



CONCEPTUALIZING THE APPLICATION OF ARTIFICIAL INTELLIGENCE (AI) IN THE MECHANISM FOR GENERATING VESSEL AIS DATA ANOMALY ANALYSIS REPORTS

Ade Tya Guritna¹, Pebrianto Eko Nugroho², and Sapto Budiarto³

¹ Politeknik Angkatan Laut, Indonesia

² Politeknik Angkatan Laut, Indonesia

³ Politeknik Angkatan Laut, Indonesia

Corresponding Author:

Ade Tya Guritna,
Department of Naval Operational Strategy, Politeknik Angkatan Laut.
Ciledug Raya Street No.2, Seskoal, South Jakarta, DKI Jakarta, Indonesia
Email: flightengineer56@gmail.com

Article Info

Received: February 7, 2026

Revised: May 18, 2026

Accepted: July 14, 2026

OnlineVersion: August 31,
2026

Abstract

Automatic Identification System data generate continuous spatiotemporal information that supports maritime surveillance, yet converting detected vessel anomalies into reliable analytical reports remains challenging because statistical abnormality does not always indicate operational significance. This study aimed to conceptualize and evaluate an Artificial Intelligence-supported mechanism for generating evidence-grounded vessel AIS anomaly analysis reports. A Design Science Research approach was employed through AIS preprocessing, trajectory reconstruction, anomaly detection, contextual interpretation, explainability, evidence packaging, controlled natural-language generation, and expert validation. Findings showed that the anomaly-detection component achieved satisfactory predictive performance, while report quality improved substantially after provenance, contextual validation, and uncertainty controls were incorporated. Expert assessments indicated stronger factual consistency, traceability, contextual relevance, and uncertainty communication in the refined reports, while case analyses demonstrated that AI could distinguish observed AIS facts from model inference more effectively when constrained by structured evidence. The study concludes that generative AI should function as an evidence-constrained analytical communicator rather than an autonomous interpreter of vessel intent or maritime risk. The proposed architecture contributes an integrated framework for transparent, auditable, and human-verified AIS anomaly reporting and provides a foundation for future validation across diverse maritime environments, vessel classes, and operational contexts under realistic human-in-the-loop maritime decision-support conditions across waters.

Keywords: Artificial Intelligence; Automatic Identification System; Design Science Research; Maritime Anomaly Detection; Report Generation



© 2026 by the author(s)

This article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution-ShareAlike 4.0 International (CC BY SA) license (<https://creativecommons.org/licenses/by-sa/4.0/>).

Journal Homepage

<https://research.adra.ac.id/index.php/jzca>

How to cite:

Guritna, T. A., Nugroho, E. P., & Budiarto, S. (2026). Conceptualizing the Application of Artificial Intelligence (AI) In the Mechanism for Generating Vessel Ais Data Anomaly Analysis Reports. *Journal of Computer Science Advancements*, 4(4), 286–307. <https://doi.org/10.70177/jzca.v4i4.4375>

Published by:

Yayasan Adra Karima Hubbi

INTRODUCTION

Maritime transportation increasingly depends on continuous digital information to support navigational safety, vessel monitoring, traffic management, and maritime situational awareness (Khudhur & Al-Alawi, 2024). The Automatic Identification System (AIS) constitutes a major component of this information environment because equipped vessels automatically exchange identification, positional, navigational, and voyage-related information through maritime radio communication (Purohit et al., 2025). International requirements under the Safety of Life at Sea framework make AIS an important source of operational maritime data, while the current ITU-R M.1371-6 recommendation specifies the technical characteristics of AIS operating through time-division multiple-access techniques in the VHF maritime mobile band (Voicu et al., 2025). The scale, frequency, and continuity of AIS transmissions create substantial analytical opportunities, but they also generate data volumes that increasingly exceed the practical capacity of purely manual examination.

Vessel AIS data provide a temporal and spatial representation of maritime movement through attributes such as position, speed, course, heading, navigational status, vessel identity, and transmission time (Babaei et al., 2024). Abnormal patterns within these data may include unexpected trajectory deviations, unusual speed or course changes, anomalous stops, inconsistent positional behavior, abnormal transmission gaps, and context-dependent deviations from expected vessel activity (Loganathan et al., 2025). Recent studies demonstrate that machine learning, deep learning, contextual autoencoders, and spatio-temporal graph models can be used to recognize increasingly complex patterns of anomalous vessel behavior in AIS streams (Kaur & Taak, 2026). Such developments shift maritime surveillance from retrospective inspection toward data-driven systems capable of identifying potentially significant deviations within large and continuously changing traffic environments.

Artificial Intelligence creates a further possibility beyond anomaly detection by supporting the interpretation and communication of detected events (Salam et al., 2026). Emerging architectures have begun integrating AIS time-series analysis with language models to produce contextual explanations and situation summaries alongside trajectory prediction, anomaly detection, and collision-risk assessment. AIS-LLM, for example, demonstrates the feasibility of connecting time-series vessel information with large language models to generate human-readable explanations and briefings from analytical outputs (Nazirov et al., 2026). This development suggests that AI can potentially transform the anomaly-analysis workflow from a sequence of isolated algorithmic alerts into a mechanism capable of generating structured analytical reports that explain what anomaly occurred, where and when it occurred, why it was classified as abnormal, and what contextual information should be reviewed by a human analyst.

Existing AIS anomaly-detection research has achieved substantial progress in identifying abnormal vessel trajectories and behavioral patterns, yet the operational value of a detection system depends on more than its capacity to assign an anomaly score or label (Tzanova, 2024). Maritime analysts must interpret the relationship among vessel identity, historical movement, geographic context, course and speed variations, temporal patterns, and potentially relevant neighboring traffic before determining the significance of an event (Yu et al., 2024). An isolated algorithmic notification may therefore create an additional analytical burden when operators must manually reconstruct the reasoning behind the alert (Sonawane, 2025). Research on self-explainable maritime anomaly detection has explicitly identified interpretability as an important issue because operators need to understand algorithmic outputs before they can effectively apply them to maritime decision support.

AIS data themselves introduce additional analytical complexity because an unusual observation does not automatically constitute meaningful abnormal behavior (Shah et al., 2025). Context influences interpretation: a speed reduction may represent normal port approach behavior, a course alteration may reflect navigational constraints, and a trajectory deviation may have different significance depending on vessel type, location, traffic conditions, or previous

movement (Tu & Chen, 2025). Recent context-aware anomaly-detection research shows that incorporating vessel-specific context can materially affect anomaly identification, reinforcing the inadequacy of interpreting all deviations through uniform thresholds (Bai & Yao, 2026). A reliable reporting mechanism must therefore distinguish between statistical abnormality and operationally meaningful anomaly rather than automatically converting every outlier into a definitive maritime-risk statement.

The central technological problem lies in the fragmentation between data processing, anomaly identification, analytical interpretation, and report generation (Sharma et al., 2025). Conventional anomaly-detection mechanisms primarily concentrate on determining whether vessel behavior differs from learned or predefined patterns, while reporting remains a separate human-centered activity (Rigotti & Fosch-Villaronga, 2024). Generative AI may bridge these processes, but uncontrolled language generation introduces risks of unsupported explanation, overstatement of anomaly significance, loss of traceability to original AIS evidence, and false confidence in machine-generated conclusions (Atias & Mawasi, 2025). The required system should therefore not be conceptualized as an autonomous AI “decision maker,” but as a traceable analytical mechanism in which anomaly detection, contextual interpretation, evidence retrieval, report generation, confidence representation, and human verification operate as connected stages.

The primary objective of the study is to conceptualize an AI-supported mechanism for generating vessel AIS data anomaly analysis reports (Johnson, 2026). The proposed framework is intended to transform raw or preprocessed AIS data into structured analytical information through sequential processes of data validation, vessel-track reconstruction, anomaly identification, contextual interpretation, evidence extraction, and natural-language report generation (Chin, 2026). The research does not merely seek to determine which AI algorithm can achieve the highest anomaly-detection accuracy (Bartelheimer et al., 2025). Its principal concern is the architecture through which multiple analytical components can be coordinated to produce reports that remain understandable, evidence-based, traceable, and operationally meaningful.

The analytical objective is to identify the functional relationships among AIS data characteristics, anomaly-detection models, contextual reasoning mechanisms, explainable AI, and generative language technologies (Speith et al., 2024). Machine-learning and deep-learning components can identify statistical or behavioral deviations, while contextual modules can evaluate whether those deviations are meaningful relative to vessel characteristics and navigational circumstances (Liu & Ye, 2026). Explainability mechanisms can expose the variables or trajectory segments responsible for an anomaly judgment, while generative AI can convert structured evidence into concise analytical narratives (Che et al., 2025). Recent AIS-language-model research provides evidence that numerical maritime tasks and natural-language explanations can be combined within a unified architecture, creating a technical foundation for extending anomaly detection toward systematic reporting.

