in its attempt to construct a conceptual AI-IMINT-EWS model that integrates three layers of work. The first layer positions AI as an analytical accelerator that processes imagery data rapidly. The second layer positions human analysts as strategic verifiers who interpret AI outputs within security, political, and geopolitical contexts. The third layer positions EWS as the integrative node that converts analytical results into early warning for decision-makers. Within this framework, the article addresses not only how AI processes imagery, but also how the resulting output can be transformed into strategic knowledge of value to the state (Maslej et al., 2024). The model is deliberately conceptual rather than technical: it does not propose a specific algorithm or software architecture, but rather a way of organizing the relationship between machine capability, human judgment, and institutional response.

The central research question is: how can AI-IMINT integration improve the effectiveness of early warning systems in supporting strategic intelligence? This question matters because an effective EWS is not merely a monitoring system but an anticipatory mechanism that enables organizations to recognize weak signals, assess threat levels, and allocate responses in a timely manner. In many domains, including public health, AI has been shown to accelerate early detection through the analysis of complex data patterns. Studies on AI-based epidemic and pandemic early warning systems indicate that AI can improve both the speed and accuracy of detection, although challenges related to data quality, bias, and model transparency remain significant (El Morr et al., 2024). The same logic is highly relevant to the security and intelligence domain, where the cost of a missed or delayed signal is measured not in infection rates but in strategic surprise, escalation, or loss of life.

Furthermore, the relevance of this topic is reinforced by growing attention to trust, risk, and security management in AI. The AI TRiSM literature emphasizes that organizations must proactively manage AI-related risk through governance, monitoring, validation, and data protection (Habbal et al., 2024). In the intelligence domain, these demands are even greater, since even minor errors in image classification or change interpretation can affect high-value strategic decisions. AI use in IMINT must therefore always be situated within a rigorous control framework, in which technological capability is matched by institutional accountability.

Beyond the technical promise of AI, it is essential to recognize that AI has evolved into a form of cognitive infrastructure across numerous sectors, including security, defense, and intelligence. Within this framework, AI is no longer understood merely as a computational tool but as a system capable of transforming the human relationship with information. The AI Index Report 2024 shows that AI development is proceeding rapidly in terms of model performance, investment, and industry adoption, while also generating an urgent need for more standardized responsible AI practices (Maslej et al., 2024). This indicates that technological gains must always be balanced against adequate governance, a point that recurs throughout this article and ultimately motivates the design of the proposed model.

In the field of information security more broadly, recent surveys show that AI supports threat detection, risk classification, and response to increasingly complex attacks. However, Hashmi et al. (2025) note that AI also introduces new vulnerabilities, such as adversarial attacks, dependence on data quality, and challenges in integrating with traditional security systems. In the intelligence context, these characteristics become more sensitive, since AI is used not only for digital defense but also for reading physical and strategic signals from imagery. A compromised model in

a purely digital security setting may cause a data breach; a compromised model in an IMINT setting may cause a state to misread another actor's intentions.

Meanwhile, Vassilev *et al.* (2025) emphasize the need for a clearer taxonomy and terminology in adversarial machine learning, including attack types such as evasion, poisoning, and privacy attacks on both predictive and generative systems. For IMINT, the implications are significant: manipulated imagery, compromised inputs, or adversarially attacked models can mislead intelligence analysis. Consequently, the security of AI models becomes inseparable from the effectiveness of intelligence itself, and any conceptual model that integrates AI into IMINT must treat adversarial risk as a first-order design consideration rather than an afterthought.

Turning specifically to imagery intelligence as a discipline, IMINT uses imagery to understand objects, activities, and changes within an area of interest. Historically, IMINT has played a crucial role in mapping military capabilities, strategic infrastructure, and activity patterns that are not directly accessible through other means. However, the primary challenge facing modern IMINT is no longer data scarcity but data surplus. Imagery generated by various sensors is now produced faster than human capacity can interpret it simultaneously (Henrico & Putter, 2025). This inversion from a discipline historically constrained by access to imagery, to one constrained by the capacity to interpret an abundance of imagery is arguably the single most important structural change affecting IMINT practice today.

At this point, AI opens new possibilities through automated detection, neural-network-based classification, and change detection. Advances in computer vision and machine learning enable faster and more consistent object identification. However, increased speed does not automatically guarantee intelligence quality. Visual data must still be linked to geographic, political, and operational context. AI-assisted IMINT should therefore be understood as a collaborative system rather than a fully automated one. Human analysts remain essential for validating results, connecting visual indicators with actor intent, and assessing implications for national security. Literature on intelligence collection disciplines indicates that intelligence collection increasingly requires cross-disciplinary integration, since modern threats are cross-domain in nature and difficult to interpret from a single information source (Henrico & Putter, 2025). IMINT can be particularly powerful when combined with other sources such as SIGINT, OSINT, and HUMINT. Nevertheless, this article focuses primarily on IMINT, since imagery is the most direct medium for detecting physical changes related to strategic threats, and AI strengthens this function by expanding the scale and speed of analysis.

Early warning, the third element of the proposed model, is a mechanism designed to recognize early signals of a threat and translate them into actionable warnings. In public health, AI-based early warning systems have demonstrated the ability to detect epidemic and pandemic patterns faster than conventional methods, although they remain constrained by data quality, bias, and limited generalizability (El Morr *et al.*, 2024). This logic is relevant to security, since the underlying principle is similar: detecting small changes before they escalate into major threats. In the strategic intelligence domain, EWS does not function merely as an alarm but as a system of interpretation. It must be able to assess risk levels, construct indicators, and connect weak signals to potential escalation. With AI assistance, EWS can process far larger volumes of data, identify anomalies that are not easily visible, and update assessments in near real time. Research on global infectious disease early warning models likewise shows that the integration of big data and AI is key to building more effective early warning models (Hu *et al.*, 2025).

However, cross-domain experience shows that AI-based EWS is vulnerable to similar problems: incomplete data, imbalanced data distribution, and limited model interpretability.

Consequently, EWS design in the security domain must not disregard human judgment. In a strategic context, warning errors can lead to overreaction or underreaction, both of which are equally detrimental. An effective EWS must therefore maintain a balance between detection sensitivity and verification reliability. It is this balance between speed and caution, automation and judgment that forms the conceptual core of the model developed in this article, and it is to the method used for developing that model that the discussion now turns.

RESEARCH METHODS

This article adopts a qualitative conceptual approach based on a systematic literature review. Its aim is not to test causal relationships empirically, but to build an analytical framework that explains how AI, IMINT, and EWS interrelate within strategic intelligence. A conceptual method of this kind is appropriate for early-stage theory development in a field where empirical data on operational AI-IMINT-EWS integration remain scarce, fragmented across institutional reports, or classified. Rather than attempting to measure the performance of a specific AI system, the article seeks to organize existing knowledge into a coherent structure that can subsequently guide empirical inquiry.

The literature used is drawn from academic journals, institutional reports, and recent scholarly publications relevant to AI, security, and early warning systems, as listed in the References. Sources were selected based on three criteria: (1) direct relevance to at least one of the three core concepts, artificial intelligence, imagery intelligence, or early warning systems; (2) recency, with priority given to publications from 2024 and 2025 in order to reflect the current state of AI capability and governance discourse; and (3) credibility, prioritizing peer-reviewed journal articles and reports issued by recognized research institutions, such as the Stanford Institute for Human-Centered Artificial Intelligence and the National Institute of Standards and Technology.

The synthesis process was conducted in four stages. First, identification of key themes such as AI, imagery intelligence, early warning systems, adversarial risk, and strategic intelligence, drawn from an initial reading of the selected literature. Second, critical reading of the role of each concept within the contemporary security context, with attention to how each source frames the relationship between technological capability and institutional practice. Third, thematic grouping to connect findings across the literature, so that recurring ideas, such as the tension between speed and reliability, or the necessity of human oversight could be identified across otherwise disconnected bodies of work spanning security studies, public health informatics, and AI governance. Fourth, conceptual synthesis to formulate the integrative AI-IMINT-EWS model, in which the thematic groupings from the third stage were organized into a layered structure representing the flow of intelligence work from data acquisition to strategic warning. This approach is appropriate for a conceptual article, as it allows for theory development without reliance on primary field data, while still grounding every element of the proposed model in existing scholarship.