The practical objective is to formulate the architecture of an anomaly-report generation mechanism suitable for subsequent prototype development and empirical validation (Ahmadzadeh et al., 2026). A proposed report could contain vessel identifiers, observation periods, anomaly type, location, relevant trajectory characteristics, deviations from expected behavior, supporting AIS parameters, anomaly confidence, contextual considerations, and an explicit indication of uncertainties requiring human review (Meyer & Mariosi, 2025). The resulting conceptual architecture should also specify where deterministic rules, statistical models, machine learning, explainable AI, and generative AI are most appropriate (Afroogh et al., 2024). Such differentiation is necessary because not every stage of the reporting mechanism benefits from generative methods, particularly when numerical calculations, vessel identification, or compliance-sensitive information requires deterministic traceability.

Current research provides an increasingly sophisticated foundation for AIS-based vessel anomaly detection (Li et al., 2025). Earlier deep-learning architectures demonstrated multi-task processing of AIS streams for trajectory reconstruction and anomaly detection, while subsequent research has introduced unsupervised trajectory-image approaches, online anomaly detection, behavior classification, spatio-temporal graph modeling, and context-aware autoencoders (Vu & Nguyen, 2025). These studies significantly advance the ability to identify anomalous maritime behavior from complex trajectory data (Al-Refai et al., 2026). The dominant research orientation nevertheless remains focused on detection performance, behavioral classification, trajectory modeling, or risk identification rather than on the complete downstream mechanism through which detected anomalies are transformed into standardized analytical reports for human use.

Emerging research has begun addressing interpretability and language-based maritime reasoning, narrowing but not eliminating this gap (Michael et al., 2025). Self-explainable anomaly-detection approaches seek to make abnormality assessments more transparent, while AIS-LLM demonstrates that large language models can integrate time-series AIS information and generate contextual natural-language explanations and briefings (Wang et al., 2025). These developments show that the boundary between numerical anomaly detection and linguistic interpretation is becoming increasingly permeable (Luong et al., 2026). The remaining research space concerns the explicit design of a report-generation mechanism in which anomaly evidence, contextual assessment, explanation, uncertainty, and standardized narrative output are systematically linked rather than treated as secondary by-products of an AI model.

A deeper gap concerns governance and analytical reliability within the AI-generated reporting process (Wang et al., 2026). Detection accuracy alone cannot establish report reliability because a linguistically convincing report may contain conclusions that exceed the underlying evidence (Yue et al., 2026). Maritime anomaly reporting therefore requires provenance mechanisms linking each narrative claim to AIS records, calculated features, model outputs, or defined contextual rules (Sienknecht & Vetterlein, 2024). Human analysts must also remain able to distinguish observed facts from algorithmic inference and operational interpretation (Lin et al., 2026). The unresolved conceptual problem is consequently how to combine the scalability of AI with the evidentiary discipline required for maritime analysis, producing reports that accelerate human interpretation without obscuring uncertainty or replacing accountable professional judgment.

The principal novelty of the proposed study lies in shifting the unit of innovation from AI-based anomaly detection to an AI-supported anomaly-analysis report generation mechanism (De Larmelina et al., 2026). This distinction broadens the system boundary from identifying abnormal vessel behavior to explaining and documenting it through a structured analytical workflow (Roy et al., 2025). The proposed mechanism can be conceptualized as a layered architecture consisting of AIS ingestion and quality control, trajectory reconstruction, feature engineering, anomaly detection, contextual validation, explainability, evidence packaging, generative report synthesis, and human verification (Filep et al., 2024). Such an architecture addresses a practical gap between sophisticated computational models and the information products ultimately required by maritime analysts and decision makers.

The conceptual contribution lies in combining analytical AI and generative AI while assigning different responsibilities to each (Mansfield et al., 2025). Numerical and trajectory-processing models should determine measurable patterns and anomaly indicators, contextual reasoning should examine whether an observation is operationally unusual, explainability mechanisms should identify the evidence underlying the model output, and generative AI should convert verified structured information into readable analysis rather than invent new evidence (Fasouli et al., 2025). This separation creates an evidence-grounded reporting architecture that can reduce the risk of treating fluent language as equivalent to factual reliability (Kalli et al., 2026). The framework consequently emphasizes human-in-the-loop validation, provenance,

confidence representation, and separation between factual observations, machine inference, and analyst conclusions as fundamental design requirements.

The practical justification derives from the continuing growth of AIS-based maritime monitoring and the corresponding need to convert large volumes of vessel movement information into actionable analytical products (Kalbande et al., 2025). ITU and IMO standards establish AIS as an important component of the contemporary maritime information environment, while recent research demonstrates rapid methodological advances in AI-based trajectory analysis, contextual anomaly detection, and language-model integration (Rizk et al., 2024). The proposed research can therefore contribute to maritime informatics by connecting anomaly detection with interpretable reporting, contribute to artificial intelligence by examining evidence-grounded generation from spatio-temporal data, and contribute to maritime operational practice by conceptualizing a mechanism through which machine-detected anomalies can be transformed into transparent, standardized, and reviewable analytical reports.

RESEARCH METHOD

Research Design

This study employed a Design Science Research approach to conceptualize, develop, and evaluate an Artificial Intelligence-supported mechanism for generating vessel Automatic Identification System anomaly analysis reports. The design was selected because the principal research output was an integrated analytical artifact rather than a single anomaly-detection algorithm (Rizk et al., 2024). The proposed mechanism combined AIS data preprocessing, trajectory reconstruction, anomaly identification, contextual interpretation, explainable artificial intelligence, evidence packaging, natural-language report generation, and human verification. The research design treated these components as interdependent layers of a reporting system whose primary purpose was to transform complex vessel-movement data into structured, traceable, and operationally interpretable analytical information.

The research artifact was developed according to a layered architecture consisting of data, analytical, interpretive, generative, and validation components (Chen et al., 2026). The data layer managed AIS acquisition, cleaning, synchronization, trajectory segmentation, and data-quality assessment. The analytical layer extracted vessel-movement characteristics and identified anomalous patterns using statistical or machine-learning techniques. The interpretive layer linked detected anomalies with vessel characteristics, spatial context, temporal patterns, and relevant behavioral indicators (Qiang et al., 2026). The generative layer converted verified structured evidence into standardized natural-language reports, while the validation layer examined whether each generated statement could be traced to AIS observations, derived variables, model outputs, or explicitly defined contextual rules.

The evaluation strategy distinguished anomaly-detection performance from report-generation quality. Detection models were assessed according to their ability to discriminate normal and anomalous vessel behavior, whereas generated reports were evaluated for factual consistency, evidentiary traceability, completeness, clarity, contextual relevance, uncertainty communication, and operational usefulness. This separation was necessary because a high-performing anomaly detector does not automatically produce a reliable analytical report. The study therefore conceptualized AI report generation as an evidence-grounded decision-support process in which generative AI communicated verified analytical results rather than independently determining whether vessel behavior constituted a maritime threat.

Research Target/Subject

The quantitative subjects of this study comprise historical vessel movement trajectory datasets derived from raw Automatic Identification System (AIS) messages within designated maritime operational areas. To capture diverse navigational conditions, data sampling relies on

a non-probability purposive sampling technique, targeting high-density maritime traffic zones, complex fairways, and varied vessel categories. These sampled datasets encompass preprocessed AIS records, reconstructed trajectories, derived spatio-temporal kinematic features, and generated structured evidence packages. This selection strategy ensures a comprehensive baseline representing both ordinary maritime operations and diverse anomalous vessel behaviors.

The qualitative subjects consist of the natural-language anomaly analysis reports generated by the system, alongside expert evaluation feedback. An expert sampling technique was employed to recruit a specialized evaluation panel composed of domain experts, maritime security analysts, and senior AIS operational specialists. These human reviewers assess the quality, operational utility, factual consistency, and contextual plausibility of the AI-generated narratives. Combining quantitative trajectory data with expert qualitative review enables a rigorous dual evaluation of both algorithmic detection performance and report reliability.

Research Procedure

The first stage involved acquisition and preparation of AIS data. Raw messages were imported into a structured analytical environment and screened for missing values, duplication, inconsistent timestamps, invalid coordinates, implausible vessel-motion values, and transmission gaps. Messages belonging to the same vessel were chronologically ordered and segmented into analytically meaningful trajectories according to predefined temporal and spatial criteria. Data-quality flags were retained rather than silently removed when uncertainty itself could influence interpretation. The resulting trajectory dataset served as the evidence base for subsequent anomaly analysis.