The analysis addresses three derivative questions that structure the discussion in the following section: how AI transforms IMINT practice; how IMINT supports EWS; and how EWS can generate strategic warnings usable by decision-makers. Within this framework, the discussion moves beyond a mere description of technology toward institutional and governance logic. In addition, this article situates security, trust, and governance as integral to system design rather than as supplementary considerations (Habbal *et al.*, 2024; Vassilev *et al.*, 2025). This positioning reflects a deliberate methodological choice: rather than treating governance as a closing caveat appended to a technical discussion, the article treats risk and accountability as constitutive elements of the model itself, discussed alongside the analytical capabilities of AI.

As with any conceptual article grounded in literature synthesis rather than primary data collection, this method carries limitations that should be acknowledged. The conclusions drawn are analytical rather than empirical, and the proposed model has not yet been tested against operational data from a specific intelligence or defense institution. The value of the method lies instead in its capacity to organize a fragmented literature into a structure that is internally coherent and that can serve as a reference point for subsequent empirical research, a point revisited in the concluding section of this article.

RESULTS AND DISCUSSION

The Transformation of IMINT through Artificial Intelligence

The central conceptual finding of this article is that AI transforms IMINT from a highly manual practice into a collaborative practice between machines and humans. In traditional systems, analysts must examine imagery one at a time, a process that demands considerable time, energy, and concentration. As the volume of imagery grows exponentially, analysts' workload rises sharply, and the risk of error increases accordingly. AI addresses this bottleneck by automatically performing preprocessing, classification, segmentation, and change detection.

The first advantage of AI is speed. AI-based systems can scan vast volumes of visual data and flag relevant areas within a short time. The second advantage is consistency, as a well-trained model can apply uniform analytical criteria across diverse datasets, reducing the variability that arises when different human analysts apply slightly different standards to similar imagery. The third advantage is scalability, as systems can be expanded to monitor multiple areas simultaneously without a proportional increase in staffing. In the security context, these three advantages are critical, since threats do not wait for slow analytical processes (Hashmi et al., 2025).

This transformation, however, should not be understood as the replacement of humans by machines. AI performs optimally when positioned as an augmentation of analytical capacity. Human analysts remain necessary to interpret context, test the validity of findings, and assess the strategic significance of visual change. Without human involvement, AI outputs risk becoming a set of predictions devoid of operational meaning. The relationship between AI and IMINT is therefore complementary rather than substitutive, a point that recurs throughout the discussion and that ultimately shapes the structure of the model proposed in Table 1.

Artificial Intelligence and the Quality of Change Detection

In the context of strategic intelligence, change detection is a critically important function. Small changes in infrastructure, transportation activity, flight patterns, or object configuration can serve as early signals of policy shifts or operational preparations. AI enables such change detection to be performed with greater sensitivity and speed. In many cases, this type of information determines whether a state is able to take preventive action before a crisis escalates.

However, high sensitivity also introduces the risk of false positives. Systems may flag changes that are, in fact, inconsequential, or misinterpret visual variation as a threat. This is where multi-layered validation becomes necessary. Vassilev et al. (2025) note that AI models remain vulnerable to adversarial manipulation, requiring systems to be designed with attention to evasion, poisoning, and data misuse. For IMINT, this means that cross-verification against other sources becomes a fundamental requirement rather than a discretionary safeguard.

In addition, AI models used in IMINT must possess an adequate level of explainability. Decision-makers need not only to know that a system has issued a warning, but also to understand why that warning was generated. Without explanation, trust in the model becomes difficult to establish. Habbal et al. (2024) emphasize that transparency, monitoring, and governance are core

elements in building trustworthy AI. This principle is especially important in security, where strategic decisions demand a high degree of legitimacy, and where an unexplained warning is more likely to be dismissed than one accompanied by a defensible rationale.

Implications for Early Warning Systems

AI-supported early warning systems serve two principal functions. First, accelerating the detection of weak signals. Second, assisting in threat prioritization based on urgency and probability. In conventional systems, weak signals are often obscured within large volumes of data or dismissed as noise due to insufficient clarity. AI helps convert noise into interpretable patterns. At this point, EWS becomes more dynamic and responsive, shifting from a passive repository of alerts toward an active instrument of anticipation.

Studies on AI-based epidemic and pandemic early warning systems show that AI is effective in detecting early patterns, though its effectiveness depends on data quality, variable completeness, and the model's capacity for generalization (El Morr et al., 2024). This lesson is highly relevant to security-oriented EWS. If imagery data is incomplete, of low resolution, or misread in context, the resulting warning will likewise be weak. In other words, model sophistication cannot fully compensate for data limitations, a caution that should inform how much confidence institutions place in any single automated output.

In a strategic context, EWS should provide not only an alarm but also actionable recommendations. The system must be able to answer the question: does this visual change indicate a routine pattern, a temporary anomaly, or genuine threat preparation? Answering this question requires combining AI output with analyst interpretation. The proposed AI-IMINT-EWS model therefore positions EWS as an integrative mechanism rather than a mere technical notification (El Morr et al., 2024; Hu et al., 2025).

Governance, Risk, and Ethics

The greatest challenge in using AI for IMINT and EWS is not purely technical but also ethical and institutional. AI can process data rapidly, but this speed may generate an illusion of certainty. If a model misreads an object or misjudges a change, the consequences can be severe. In the security domain, such errors may affect defense policy, diplomacy, or even conflict escalation, particularly when a warning is acted upon before it has been adequately verified.

Consequently, AI governance must form a core component of intelligence systems. The NIST and AI TRiSM literature emphasizes the importance of clear terminology, risk classification, model auditing, continuous monitoring, and protection against adversarial attacks (Habbal et al., 2024; Vassilev et al., 2025). Furthermore, the AI Index Report 2024 indicates that responsible AI practices remain highly varied and insufficiently standardized (Maslej et al., 2024). This implies that user organizations must establish internal procedures stricter than simply relying on model capability, treating governance as an operational requirement rather than an aspirational goal.

Ethics cannot be separated from this discussion. In the intelligence context, AI use raises questions of privacy, proportionality, and accountability. Human-in-the-loop is not merely a technical solution but also an ethical mechanism that ensures strategic decisions remain under human control. Technology thus does not stand as a final authority, but as a tool that supports more prudent judgment, and this ethical framing underlies the placement of human verification as a distinct layer in the model presented below.

The Conceptual AI-IMINT-EWS Model

Based on the synthesis above, the AI-IMINT-EWS model can be described as a four-layer flow, as summarized in Table 1.

Table 1. The four layers of the conceptual AI-IMINT-EWS model

Layer	Main Function	Actors/Components
1. Acquisition	Generating imagery from various sensing platforms	Satellites, UAVs, other optical sensors
2. Machine Analysis	Preprocessing, classification, object detection, change detection, anomaly flagging	AI/computer vision models
3. Human Verification	Reviewing AI output, assessing context, testing strategic relevance	Intelligence analysts
4. Early Warning	Compiling analytical results into indicators and delivering them to decision-makers	EWS, policymakers

Source: Compiled by the author, 2026 (synthesized from the literature on Artificial Intelligence, Imagery Intelligence, and Early Warning Systems).

The model affirms that the value of AI lies in acceleration rather than absolutization. AI accelerates the process, but humans retain the decisive role in interpretation and final decision-making. In practice, this model enables intelligence institutions to handle large data volumes, enhance vigilance toward strategic change, and reduce response delays. The model is also compatible with the principles of trust, risk, and security, as each layer performs a distinct control function (Habbal et al., 2024): the acquisition layer controls data quality at the point of collection, the machine analysis layer controls consistency and scale, the human verification layer controls contextual accuracy, and the early warning layer controls the translation of analysis into decision-relevant communication.