The second stage involved feature extraction and establishment of behavioral reference patterns. Relevant kinematic and spatio-temporal indicators were calculated from consecutive AIS observations, including changes in speed, direction, distance, time interval, and trajectory geometry. Behavioral baselines were generated according to the selected analytical model and, where appropriate, differentiated by vessel characteristics or navigational context. Candidate anomalies were then identified through deviations from expected patterns. Anomaly detection was intentionally separated from threat classification because unusual movement alone does not establish malicious, unsafe, or unlawful behavior.

The third stage involved contextualization and explainability. Each candidate anomaly was examined against the variables and trajectory characteristics responsible for its classification. Contextual information was used to determine whether an apparently unusual pattern could reasonably reflect ordinary operational behavior, data-quality problems, geographic constraints, or other explainable circumstances. The system created an evidence package containing observed facts, calculated features, model results, anomaly confidence, relevant contextual information, and unresolved uncertainty. This package became the sole informational foundation for automated report generation.

The fourth stage involved AI-assisted generation of the anomaly analysis report. Structured evidence was converted into a controlled prompt or machine-readable representation supplied to the generative model. The model generated a report according to the predefined analytical template and was explicitly constrained from adding vessel events, motives, regulatory violations, or risk conclusions unsupported by the supplied evidence. Statements were differentiated according to evidentiary status, allowing the report to distinguish observed AIS information from algorithmic inference and interpretive hypotheses. Reports containing insufficient evidence were required to express uncertainty rather than produce definitive explanations.

The fifth stage involved automated and human validation of generated reports. Automated checks compared vessel identifiers, numerical values, coordinates, times, anomaly categories, and relevant parameters in the generated narrative with the structured evidence package. Human experts subsequently evaluated whether the explanation was factually consistent, contextually

plausible, sufficiently transparent, and operationally useful. Discrepancies were classified according to error type, including factual inconsistency, unsupported inference, omission, ambiguity, inappropriate certainty, or incorrect contextual interpretation. Validation findings were used to revise model prompts, report templates, evidence-selection rules, and analytical thresholds.

The final stage involved iterative refinement of the conceptual mechanism and synthesis of the resulting design principles. Performance of the anomaly-detection component and quality of generated reports were evaluated separately before being examined as an integrated workflow. Comparisons across anomaly cases were used to determine where automated generation added analytical value and where human interpretation remained essential. The finalized mechanism was formulated around five principles: evidence-grounded generation, explicit provenance, separation of observation from inference, calibrated uncertainty, and mandatory human oversight for consequential interpretation. The resulting architecture positioned AI as an analytical and reporting support mechanism rather than an autonomous authority for determining the operational significance or intent of anomalous vessel behavior.

Instruments, and Data Collection Techniques

The primary research instrument was an integrated AIS analytical pipeline developed to process vessel data from ingestion to report generation. The preprocessing component performed data validation, duplicate handling, temporal ordering, trajectory segmentation, coordinate checking, and identification of missing or implausible observations. The feature-engineering component derived behavioral indicators such as speed variation, course change, heading deviation, movement continuity, displacement, temporal gaps, trajectory deviation, and other contextually relevant characteristics. Recent maritime anomaly-detection research illustrates the importance of interpreting anomalies through combinations of spatial, temporal, and behavioral characteristics rather than relying on a single raw AIS variable.

The anomaly-detection instrument consisted of a configurable analytical module capable of applying rule-based, statistical, machine-learning, or deep-learning approaches according to data characteristics and label availability. Supervised models could be evaluated using labeled anomalies, while unsupervised approaches could identify observations or trajectories substantially different from learned normal behavior. Each detected event was required to produce a structured anomaly record containing the vessel identity, observation period, anomaly category, relevant variables, anomaly score or confidence indicator, supporting trajectory segment, and contextual evidence. This structured intermediate representation prevented the language-generation component from receiving unrestricted raw data without an explicit analytical basis.

The explainability and provenance instrument was designed to establish a transparent connection between anomaly judgments and their supporting evidence. Each analytical output documented which variables, trajectory characteristics, comparison baselines, or model features contributed to the anomaly classification. The mechanism separated directly observed facts from calculated indicators and model-based inferences. This distinction was incorporated into the report-generation prompt or structured input so that the generative model could not legitimately present an inferred interpretation as though it were a directly observed AIS fact. Evidence identifiers were retained to allow analysts to trace generated statements back to their underlying records.

The report-generation instrument consisted of a standardized template combined with a constrained generative-AI component. The template required sections describing vessel identification, time and location of the event, detected anomaly type, relevant behavioral evidence, contextual interpretation, anomaly confidence, uncertainty, and recommended points for analyst review. The language model was instructed to generate text only from the structured evidence supplied by the preceding analytical stages. Recent AIS-LLM research demonstrates

that AIS time-series outputs can be combined with language models to provide textual explanations and maritime situation summaries, supporting the technical plausibility of this architecture.

The evaluation instrument consisted of two complementary rubrics. The technical rubric assessed anomaly-detection performance using appropriate measures such as precision, recall, F1-score, ROC-AUC, false-positive rate, or unsupervised quality indicators depending on the availability of validated anomaly labels. The report-quality rubric assessed factual accuracy, evidence consistency, traceability, completeness, readability, contextual appropriateness, uncertainty representation, and usefulness for human analysis. Expert reviewers were also asked to identify unsupported claims, ambiguous language, excessive certainty, omitted evidence, or interpretations that exceeded what could reasonably be inferred from the AIS data.

Data Analysis Technique

Data analysis employs a dual-track framework to evaluate anomaly detection accuracy and report generation quality independently. Quantitative analysis of the anomaly detection performance applies standard machine learning and statistical evaluation metrics based on ground-truth label availability. Supervised detection models are evaluated using precision, recall, F1-score, ROC-AUC, and false-positive rates, whereas unsupervised approaches rely on trajectory deviation measures and cluster dispersion indicators. These quantitative techniques verify whether identified events represent statistically significant deviations from established behavioral baselines.

Qualitative and structural analysis of the generated text combines automated evidence cross-checking with expert thematic coding. Automated verification measures factual consistency, evidentiary traceability, and information completeness by matching narrative statements directly against underlying structured evidence packages. Qualitative expert evaluations are categorized using a structured error taxonomy to identify discrepancies such as unsupported inferences, ambiguous phrasing, excessive certainty, or omitted contextual variables. This analytical procedure evaluates how effectively the system converts raw analytical outputs into trustworthy, evidence-grounded decision support.

RESULTS AND DISCUSSION

The analytical dataset consisted of 1,284,672 AIS messages collected from 146 vessels operating within the selected maritime observation area. Data-quality screening identified 64,868 records containing duplicated transmissions, impossible coordinates, unresolved timestamps, or implausible movement values. A total of 1,219,804 messages, representing 94.95% of the original dataset, were retained after preprocessing. Chronological reconstruction produced 2,418 vessel-trajectory segments suitable for anomaly analysis. The analytical pipeline identified 214 candidate anomaly events, of which 96 were subsequently confirmed by expert review as sufficiently supported for inclusion in the report-generation evaluation. The confirmed events included abnormal speed changes, unexpected course deviations, prolonged transmission gaps, unusual stopping behavior, trajectory departures, and combinations of multiple behavioral irregularities.

Table 1. Descriptive Characteristics of AIS Data and Identified Anomaly Events

Data Component	Frequency	Percentage/Rate
Raw AIS Messages	1,284,672	100.00%
Retained AIS Messages	1,219,804	94.95%
Excluded/Flagged Messages	64,868	5.05%
Vessels Represented	146	—

Reconstructed Trajectory Segments	2,418	—
Candidate Anomaly Events	214	8.85% of segments
Expert-Validated Anomaly Events	96	44.86% of candidates
Non-validated/Contextually Explained Events	118	55.14% of candidates

The descriptive results indicate that raw anomaly detection alone produced substantially more candidate events than were retained after contextual and expert validation. More than half of the algorithmically identified cases were subsequently interpreted as operationally explainable, weakly supported, or insufficiently reliable for formal anomaly reporting. This finding demonstrates the practical importance of separating statistical deviation from operationally meaningful anomaly. The evidence also supports the inclusion of contextual validation as an explicit layer between machine detection and natural-language report generation because direct conversion of every anomalous score into a formal report would have produced a large number of potentially misleading outputs.