Artificial Intelligence as a Force Multiplier for Intelligence Capacity

In developing the AI-IMINT-EWS model, the most important point is not merely that AI can process imagery faster, but that AI transforms the structure of intelligence work itself. In traditional systems, the analytical process depends heavily on individual expertise and available time. As imagery volume increases, organizational capacity often lags behind, placing important signals at risk of being missed. AI addresses this problem by acting as a force multiplier, an analytical capacity multiplier that allows a small team to handle data at scale (Maslej et al., 2024).

This conceptual maturity matters because AI should not be viewed as a stand-alone automated solution. Habbal et al. (2024) show that organizations need to manage AI through a comprehensive trust, risk, and security framework, including data governance, model oversight, and operational risk mitigation. In strategic intelligence, this principle becomes even more important, since even minor errors can affect threat perception, national priorities, and policy direction.

Accordingly, AI's contribution to IMINT is best understood as an increase in analytical throughput rather than a substitute for strategic reasoning. AI helps filter noise, prioritize areas of attention, and flag changes requiring further verification. In operational terms, AI frees analysts from repetitive tasks, allowing them to focus on high-value work such as interpretation, synthesis,

and implication assessment, work that remains by its nature, a distinctly human contribution to the intelligence cycle.

Change Detection as the Core Value-Added Function

Within IMINT, one of the most important functions is change detection. Many strategic threats do not manifest as immediately visible major events, but rather as accumulations of small changes: new structures, vehicle movement, increased logistical activity, or modified spatial usage patterns. AI is highly effective in detecting such patterns, as machine learning models can consistently compare large numbers of images across time (Hashmi et al., 2025).

Findings from research on AI-based epidemic and pandemic early warning systems show that AI performs strongly in recognizing early patterns, though the quality of its outputs depends heavily on data quality, variable completeness, and bias handling (El Morr et al., 2024). This analogy is highly relevant to IMINT: if input imagery is excessively variable, of low resolution, or lacking geographic context, AI may produce false positives or false negatives. Change detection must therefore be accompanied by contextual validation rather than treated as a self-sufficient analytical output.

In the security context, change detection is not merely a matter of visual detection but also of meaning. A new structure in a given area may not be significant when read in isolation, but becomes important when connected to military training patterns, logistical flows, or regional political dynamics. AI output should therefore be viewed as an initial signal rather than a final conclusion. Human analysts remain responsible for transforming visual signals into fuller strategic knowledge.

The Need for Multi-Source Verification

One of the greatest weaknesses of AI systems in the intelligence domain is their tendency to generate false certainty. A system may express high confidence in a particular classification result, even though that result may not be substantively accurate. This is why multi-source verification is an essential component of AI-IMINT-EWS (Hashmi et al., 2025). Visual results should ideally be contextualized against other sources, such as field reports, open-source data, SIGINT, or relevant geopolitical information.

Vassilev et al. (2025) emphasize that adversarial machine learning must be understood through a clear threat taxonomy, including attacks targeting the training, inference, and deployment stages of a model. For IMINT, this means systems must be prepared to face image manipulation, data interference, or deliberately falsified inputs. Verification must therefore occur not only at the technical level but also at the methodological level. A robust system requires dual confirmation pathways so that AI output is not immediately treated as established fact.

A multi-source approach also strengthens resilience against misinterpretation. If one AI model flags an area as a threat but other sources do not support this assessment, analysis can be delayed or deepened. Conversely, if multiple sources converge on the same pattern, confidence in the warning increases. A mature AI-IMINT-EWS system therefore does not rest on a single model but on a mutually reinforcing ecosystem of evidence.

Connections to Cross-Domain Early Warning Systems

One reason this model is conceptually robust is that early warning systems have proven relevant across many other domains. El Morr et al. (2024) show that AI can help detect threats earlier through data integration and pattern analysis. Hu et al. (2025) study on global infectious disease early warning models likewise confirms that big data and AI have become key approaches

for building more effective early warning models. These findings lend legitimacy to the claim that AI-based EWS logic is not speculative but an increasingly established approach across multiple fields of application.

Nevertheless, transferring this logic from public health to security must be undertaken carefully. In health, threat indicators relate to disease cases, transmission trends, and epidemiological data. In security, indicators may include shifts in military posture, infrastructure construction, asset mobilization, or unusual logistical activity. Although the objects differ, the underlying logical structure of EWS remains the same: recognizing anomalies, measuring urgency, and triggering timely response. For this reason, cross-domain literature is highly valuable in enriching the conceptual model developed in this article.