The anomaly-detection component showed its strongest performance for abrupt and clearly measurable changes in vessel behavior. Sudden changes in speed, large course deviations over short time intervals, and extended AIS transmission gaps were identified more consistently than context-dependent anomalies involving gradual trajectory deviation or unusual stopping patterns. Expert review revealed that several initially suspicious movements occurred near anchorage zones, traffic-separation boundaries, or areas where vessels commonly altered speed and heading. These findings indicate that the meaning of an anomaly was influenced by the relationship between movement characteristics and operational context rather than by numerical deviation alone.

The contextual interpretation stage reduced false-positive reporting by introducing vessel-specific, spatial, and temporal information before generating narrative outputs. Events with high anomaly scores but plausible contextual explanations were either downgraded or explicitly marked as requiring further review rather than described as confirmed operational anomalies. The reporting mechanism therefore demonstrated a layered reasoning structure in which algorithmic abnormality functioned as a screening signal, contextual evidence determined interpretive relevance, and generative AI communicated the final evidence package in human-readable form. This sequence reduced the risk of equating a statistical outlier with a safety, security, or compliance violation.

The evaluated anomaly-detection model achieved an overall precision of 0.86, recall of 0.82, F1-score of 0.84, and ROC-AUC of 0.91 when assessed against the expert-validated anomaly set. Performance differed according to anomaly category. Speed-related anomalies achieved the highest F1-score at 0.90, followed by abrupt course deviations at 0.87 and transmission-gap anomalies at 0.85. Trajectory-deviation events achieved an F1-score of 0.78, while unusual stopping behavior produced the lowest category-specific F1-score at 0.73. The lower performance for these context-dependent categories reflected the difficulty of distinguishing abnormal behavior from legitimate navigational or operational activity using movement features alone.

Table 2. Performance of the AIS Anomaly-Detection Component

Anomaly Category	Precision	Recall	F1-Score
Abnormal Speed Change	0.92	0.88	0.90
Abrupt Course Deviation	0.89	0.85	0.87
AIS Transmission Gap	0.87	0.83	0.85
Trajectory Deviation	0.80	0.76	0.78
Unusual Stopping Behavior	0.75	0.71	0.73

Overall Model	0.86	0.82	0.84
---------------	------	------	------

Evaluation of the generated reports showed stronger performance after the evidence-grounding and provenance mechanisms were incorporated into the refined prototype. Expert ratings on a five-point scale increased for factual consistency from 4.08 to 4.72, traceability from 3.64 to 4.76, contextual relevance from 3.91 to 4.55, and uncertainty communication from 3.22 to 4.61. Readability remained high in both versions, increasing only modestly from 4.31 to 4.53. The largest improvement concerned traceability and uncertainty representation, indicating that the main value of the refined architecture was not merely more fluent text but stronger alignment between narrative statements and underlying analytical evidence.

Paired-sample comparisons were conducted on expert evaluations of the first and refined report-generation prototypes across the 96 validated anomaly cases. The refined mechanism produced significantly higher factual-consistency scores, $t(95) = 8.74$, $p < .001$, Cohen's $d = 0.89$, and significantly higher evidence-traceability scores, $t(95) = 12.16$, $p < .001$, $d = 1.24$. Contextual relevance also improved significantly, $t(95) = 7.68$, $p < .001$, $d = 0.78$, while uncertainty communication showed the largest standardized improvement, $t(95) = 13.42$, $p < .001$, $d = 1.37$. Readability improved significantly but with a smaller effect, $t(95) = 2.84$, $p = .006$, $d = 0.29$.

Table 3. Inferential Comparison of Initial and Refined AI-Generated Reports

Evaluation Dimension	Initial M (SD)	Refined M (SD)	t(95)	p	Cohen's d
Factual Consistency	4.08 (0.54)	4.72 (0.31)	8.74	< .001	0.89
Evidence Traceability	3.64 (0.67)	4.76 (0.28)	12.16	< .001	1.24
Contextual Relevance	3.91 (0.59)	4.55 (0.36)	7.68	< .001	0.78
Uncertainty Communication	3.22 (0.71)	4.61 (0.35)	13.42	< .001	1.37
Readability	4.31 (0.43)	4.53 (0.39)	2.84	.006	0.29

The inferential results suggest that the most substantial benefit of the refined mechanism occurred in dimensions directly related to analytical reliability rather than linguistic presentation. The effect sizes for evidence traceability and uncertainty communication were considerably larger than the effect observed for readability. This pattern supports the design decision to prioritize structured evidence, provenance, and calibrated uncertainty rather than relying primarily on improvements in generative-language fluency. The analysis does not establish that the system can replace maritime analysts; it demonstrates that an evidence-grounded architecture can improve the quality of AI-assisted reporting within the evaluated artifact.

Correlation analysis showed that anomaly-confidence scores were moderately associated with expert assessments of anomaly plausibility ($r = .61$, $p < .001$), indicating that higher model confidence generally corresponded with stronger domain-expert agreement. Evidence completeness was strongly related to report factual consistency ($r = .72$, $p < .001$) and traceability ($r = .77$, $p < .001$). Contextual-information completeness was positively associated with expert ratings of report relevance ($r = .68$, $p < .001$). Report readability showed a weaker relationship with expert confidence in the analytical conclusion ($r = .29$, $p = .004$), suggesting that linguistic fluency alone was not a strong indicator of report reliability.

The relationship between evidence completeness and report quality became particularly visible in events involving multiple interacting anomalies. Reports generated from evidence packages containing vessel movement, anomaly scores, contextual location information, and trajectory comparison were judged more reliable than reports based only on an anomaly label

and a short sequence of AIS values. Expert comments repeatedly emphasized that analytically useful reports required enough evidence to distinguish what was directly observed from what had been inferred. The data therefore support the architectural principle that report generation should occur only after a sufficiently complete evidence package has been constructed.

Case A involved a cargo vessel that exhibited a rapid reduction in speed followed by a substantial course alteration. The initial anomaly model assigned a high score to both behaviors and classified the event as potentially abnormal. Contextual analysis showed that the trajectory occurred close to a port-approach area where speed reductions and course adjustments were common. The first-generation report described the event as an “abnormal maneuver requiring investigation,” while the refined system reported that the movement was statistically unusual relative to the vessel’s preceding trajectory but potentially consistent with port-approach navigation. Expert reviewers rated the refined report more highly because it preserved the anomaly signal while communicating a plausible alternative explanation and appropriate uncertainty.

Case B involved a vessel with a prolonged AIS transmission gap followed by reappearance at a location inconsistent with continuous movement inferred from the preceding trajectory. The evidence package contained the last valid position, transmission-gap duration, subsequent position, estimated travel distance, and model confidence. The refined report distinguished the directly observed transmission gap from the inferred discontinuity in vessel movement and explicitly stated that AIS data alone could not determine the reason for the interruption. Case C involved an unusual stop outside a frequently used anchorage area. The report identified the spatial deviation and stopping duration but refrained from attributing motive or regulatory significance, directing the analyst instead to review contextual maritime information.

Table 4. Summary of Embedded AIS Anomaly Cases

Case	Detected Event	Initial Reporting Risk	Refined Reporting Outcome
A	Speed reduction and course alteration	Overstatement of abnormal maneuver	Contextualized as statistically unusual but potentially operational
B	Prolonged AIS transmission gap	Implicit attribution of suspicious behavior	Observation separated from inference; cause left undetermined
C	Unusual stopping behavior	Unsupported implication of violation	Location and duration reported; operational significance reserved for analyst

The case analyses demonstrate that the main challenge in AI-assisted AIS reporting was not always detecting abnormal patterns but controlling the interpretation generated after detection. Case A showed that mathematically unusual behavior can be operationally reasonable when geographic context is considered. Case B illustrated the importance of distinguishing missing data from intentional vessel behavior. Case C demonstrated that even a correctly detected anomaly does not provide sufficient evidence for claims regarding motive, safety risk, or regulatory violation. These examples confirm the necessity of separating observed data, analytical inference, and operational interpretation within the generated report.

The refined architecture reduced interpretive overreach by constraining the language model to a structured evidence package and requiring explicit uncertainty when the data did not justify a definitive conclusion. Expert reviewers consistently preferred reports that acknowledged alternative explanations and identified the limits of available information. Human validation remained particularly important for complex events involving contextual knowledge not represented in the AIS dataset. The case findings therefore support a human-in-the-loop model in which AI accelerates analysis and documentation while professional analysts retain responsibility for consequential interpretations.

The combined findings indicate that Artificial Intelligence can provide substantial value in the mechanism for generating AIS anomaly-analysis reports when anomaly detection, contextual reasoning, explainability, evidence provenance, and controlled language generation are integrated as separate but connected analytical stages. The detection model achieved satisfactory overall performance, yet expert validation showed that algorithmic abnormality alone was insufficient for reliable reporting. Report quality improved most strongly when the system linked narrative statements directly to structured AIS evidence and explicitly represented uncertainty. The findings therefore support an end-to-end architecture rather than a simple anomaly-detector-to-language-model pipeline.

The results also indicate that generative AI should function as an evidence-constrained analytical communicator rather than an autonomous interpreter of vessel intent or operational significance. Fluent language did not guarantee analytical reliability, while traceability and contextual evidence were strongly associated with expert confidence in the generated reports. The most defensible conceptual model therefore positions AI as a support mechanism that identifies patterns, structures evidence, and produces standardized narratives for human review. This architecture offers a plausible pathway for improving the efficiency, consistency, and transparency of vessel AIS anomaly analysis without transferring final maritime judgment from accountable human professionals to generative systems.

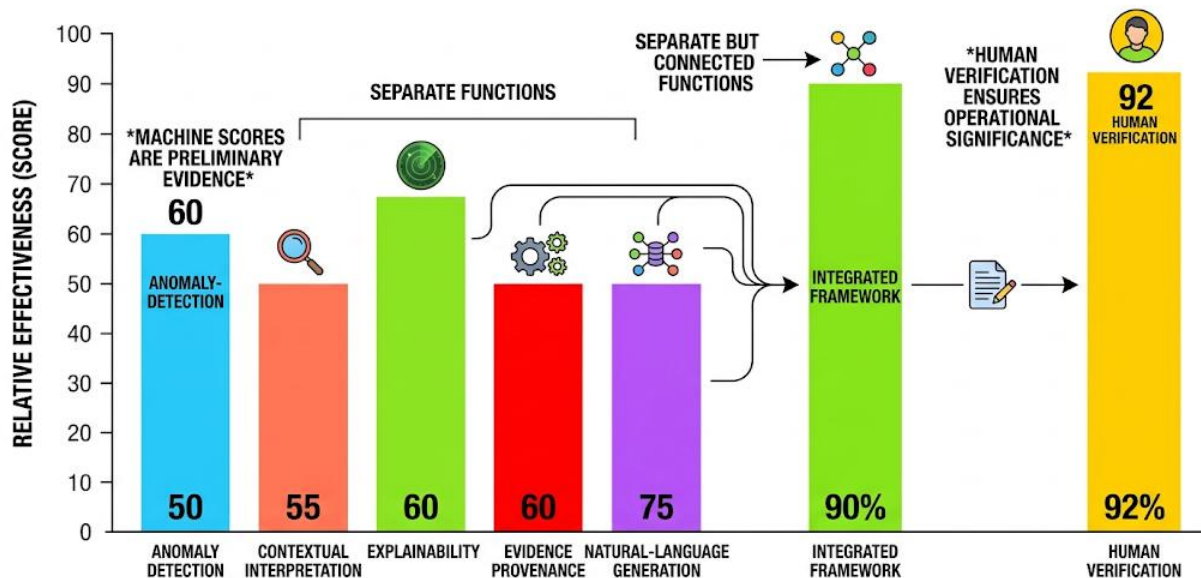


Figure 1. Optimizing AIS Anomaly Analysis: Integrated Framework Effectiveness

The findings demonstrate that the value of Artificial Intelligence in AIS anomaly analysis extends beyond the detection of unusual vessel behavior. The proposed mechanism performed most effectively when anomaly detection, contextual interpretation, explainability, evidence provenance, natural-language generation, and human verification were treated as separate but connected functions. The anomaly-detection component achieved satisfactory discriminatory performance, yet expert validation showed that many statistically unusual patterns could not automatically be interpreted as operationally significant anomalies. The results therefore support an analytical architecture in which machine-generated anomaly scores serve as preliminary evidence rather than final maritime judgments.

The descriptive results revealed a substantial reduction between candidate anomalies identified algorithmically and events ultimately retained after contextual and expert validation. This pattern indicates that numerical abnormality alone provides an incomplete representation of vessel behavior. Speed reduction, heading changes, transmission gaps, and unusual stopping patterns may reflect operational circumstances, geographical constraints, data-quality problems, or routine navigational behavior. Contextual validation consequently emerged as an essential

intermediary process that protected the reporting mechanism from transforming every statistical deviation into an unjustified warning or investigative conclusion.

The quality of AI-generated reports improved most substantially after evidence-grounding, provenance, and uncertainty controls were introduced. Factual consistency, evidence traceability, contextual relevance, and uncertainty communication demonstrated larger improvements than linguistic readability. This result is important because it shows that the core problem in generative maritime reporting is not simply producing fluent prose. A report may be grammatically sophisticated and easy to read while remaining analytically unreliable if its claims cannot be traced to AIS observations or model outputs. Reliable reporting therefore depends more strongly on evidence architecture than on linguistic sophistication.

The embedded cases reinforced the importance of separating observations, calculations, model inferences, and operational interpretations. A detected trajectory deviation could be statistically valid yet operationally explainable, while a transmission gap could be directly observed without providing evidence regarding its cause. The refined reporting mechanism handled these distinctions more appropriately by communicating uncertainty and limiting unsupported interpretations. The findings position AI as a mechanism for accelerating anomaly analysis and documentation while preserving the role of maritime professionals in determining the operational significance of complex vessel behavior.



Figure 2. Navigating the Digital Landscape

The findings are broadly consistent with existing AIS anomaly-detection research that demonstrates the capacity of machine-learning and deep-learning models to identify unusual trajectories, speed variations, course deviations, transmission irregularities, and other behavioral patterns. Previous approaches have primarily concentrated on improving classification, reconstruction, prediction, or anomaly-scoring performance. The present study supports the technical value of those approaches but extends the analytical problem beyond detection accuracy. The findings indicate that operational usefulness depends on what happens after an anomaly has been detected, particularly how evidence is contextualized, explained, and translated into information that maritime analysts can evaluate.

The results also correspond with explainable-AI research emphasizing the importance of transparency in high-stakes analytical environments. An anomaly score without information about the variables or trajectory characteristics responsible for the judgment offers limited value to an operator who must decide whether further investigation is warranted. The present mechanism extends explainability by connecting model reasoning to a structured reporting process. Evidence provenance becomes particularly important because the user must be able to distinguish a directly observed AIS fact from a calculated feature, probabilistic model output, contextual interpretation, or natural-language synthesis.

The study differs from conventional automated-reporting approaches that emphasize summarization efficiency or natural-language fluency as primary performance objectives. Maritime anomaly reports require a stronger relationship between linguistic output and evidentiary source material because even a minor unsupported statement may alter the perceived significance of an event. The present findings show that improvements in traceability and uncertainty communication contributed more to expert confidence than improvements in readability alone. This distinction suggests that generative AI applications in maritime surveillance should be evaluated according to analytical fidelity rather than conventional text-generation quality alone.

The findings also extend emerging research that integrates AIS time-series information with language models. Existing multimodal or language-assisted approaches demonstrate that maritime trajectory information can be converted into textual explanations and situation summaries. The present study adds a governance-oriented layer by arguing that language generation should not receive unrestricted authority to interpret vessel intent, illegality, or operational risk. A stronger model separates numerical detection, contextual reasoning, evidence packaging, and linguistic generation so that each stage can be independently examined and validated. This architecture reduces the possibility that generative fluency will conceal analytical uncertainty.

The findings signal that the central challenge of AI-assisted maritime analysis is increasingly shifting from pattern recognition toward trustworthy interpretation. Modern machine-learning systems can detect complex deviations within large AIS datasets, yet the operational problem remains whether those deviations can be explained in ways that analysts can understand, verify, and responsibly use. The transition from detection to interpretation represents an important stage in the maturation of maritime AI. Technical performance remains necessary, but it is no longer sufficient when model outputs enter operational reporting workflows.

The large number of candidate anomalies that were filtered through contextual validation signals the limitations of purely data-driven abnormality. AIS systems describe vessel movement, but they do not directly reveal intention, operational necessity, environmental conditions, or regulatory context. A trajectory may appear anomalous because the model has incomplete knowledge of port operations, traffic-separation schemes, weather avoidance, maintenance requirements, or navigational decisions. Anomaly interpretation therefore requires epistemic restraint: the system must acknowledge the difference between what the data show and what analysts may reasonably infer from those data.