AI can also serve as a bridge between big data and rapid decision-making. Research on AI-based epidemic intelligence shows that AI-based systems are increasingly positioned as complements to traditional surveillance rather than replacements for it (El Morr et al., 2024; Hu et al., 2025). This principle aligns with the model proposed in this article. Security-oriented EWS should likewise be understood as a complement to existing intelligence mechanisms, rather than a wholesale replacement for traditional analytical processes. Its strength lies precisely in its ability to accelerate interpretation and expand the scope of surveillance, not in its capacity to operate independently of established institutional practice.

Implications for Strategic Intelligence

In the strategic intelligence framework, time is a decisive factor. Threats detected too late complicate policy response. AI-based EWS therefore makes a direct contribution to high-level decision-making. Systems capable of detecting change at an early stage provide policymakers with the space to weigh options, allocate resources, and reduce the likelihood of strategic surprise (Hashmi et al., 2025).

At the same time, strategic intelligence requires a high degree of trust from decision-makers. If a system frequently produces false alarms, warnings will be ignored. If a system is too slow, its benefits are lost. AI-IMINT-EWS design must therefore balance sensitivity and specificity. An overly sensitive model floods users with signals, while an overly conservative model risks missing significant threats. This balance lies at the core of an effective EWS, and it cannot be resolved through technical tuning alone; it requires ongoing institutional judgment about acceptable error rates in a given strategic environment.

It is equally important to position human-in-the-loop as the primary control structure. In practical terms, analysts do not merely review AI output but also determine whether a warning warrants escalation to the policy level. This human role is not a weakness but a strength of the system, since strategic decisions require judgment that cannot be fully quantified. AI can calculate probability, but it is humans who connect probability to political and ethical consequence.

Bias Risk and Governance Improvement

Algorithmic bias is a central issue in modern AI. If training data is unrepresentative, a model will tend to reproduce inaccurate or discriminatory patterns. In IMINT, bias may emerge when a model is overly sensitive to particular object types, regions, or lighting conditions, resulting in unstable classification and erroneous inference. Model testing must therefore be conducted across a range of operational scenarios (Vassilev, 2025).

Habbal et al. (2024) emphasize the importance of proactive, rather than reactive, risk management. For intelligence systems, this means model audits must be conducted from the design stage onward, not only after incidents occur. Organizations need clear documentation of data

sources, model parameters, usage limitations, and update procedures. Without such governance, AI itself may become a new source of vulnerability rather than a solution to existing ones.

Furthermore, a well-designed model must support continuous learning. The security landscape changes rapidly, meaning a model trained today may quickly become outdated. Systems must therefore incorporate mechanisms for retraining, re-evaluation, and periodic database updates. AI-IMINT-EWS is thus not a static product but an adaptive system that evolves alongside the dynamics of threat, requiring sustained institutional investment rather than a one-time deployment.

Theoretical Contribution of the Article

From a theoretical standpoint, this article offers three main contributions. First, it broadens the understanding of IMINT by situating it within the AI and EWS ecosystem rather than treating it as a stand-alone discipline. Second, it demonstrates that early warning is not solely a concept within public health or disaster management, but a principle highly relevant to strategic intelligence. Third, it asserts that AI governance is an inherent component of security analysis rather than a supplementary technical issue (El Morr et al., 2024; Habbal et al., 2024).

This theoretical contribution matters because security literature often separates technology, ethics, and strategy into distinct domains. In practice, however, these three dimensions are closely interconnected. AI determines the speed of analysis, governance determines trust, and strategy determines ultimate meaning. By integrating these three dimensions, the AI-IMINT-EWS model offers a more comprehensive conceptual framework for reading contemporary threats, one that treats speed, accountability, and strategic relevance as jointly necessary rather than as competing priorities to be traded off against one another.