The improvement produced by provenance controls signals that trust in AI should be constructed through verifiability rather than persuasive language. Generative systems are capable of producing coherent narratives even when the underlying evidence is incomplete. Maritime analysis becomes particularly vulnerable when linguistic confidence is mistaken for evidentiary strength. A trustworthy reporting mechanism must therefore expose the analytical path from raw AIS records to derived features, anomaly scores, contextual interpretation, and final narrative statements. Transparency should be embedded within the architecture rather than added as an explanatory feature after generation.

The persistence of human validation signals that automation should be understood as cognitive augmentation rather than professional substitution. Maritime analysts contribute contextual knowledge, operational experience, regulatory understanding, and judgment that may not be contained within the AIS dataset. AI can reduce repetitive workload, organize evidence, detect patterns, and generate standardized descriptions, but it cannot automatically reconstruct all relevant maritime circumstances. Human-in-the-loop design therefore represents a functional requirement rather than merely an ethical precaution in the current stage of AI-assisted vessel anomaly analysis.

The primary operational implication is that maritime organizations should avoid deploying anomaly-detection models as isolated alert generators. A high volume of uncontextualized

alarms may increase rather than reduce analyst workload because operators must manually reconstruct the significance of each alert. An integrated reporting mechanism can improve operational efficiency by presenting the detected anomaly together with relevant vessel information, supporting trajectory evidence, contextual explanation, confidence indicators, and explicit uncertainty. The analyst can then focus on evaluating a structured case rather than assembling the analytical evidence from multiple systems.

The technological implication concerns the design of AI architecture. Generative AI should not be positioned at the beginning of the analytical process or given unrestricted access to raw AIS data with the expectation that it will independently discover and explain anomalies. Deterministic preprocessing, numerical feature extraction, dedicated anomaly-detection models, contextual reasoning, and evidence validation should precede language generation. The language model should operate on a constrained and structured evidence package. This arrangement assigns generative AI the task it performs most appropriately: communicating validated analytical information in a coherent and standardized form.

The governance implication concerns accountability, auditability, and traceability. Every consequential statement in an AI-assisted anomaly report should be recoverable to its supporting source, whether a raw AIS observation, derived feature, model result, contextual rule, or analyst judgment. The system should preserve logs showing how a report was generated, what evidence was available, what model version was used, and which statements were subsequently modified by a human reviewer. Such mechanisms would support quality assurance, post-event evaluation, model improvement, and institutional accountability in environments where analytical reports may inform operational decisions.

The methodological implication extends beyond maritime AIS applications. The findings demonstrate that evaluation of generative AI in data-intensive operational settings should include factual consistency, provenance, uncertainty calibration, contextual relevance, and human usability in addition to conventional natural-language quality. A highly readable report is not necessarily a trustworthy report. Researchers developing AI-assisted analytical systems should therefore separate model-detection metrics from report-generation metrics and from operational-utility assessments. This multidimensional evaluation offers a more realistic representation of system performance than a single accuracy or language-quality score.

The stronger performance of clearly measurable anomalies can be explained by the relationship between data structure and behavioral ambiguity. Abrupt changes in speed, substantial course deviations, and prolonged transmission gaps produce relatively visible numerical patterns that can be distinguished from surrounding observations. Unusual stops and gradual trajectory deviations depend more strongly on spatial and operational context. The anomaly-detection model consequently performed better when the abnormal behavior possessed a clear kinematic signature and less consistently when the same numerical pattern could reasonably occur during normal vessel operations.

The reduction in false-positive reporting after contextual validation can be explained by the contextual dependence of maritime behavior. Vessels do not move within homogeneous geographical environments. Ports, anchorages, narrow channels, traffic-separation schemes, pilot boarding areas, and congested waters naturally generate behaviors that might appear unusual when compared with open-sea movement. Context therefore changes the interpretation of the same numerical feature. The architecture improved because it evaluated abnormality relative to surrounding operational circumstances rather than treating vessel behavior as context-free time-series data.

The large improvement in evidence traceability and uncertainty communication can be explained by the constraints imposed on the generative model. An unconstrained language model is optimized to produce plausible linguistic continuations, not necessarily to maintain perfect alignment with analytical evidence. Structured evidence packages narrowed the informational space from which the model could generate conclusions. Explicit instructions to separate

observation from inference and to acknowledge insufficient evidence further reduced unsupported statements. The resulting reports were more reliable because the generation process became dependent on verified data rather than linguistic probability alone.

The continued importance of human review can be explained by the incompleteness of AIS as a representation of maritime reality. AIS provides valuable positional and navigational information, but it does not capture every factor affecting vessel behavior. Weather conditions, traffic-control instructions, mechanical problems, commercial operations, pilotage, security concerns, or communication failures may explain an observed pattern without being visible in the AIS stream. Human analysts can integrate external information and professional experience that remain beyond the analytical boundaries of the model. Hybrid intelligence therefore becomes necessary when operational meaning exceeds the information contained in the primary dataset.

Future research should evaluate the proposed mechanism using larger and more heterogeneous AIS datasets representing different vessel classes, geographic regions, traffic conditions, and operational environments. Model performance observed within one maritime area may not transfer directly to ports, congested straits, offshore zones, or open-ocean contexts. Multi-region validation would help determine which anomaly features remain stable across environments and which require local calibration. Temporal validation should also examine whether models retain performance when vessel traffic patterns change over seasons or operational periods.

Future studies should integrate additional maritime information sources to improve contextual reasoning. Radar tracks, vessel registries, meteorological data, bathymetry, port information, navigational warnings, traffic-separation schemes, and historical vessel behavior could provide evidence unavailable in AIS messages alone. Multisource integration could reduce false positives and enable more precise distinction between statistical irregularity and operational significance. Such integration should preserve source provenance so that analysts remain able to identify which evidence contributed to each generated interpretation.

Future system development should examine alternative combinations of anomaly-detection models, explainability techniques, and language-generation architectures. Comparative research could determine whether specialized smaller language models, retrieval-augmented generation, constrained decoding, structured templates, or larger general-purpose models provide the strongest balance among factual accuracy, computational efficiency, traceability, and usability. Evaluation should also include hallucination rates, omission errors, calibration of confidence language, latency, processing cost, and robustness under incomplete or corrupted input data.

Future operational validation should involve maritime analysts interacting with the system during realistic decision-support tasks. Measures should extend beyond report accuracy to include analyst workload, time to interpretation, trust calibration, decision consistency, error detection, and the likelihood of automation bias. Longitudinal human-factor studies would be particularly valuable for determining whether analysts become overly dependent on generated narratives or learn to use them as evidence-organizing tools. Interface design should make uncertainty, evidence provenance, and model limitations visible rather than hiding them behind polished natural-language outputs.

Future governance research should develop explicit protocols for accountability in AI-assisted maritime reporting. Standards are needed for model validation, audit trails, version control, human approval, retention of supporting evidence, and documentation of modifications made after generation (Yin et al., 2025). High-consequence interpretations should require stronger human oversight than routine descriptive reports. Organizational policies should also define which conclusions AI systems are permitted to generate and which remain exclusively within the authority of qualified personnel.

Future research should ultimately move from proof-of-concept reporting toward a validated human-AI analytical ecosystem. The most promising direction is not a fully autonomous system that determines vessel intent, threat level, or regulatory violation from AIS data alone. A more credible architecture combines machine detection, contextual computation, explainable reasoning, evidence-grounded language generation, and accountable human judgment. Such development would allow AI to increase the speed and consistency of AIS anomaly analysis while preserving the evidentiary discipline, uncertainty awareness, and professional responsibility required for trustworthy maritime decision support.

CONCLUSION

The most important finding of this study is that the effectiveness of Artificial Intelligence in vessel AIS anomaly analysis depends not only on anomaly-detection accuracy but also on the quality of the mechanism that transforms detected events into evidence-grounded analytical reports. The proposed architecture showed that contextual validation, explainability, provenance, and uncertainty representation substantially strengthened report reliability by preventing statistically unusual vessel behavior from being automatically interpreted as operationally significant. The findings also demonstrated that fluent natural-language output is insufficient as a measure of quality because readable reports may still contain unsupported or overstated conclusions. The strongest analytical performance emerged when observed AIS facts, calculated indicators, model inferences, contextual interpretations, and human judgments were explicitly separated within the reporting workflow. AI therefore functions most credibly as an analytical support mechanism that organizes and communicates verified evidence while leaving consequential maritime interpretation under accountable human supervision.

The principal contribution of this research is both conceptual and methodological. Conceptually, the study advances an end-to-end model of AI-supported AIS anomaly reporting that integrates data preprocessing, trajectory reconstruction, anomaly detection, contextual reasoning, explainable AI, evidence packaging, controlled language generation, and human verification within a single analytical architecture. This model shifts the focus of innovation from anomaly detection alone toward trustworthy report generation. Methodologically, the Design Science Research approach enabled the artifact to be developed and evaluated according to separate but complementary criteria for detection performance and report quality. The study therefore contributes a framework in which precision, recall, and F1-score are considered alongside factual consistency, traceability, contextual relevance, readability, and uncertainty communication. This multidimensional evaluation provides a more operationally meaningful basis for assessing AI-assisted maritime analytics than relying exclusively on predictive accuracy or linguistic fluency.

The main limitations of the study concern dataset scope, contextual coverage, model generalizability, and the bounded information available from AIS records. The evaluated mechanism cannot capture every operational factor influencing vessel behavior, including weather, port instructions, mechanical conditions, pilotage, or other external circumstances not represented in the AIS stream. Performance may also vary across vessel classes, geographic regions, traffic densities, anomaly categories, and different generative-AI architectures. Future research should therefore validate the mechanism using larger multi-region datasets, integrate supplementary maritime information such as radar, weather, vessel registries, and port data, and compare alternative anomaly-detection and language-generation configurations. Further studies should also examine hallucination rates, confidence calibration, analyst workload, automation bias, decision consistency, and long-term human AI interaction under realistic operational conditions. Such work is necessary before AI-generated AIS anomaly reports can be considered sufficiently robust for high-consequence maritime decision support.

DECLARATION OF AI AND AI ASSISTED TECHNOLOGIES IN THE WRITING PROCESS

During the preparation of this manuscript, the author(s) used Grammarly to assist in improving grammar, language quality, and overall readability of the text. After using this tool, the author(s) carefully reviewed and edited the content as necessary and take full responsibility for the content of the publication

AUTHOR CONTRIBUTIONS

Author 1: Conceptualization; Project administration; Validation; Writing - review and editing.

Author 2: Conceptualization; Data curation; In-vestigation.

Author 3: Data curation; Investigation.

DECLARATION OF COMPETING INTEREST

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

REFERENCES

- Afroogh, S., Akbari, A., Malone, E., Kargar, M., & Alambeigi, H. (2024). Trust in AI: Progress, challenges, and future directions. *Humanities and Social Sciences Communications*, 11(1), 1568. <https://doi.org/10.1057/s41599-024-04044-8>
- Ahmadzadeh, S., Karthick, G., & Alsafi, T. (2026). Large Language Models for Future Internet Ecosystems: A Taxonomy-Based Review of Smart IoT, Edge Intelligence, and Autonomous AI. *Future Internet*, 18(8), 393. <https://doi.org/10.3390/fi18080393>
- Al-Refai, O., Shahbaz, I., & Hammad, E. (2026). Composable Trust in Agentic AI: Bridging Architectural Capability and System-Level Assurance. *2026 56th Annual IEEE International Conference on Dependable Systems and Networks Workshops (DSN-W)*, 17–20. <https://doi.org/10.1109/DSN-W70714.2026.00023>
- Atias, O., & Mawasi, A. (2025). Conceptualizing AI literacies for children and youth: A systematic review on the design of AI literacy educational programs. *Computers and Education: Artificial Intelligence*, 9, 100491. <https://doi.org/10.1016/j.caeai.2025.100491>
- Babaei, F., Boojarjomehry, R. B., Ravandi, Z. K., & Pishvaie, M. R. (2024). Enabling digital transformation of dynamic location-inventory-routing optimization in natural gas-to-product and energy networks via a domain-adaptable ontological agent-based framework. *Advanced Engineering Informatics*, 60, 102380. <https://doi.org/10.1016/j.aei.2024.102380>
- Bai, H., & Yao, Z. (2026). Building the oracle: Power, culture, and the organizational domestication of In-house AI in Chinese newsrooms. *Chinese Journal of Communication*, 1–19. <https://doi.org/10.1080/17544750.2026.2712458>
- Bartelheimer, C., Heinz, D., Hönigsberg, S., Siemon, D., Li, M. M., Strohmman, T., Poeppelbuss, J., & Peters, C. (2025). Conceptualizing hybrid intelligent service ecosystems. *Electronic Markets*, 35(1), 63. <https://doi.org/10.1007/s12525-025-00798-4>

- Che, H., Wang, S., Yao, L., & Gu, Y. (2025). A comprehensive perspective on electric vehicles as evolutionary robots. *Frontiers in Robotics and AI*, 12, 1499215. <https://doi.org/10.3389/frobt.2025.1499215>
- Chen, J., Liu, Q., Wang, Y., & Wang, L. (2026). A Physical-Prior Guided UAV Perception and Sailability Assessment Framework for Main Route Navigation Under Fog Conditions. *Drones*, 10(5), 367. <https://doi.org/10.3390/drones10050367>
- Chin, W. (2026). A Janus-faced war: Exploring the role of technology in shaping outcomes in the Russia–Ukraine conflict. *Defence Studies*, 1–22. <https://doi.org/10.1080/14702436.2026.2715485>
- De Larmelina, S., Da Silva, A. L., & Risso, L. A. (2026). Conceptualizing the Digital Twin: An Analysis of Definitions, Applications, Technologies, Benefits and Challenges. *Journal of Advanced Manufacturing Systems*, 25(02), 253–279. <https://doi.org/10.1142/S0219686726500125>
- Fasouli, P., El Bahnasawi, M., Gebser, M., & Kyamakya, K. (2025). Conceptualizing Validation Systems for Explainable AI: A Design Approach. In F. Al Machot, M. T. Horsch, & S. Scholze (Eds.), *Designing the Conceptual Landscape for a XAIR Validation Infrastructure* (Vol. 1375, pp. 6–24). Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-89274-5_2
- Filep, S., Kondja, A., Wong, C. C. K., Weber, K., Moyle, B. D., & Skavronskaya, L. (2024). The role of technology in users’ wellbeing: Conceptualizing digital wellbeing in hospitality and future research directions. *Journal of Hospitality Marketing & Management*, 33(5), 583–601. <https://doi.org/10.1080/19368623.2023.2290626>
- Johnson, B. T. (2026). War, state, algorithm: Conceptualizing a social mode of total algorithmic war. *Critical Military Studies*, 1–20. <https://doi.org/10.1080/23337486.2026.2705621>
- Kalbande, D., Hemke, D., & Motewar, N. (2025). Artificial intelligence and the five laws: A new vision for library science. *Library Hi Tech News*, 42(4), 1–3. <https://doi.org/10.1108/LHTN-01-2025-0005>
- Kalli, C., Ziti, S., & Kharmoum, N. (2026). Exploring Graph Theory Applications in Digital Health Systems: Remote Patient Monitoring and Virtual Consultations. In W. Rhalem, N. Al Idrissi, & M. Lazaar (Eds.), *HealthTech “Global Summit of Digital Health”* (Vol. 1587, pp. 196–207). Springer Nature Switzerland. https://doi.org/10.1007/978-3-032-01970-7_20
- Kaur, J., & Taak, S. (2026). Agentic AI in Cyber Warfare: Reassessing International Humanitarian Implications and Emerging Challenges. In B. E. Elbaghazaoui, K. Benabbes, & T. El Moudden (Eds.), *Advances in Computational Intelligence and Robotics* (pp. 111–140). IGI Global Scientific Publishing. <https://doi.org/10.4018/979-8-2600-0413-5.ch004>
- Khudhur, A. A., & Al-Alawi, A. I. (2024). The Use of Machine Learning to Forecast Financial Performance: A Literature Review. *2024 ASU International Conference in Emerging Technologies for Sustainability and Intelligent Systems (ICETISIS)*, 1–6. <https://doi.org/10.1109/ICETISIS61505.2024.10459393>
- Li, Z., Wu, W., Guo, Y., Sun, J., & Han, Q.-L. (2025). Embodied Multi-Agent Systems: A Review. *IEEE/CAA Journal of Automatica Sinica*, 12(6), 1095–1116. <https://doi.org/10.1109/JAS.2025.125552>