CONCLUSION

This article demonstrates that the integration of AI, IMINT, and EWS offers significant opportunities to strengthen strategic intelligence amid contemporary threats. AI accelerates imagery interpretation, IMINT provides strategic visual signals, and EWS converts those signals into actionable early warnings. The proposed AI-IMINT-EWS model comprising an acquisition layer, a machine analysis layer, a human verification layer, and an early warning layer, positions AI as an analytical accelerator, human analysts as strategic verifiers, and EWS as the integrative node for early warning.

The effectiveness of this model depends on three key factors: proactive governance of algorithmic and adversarial risk, consistent human verification of AI output, and layered protection against data and imagery manipulation. Without these three elements, the speed offered by AI risks producing false certainty capable of misleading strategic decision-making. These three factors are not independent safeguards to be pursued in isolation; rather, they reinforce one another, since governance without human verification remains largely procedural, and human verification without adequate governance lacks the institutional infrastructure needed to be applied consistently.

Theoretically, this article contributes by situating IMINT within a broader AI and EWS ecosystem, extending the relevance of early warning logic from the public health domain to the security domain, and affirming AI governance as an inherent component of intelligence analysis. Practically, the model offers intelligence and defense institutions a structured way to think about where AI adds value, where human judgment remains irreplaceable, and where institutional safeguards must be strengthened before AI-IMINT-EWS integration is scaled up operationally.

Further research is needed to empirically test this model, for instance through case studies of AI-IMINT-EWS implementation within specific intelligence or defense institutions, in order to validate its operational feasibility beyond the conceptual level. Future studies might also examine how the four-layer model performs under varying data quality conditions, how different institutional cultures shape the balance between automation and human verification, and how the governance mechanisms proposed here can be operationalized into concrete standard operating procedures. In doing so, subsequent research can move the AI-IMINT-EWS model from a conceptual proposition toward a tested framework capable of guiding practice in an increasingly complex and fast-moving security environment.

REFERENCES

- El Morr, C., Ozdemir, D., Asdaah, Y., Saab, A., El-Lahib, Y., & Sokhn, E. S. (2024). AI-based epidemic and pandemic early warning systems: A systematic scoping review. In *Health Informatics Journal* (Vol. 30, Number 3). SAGE Publications Ltd.
<https://doi.org/10.1177/14604582241275844>
- Habbal, A., Ali, M. K., & Abuzaraida, M. A. (2024). Artificial Intelligence Trust, Risk and Security Management (AI TRiSM): Frameworks, applications, challenges and future research directions. In *Expert Systems with Applications* (Vol. 240). Elsevier Ltd.
<https://doi.org/10.1016/j.eswa.2023.122442>
- Hashmi, E., Yamin, M. M., & Yayilgan, S. Y. (2025). Securing tomorrow: a comprehensive survey on the synergy of Artificial Intelligence and information security. *AI and Ethics*, 5(3), 1911–1929. <https://doi.org/10.1007/s43681-024-00529-z>
- Henrico, S., & Putter, D. (2025). Intelligence Collection Disciplines—A Systematic Review. *Journal of Applied Security Research*, 20(1), 46–70.
<https://doi.org/10.1080/19361610.2023.2296765>
- Hu, W. H., Sun, H. M., Wei, Y. Y., & Hao, Y. T. (2025). Global infectious disease early warning models: An updated review and lessons from the COVID-19 pandemic. In *Infectious Disease Modelling* (Vol. 10, Number 2, pp. 410–422). KeAi Communications Co.
<https://doi.org/10.1016/j.idm.2024.12.001>
- Kimaita, S., & Irungu, E. J. (n.d.). *New Frontiers in Conflict Prevention: Integrating Artificial Intelligence in Early Warning and Response Systems in Kenya*.
<https://doi.org/10.47772/IJRISS>
- Maslej, N., Fattorini, L., Perrault, R., Parli, V., Reuel, A., Brynjolfsson, E., Etchemendy, J., Ligett, K., Lyons, T., Manyika, J., Niebles, J. C., Shoham, Y., Wald, R., & Clark, J. (2024). The AI Index 2024 annual report. Stanford Institute for Human-Centered Artificial Intelligence..
- Vassilev, A. (2025). *Adversarial Machine Learning*: <https://doi.org/10.6028/NIST.AI.100-2e2025>