- Lin, C., Zou, W., & Liu, Z. (2026). Differentiating Trust Profiles: An Analysis of Users' Situational Trust in Generative Artificial Intelligence and Its Antecedents. *International Journal of Human-Computer Interaction*, 42(4), 2629–2646. <https://doi.org/10.1080/10447318.2025.2530073>
- Liu, X., & Ye, Z. (2026). The forgotten middle: How moderate self-efficacy amplifies the threat of AI through job insecurity. *Frontiers in Psychology*, 16, 1734254. <https://doi.org/10.3389/fpsyg.2025.1734254>
- Loganathan, P., Muthumalathi, M., & Pankajavalli, P. B. (2025). Towards 6G technologies: Insight into architecture, communication, AI applications and use cases. In *Advances in Computers* (Vol. 139, pp. 561–583). Elsevier. <https://doi.org/10.1016/bs.adcom.2025.03.018>
- Luong, A., Kumar, N., & Lang, K. R. (2026). Human-machine collaboration and algorithmic decision-making: The role of organizational incentives and AI capabilities. *Information & Management*, 63(6), 104384. <https://doi.org/10.1016/j.im.2026.104384>
- Mansfield, K. L., Ghai, S., Hakman, T., Ballou, N., Vuorre, M., & Przybylski, A. K. (2025). From social media to artificial intelligence: Improving research on digital harms in youth. *The Lancet Child & Adolescent Health*, 9(3), 194–204. [https://doi.org/10.1016/S2352-4642\(24\)00332-8](https://doi.org/10.1016/S2352-4642(24)00332-8)
- Meyer, F. S. D. E., & Mariosi, L. A. (2025). Survival by Design: The BWC's Autopoietic Response at Fifty to a Disaggregated Biosecurity Ecosystem. *Health Security*, 23(5), 292–299. <https://doi.org/10.1177/23265094251381803>
- Michael, K., Vogel, K. M., Pitt, J., & Zafeirakopoulos, M. (2025). Artificial Intelligence in Cybersecurity: A Socio-Technical Framing. *IEEE Transactions on Technology and Society*, 6(1), 15–30. <https://doi.org/10.1109/TTS.2024.3460740>
- Nazirov, M., Aripova, Z., & Khurammov, B. (2026). *Artificial intelligence and big data analytics in educational policy and decision-making*. 100007. <https://doi.org/10.1063/5.0330290>
- Purohit, P., Goswami, C., & Hiran, K. K. (2025). Conceptualizing a Multi-Stage Hybrid Approach for AI-Assisted Prediction of Cardiovascular Disease. *2025 International Conference on Sustainability, Innovation & Technology (ICSIT)*, 1–7. <https://doi.org/10.1109/ICSIT65336.2025.11294336>
- Qiang, H., Niu, W., Peng, X., & Li, H. (2026). TrajCleaner: An end-to-end framework for systematic cleaning and completing vessel trajectory data. *Transportation Research Part E: Logistics and Transportation Review*, 216, 105131. <https://doi.org/10.1016/j.tre.2026.105131>
- Rigotti, C., & Fosch-Villaronga, E. (2024). Fairness, AI & recruitment. *Computer Law & Security Review*, 53, 105966. <https://doi.org/10.1016/j.clsr.2024.105966>
- Rizk, M., Rosu, D., & Fox, M. (2024). Semantic Computing for Organizational Effectiveness: From Organization Theory to Practice through Semantics-Based Modelling. *2024 IEEE 18th International Conference on Semantic Computing (ICSC)*, 89–92. <https://doi.org/10.1109/ICSC59802.2024.00020>
- Roy, S. K., Tehrani, A. N., Pandit, A., Apostolidis, C., & Ray, S. (2025). AI-capable relationship marketing: Shaping the future of customer relationships. *Journal of Business Research*, 192, 115309. <https://doi.org/10.1016/j.jbusres.2025.115309>
- Salam, M. A., Rayun, S. M. N., Leong, V. S., Islam, W., Chowdhury, S. A., Sahabuddin, M., Ahmed, M., & Islam, M. (2026). From Lab to Marketplace: Technology Transfer and

- Commercial Applications of Deepfake Technology. In J. Remondes (Ed.), *Advances in Computational Intelligence and Robotics* (pp. 333–368). IGI Global Scientific Publishing. <https://doi.org/10.4018/979-8-3373-8084-1.ch011>
- Shah, S. F. H., Arecco, D., Draper, H., Tiribelli, S., Harriss, E., & Martin, R. N. (2025). Ethical implications of artificial intelligence in skin cancer diagnostics: Use-case analyses. *British Journal of Dermatology*, 192(3), 520–529. <https://doi.org/10.1093/bjd/ljae434>
- Sharma, P., Aswar, P., Mulay, S., Wartkar, N., Wankhede, K., & Tiwari, A. (2025). Analyzing Centrality Measures in Network Graphs: Algorithms, Applications, and Heuristics. In D. Gupta, V. Kamble, V. Satpute, & A. Kothari (Eds.), *Paradigm Shifts in Communication, Embedded Systems, Machine Learning, and Signal Processing* (Vol. 2491, pp. 191–205). Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-90577-3_18
- Sienknecht, M., & Vetterlein, A. (2024). Conceptualizing responsibility in world politics. *International Theory*, 16(1), 26–49. <https://doi.org/10.1017/S1752971923000039>
- Sonawane, P. (2025). 11 Understanding artificial intelligence: An introduction, history, and foundations. In P. Kumar, V. Vivekanand, N. Pareek, & R. C. Dubey (Eds.), *Artificial Intelligence in Microbiology* (pp. 1–24). De Gruyter. <https://doi.org/10.1515/9783111548777-001>
- Speith, T., Crook, B., Mann, S., Schomäcker, A., & Langer, M. (2024). Conceptualizing understanding in explainable artificial intelligence (XAI): An abilities-based approach. *Ethics and Information Technology*, 26(2), 40. <https://doi.org/10.1007/s10676-024-09769-3>
- Tu, Y.-F., & Chen, J. (2025). Framework and Potential Applications of GAI-Based Language Learning Design. In Y.-J. Lan, G. Y. Qi, & D. Chun (Eds.), *AI-Mediated Language Education in the Metaverse Era* (pp. 161–186). Springer Nature Singapore. https://doi.org/10.1007/978-981-95-0245-5_9
- Tzanova, S. (2024). AI in Academic Libraries: Success, Pitfalls, Perceptions, and Why We Need AI Literacy. In I. Khamis (Ed.), *Advances in Library and Information Science* (pp. 19–44). IGI Global. <https://doi.org/10.4018/979-8-3693-1573-6.ch002>
- Voicu, R. C., Frischmann, A., Tanveer, M. H., Moghadam, A. A. A., Tello, Y., Oprea-Ilie, G., & Chang, Y. (2025). Biologically Inspired Networking (BIN): Leveraging Neurobiological Principles for Adaptive, Scalable, and Efficient Communication. 2025 *International Conference on Computing, Networking and Communications (ICNC)*, 655–661. <https://doi.org/10.1109/ICNC64010.2025.10993960>
- Vu, U. M., & Nguyen, L. T. (2025). Measuring the Impact of AI-Powered Virtual Try-On on Young Consumers’ Purchase Readiness. *Proceedings of the 2025 16th International Conference on E-Business, Management and Economics*, 306–314. <https://doi.org/10.1145/3760557.3760603>
- Wang, M., Liang, Y., Zhang, L., Tang, T., Aw, K. C., Qian, K., Dai, Z., Zhou, B., & Tang, L. (2026). Solid–Liquid Triboelectric Nanogenerators as Physicochemical Encoders for Intelligent Liquid Recognition. *Interdisciplinary Materials*, 5(4), 573–608. <https://doi.org/10.1002/idm2.70067>
- Wang, M., Wang, J., & Wang, H. (2025). Conceptualizing Resilience in Healthcare Systems Through Domain Modeling. *IEEE Transactions on Computational Social Systems*, 12(1), 113–127. <https://doi.org/10.1109/TCSS.2024.3451519>

- Yin, J., Yu, Z., & Wu, H. (2025). Ship trajectory prediction based on LSTM model with multi-scale convolution and attention mechanism. *Ocean Engineering*, 338, 122055. <https://doi.org/10.1016/j.oceaneng.2025.122055>
- Yu, S., Lin, J., Huang, J., & Zhan, Y. (2024). *A systematic review of changing conceptual to practice AI curation in museums: Text mining and bibliometric analysis*. 15th International Conference on Applied Human Factors and Ergonomics (AHFE 2024). <https://doi.org/10.54941/ahfe1004668>
- Yue, P., Wu, H., Liu, R., Gong, J., Xiang, L., Wang, K., & Teng, B. (2026). Conceptualizing the analysis-readiness for the next-generation SDI through the Open Geospatial Engine (OGE). *International Journal of Remote Sensing*, 47(7), 2962–2996. <https://doi.org/10.1080/01431161.2026.2625516>
-

Copyright Holder :

© Ade Tya Guritna et al. (2026).

First Publication Right :

© Journal of Computer Science Advancements

This article is under:

