

Intrusion Detection with Anomaly Representation Learning and Semantic Grounding across Enterprise Network Traffic

Daniel E. King, Anthony Gray, Harper I. Lewis

ILR School Statistics and Data Science, Cornell University, Ithaca, New York, USA

Corresponding author: daniel.e.king@cornell.edu

Abstract

The proliferation of sophisticated cyber threats within enterprise network environments has necessitated the deployment of advanced intrusion detection systems. While contemporary deep learning methodologies exhibit remarkable detection capabilities, their inherent opacity creates a significant semantic gap between statistical anomaly detection and human-interpretable security intelligence. This paper presents a comprehensive framework for explaining intrusion detection through the synergistic integration of anomaly representation learning and semantic grounding. By transforming raw, high-dimensional network traffic features into a structured latent space, the proposed methodology captures the underlying distributions of normal enterprise behavior. Subsequently, a semantic grounding mechanism maps these opaque latent representations onto a predefined ontology of human-understandable network concepts, enabling the automatic generation of transparent, context-aware explanations for identified anomalies. The architectural design leverages a dual-phase training paradigm, wherein unsupervised contrastive learning is first utilized to isolate anomalous vectors, followed by an alignment process that associates specific vector dimensions with semantic attributes such as irregular payload structures or anomalous temporal flow patterns. Extensive theoretical analysis and conceptual evaluation indicate that bridging the semantic gap not only enhances the trustworthiness of automated security systems but also substantially reduces the cognitive load imposed on security analysts during incident response. This research contributes a foundational paradigm shift in cybersecurity, moving beyond mere threat detection toward holistic, interpretable threat comprehension.

Keywords: Intrusion Detection; Representation Learning; Semantic Grounding; Explainable Artificial Intelligence; Network Security

1. Introduction

1.1 The Challenge of Enterprise Network Security

The contemporary enterprise network is a highly dynamic and heterogeneous ecosystem characterized by a massive influx of data packets, diverse communication protocols, and rapidly fluctuating user behaviors. As organizations increasingly adopt cloud computing infrastructures, Internet of Things devices, and distributed workforce models, the traditional enterprise perimeter has effectively dissolved. In this environment, identifying malicious activities requires continuous monitoring and analysis of network traffic at an unprecedented scale. Intrusion detection systems serve as the primary defensive mechanism, tasked with differentiating legitimate administrative actions and standard user operations from sophisticated cyber threats such as advanced persistent threats, lateral movement maneuvers, and zero-day exploits. The sheer volume and velocity of modern network traffic render manual inspection impossible, thereby driving the widespread adoption of automated, algorithm-driven security solutions [1]. However, the complexity of enterprise traffic poses significant challenges. Normal network behavior is not a static baseline but rather a continuously evolving distribution influenced by business cycles, software updates, and organizational changes. Consequently, security mechanisms must be inherently adaptable, capable of learning the intrinsic patterns of benign traffic while remaining highly sensitive to deviations indicative of malicious intent. Early rule-based systems, while transparent, struggle to scale against novel attack vectors, leading to the

integration of machine learning techniques designed to autonomously discover anomalous patterns within the data stream [2].

1.2 Limitations of Current Intrusion Detection Systems

The transition from signature-based detection to machine learning and specifically deep learning approaches has significantly improved the detection efficacy of intrusion detection systems [3]. Deep neural networks, particularly autoencoders, recurrent neural networks, and convolutional architectures, excel at mapping high-dimensional network flow data into lower-dimensional representations, effectively identifying complex, non-linear patterns of anomalous behavior [4]. Despite their high detection accuracy, these models suffer from a fundamental limitation: they operate as black boxes. When a deep learning model flags a specific sequence of network packets as an anomaly, it provides an output score or a probability metric but fails to articulate the underlying rationale for its decision. This lack of interpretability creates a substantial barrier to operational deployment in real-world security operations centers. Security analysts presented with binary alerts or continuous anomaly scores without context are forced to manually investigate the raw network logs to ascertain the nature of the threat [5]. This process is labor-intensive, error-prone, and exacerbates the well-documented phenomenon of alert fatigue. Furthermore, the inability to explain why an event was classified as anomalous severely undermines user trust in the system, particularly when dealing with high false positive rates common in anomaly detection architectures. Traditional explainable artificial intelligence techniques, such as feature attribution methods, often fall short in this domain because explaining a decision in terms of raw network features, such as the exact byte count of a packet or a specific flag setting, does not inherently convey the operational intent or the security context of the anomaly [6].

1.3 Proposed Approach and Contributions

To address the profound limitations of opaque anomaly detection, this paper introduces a novel conceptual framework that integrates anomaly representation learning with rigorous semantic grounding. The core thesis posits that effective explainability in intrusion detection requires translating statistical deviations in high-dimensional feature spaces into comprehensible, domain-specific concepts. Our methodology approaches this by first employing advanced representation learning techniques to encode network traffic into a disentangled latent space [7]. Unlike standard dimensionality reduction, this process is optimized to ensure that distinct types of network behavior occupy distinct, well-separated regions within the representation space. Following this, we introduce a semantic grounding layer. This mechanism maps specific regions or directional vectors within the latent space to an ontology of human-understandable security concepts, such as port scanning behavior, uncharacteristic data exfiltration volumes, or anomalous protocol handshakes [8]. By anchoring the abstract mathematical representations to a concrete semantic dictionary, the system can automatically generate descriptive explanations for any detected anomaly. This paper extensively details the theoretical underpinnings of this approach, the mathematical formulations required for semantic alignment, and the architectural design necessary for processing enterprise-scale network traffic. The primary contribution of this work is the formalization of a bridge between statistical anomaly detection and human cognitive models of network security, thereby offering a paradigm where machine learning systems act not merely as automated alerting mechanisms but as intelligent, communicative analytical partners.

2. Background and Related Work

2.1 Evolution of Intrusion Detection

The historical trajectory of intrusion detection reflects a continuous arms race between defensive engineering and offensive innovation. The earliest systems relied almost exclusively on signature-based methodologies, wherein security experts manually crafted rules corresponding to known attack patterns [9]. These systems were highly deterministic, transparent, and computationally efficient, yet they possessed a fatal flaw: they were inherently reactive. A signature-based system cannot detect a novel attack, commonly referred to as a zero-day exploit, until a corresponding rule has been formulated and deployed. Recognizing this limitation, the academic community and the security industry began exploring anomaly-based detection methodologies in the late 1990s and early 2000s [10]. Anomaly detection operates on the premise that malicious activity will inevitably manifest as a deviation from established baseline behavior. Early statistical methods, such as those utilizing principal component analysis and Markov models, demonstrated the viability of this approach but often struggled with the high dimensionality and non-stationary nature of

network traffic [11]. The advent of big data technologies and the exponential increase in computational power subsequently facilitated the application of shallow machine learning algorithms, including support vector machines, random forests, and k-nearest neighbors [12]. While these algorithms provided superior generalization capabilities compared to purely statistical methods, they heavily relied on extensive manual feature engineering. Security researchers were required to manually extract and select features such as connection duration, protocol types, and byte transfer rates, a process that was not only time-consuming but also susceptible to human bias and oversight [13].

2.2 Deep Representation Learning for Anomalies

The limitations of manual feature engineering in traditional machine learning catalyzed the adoption of deep learning techniques within the cybersecurity domain. Deep learning architectures possess the intrinsic capability to automatically discover hierarchical feature representations from raw or minimally processed data [14]. In the context of anomaly detection, representation learning typically employs unsupervised or semi-supervised paradigms, given the pervasive scarcity of comprehensively labeled attack data in real enterprise environments. Autoencoders represent the most prevalent architecture utilized for this purpose. An autoencoder is trained to compress input data into a low-dimensional bottleneck layer and subsequently reconstruct the original input. When trained exclusively on benign network traffic, the model learns the foundational manifold of normal behavior [15]. Consequently, when anomalous traffic is processed, the autoencoder fails to reconstruct it accurately, resulting in a high reconstruction error that serves as an anomaly indicator. More recent advancements have introduced variational autoencoders and generative adversarial networks to model the probability distribution of the latent space more robustly, thereby improving the distinction between subtle anomalies and rare but benign network events [16]. Furthermore, contrastive learning approaches have emerged as a powerful technique for representation learning. By explicitly optimizing the network to minimize the distance between similar benign samples while maximizing the distance between dissimilar samples or synthetically generated perturbations, contrastive learning produces highly discriminative latent spaces [17]. However, despite these algorithmic sophisticated advancements, the representations generated remain fundamentally abstract vectors of continuous numerical values, devoid of inherent semantic meaning.

2.3 Explainable Artificial Intelligence in Cybersecurity

As the reliance on opaque deep learning models grew, the critical need for explainable artificial intelligence became paramount. In the cybersecurity domain, decisions made by automated systems can lead to significant operational disruptions, such as the isolation of critical servers or the blocking of legitimate business communications [18]. Therefore, understanding the rationale behind an algorithmic decision is not merely an academic exercise but an operational necessity. Existing explainable artificial intelligence methods are generally categorized into intrinsic and post-hoc approaches. Intrinsic explainability involves utilizing models that are transparent by design, such as decision trees or linear regression; however, these models often lack the expressive capacity required to model complex network traffic [19]. Post-hoc methods attempt to explain the decisions of black-box models after the fact. Techniques such as Local Interpretable Model-agnostic Explanations and SHapley Additive exPlanations are widely employed to assign importance scores to the input features contributing to a specific prediction [20]. While these feature attribution methods have proven useful in domains like image recognition, their application in network intrusion detection presents unique challenges. Highlighting that a specific source port or a precise packet inter-arrival time was the primary contributor to an anomaly score often fails to provide actionable intelligence to a security analyst. The semantic gap remains unbridged because a collection of statistical feature weights does not equate to a coherent narrative regarding the attacker's methodology or objectives [21].

2.4 Semantic Grounding Methodologies

To transcend the limitations of feature-level attribution, recent research paradigms have begun exploring the concept of semantic grounding. Originating primarily from the fields of natural language processing and computer vision, semantic grounding involves linking abstract machine representations to structured, domain-specific knowledge bases or ontologies [22]. In the context of network security, this means associating the high-dimensional vectors produced by a representation learning model with concrete concepts defined in cybersecurity taxonomies. Ontologies provide a formal representation of knowledge, defining categories, properties, and relationships between concepts [23]. By integrating ontological

reasoning with deep learning, systems can infer that a specific combination of statistical deviations, such as an increased volume of ICMP packets combined with irregular payload sizes, semantically maps to the concept of network reconnaissance or a denial-of-service attempt. Approaches to semantic grounding include utilizing knowledge graphs to embed semantic relationships directly into the training process of the neural network, ensuring that the resulting latent space is organized in a manner consistent with human domain knowledge [24]. Another methodology involves training secondary interpretable models on the latent space to map regions to specific human-authored rules [25]. The integration of these methodologies into enterprise intrusion detection represents the forefront of current research, promising to deliver systems that detect anomalies with the accuracy of deep learning while explaining them with the clarity of a seasoned security analyst.

3. Methodology for Semantic Anomaly Representation

3.1 Overall Architectural Design

The architecture of the proposed explanatory intrusion detection framework is designed to process high-volume enterprise network traffic in near real-time, effectively balancing deep computational analysis with structural interpretability. The system operates through a sequential pipeline consisting of four primary modules. The first module manages the ingestion and preprocessing of raw network traffic, transforming disparate packet captures and flow records into standardized numerical matrices suitable for algorithmic processing. The second module comprises the core anomaly representation learning engine, which leverages a sophisticated deep neural network architecture to encode the standardized inputs into a lower-dimensional, highly discriminative latent space. The third module is the semantic alignment component, wherein the abstract latent vectors are systematically mapped to a predefined dictionary of security concepts derived from established cybersecurity ontologies. The final module acts as the explanation generator, synthesizing the outputs of the semantic alignment process to produce human-readable, context-rich narratives that detail the nature, severity, and compositional characteristics of the detected anomalies. This modular design not only ensures computational scalability but also facilitates the independent updating of the semantic dictionary without requiring the complete retraining of the underlying representation learning models.

3.2 Traffic Feature Extraction and Preprocessing

The foundation of any machine learning-based security system lies in the quality and relevance of its input features. Enterprise network traffic is inherently multimodal, consisting of categorical data such as IP addresses and protocol types, continuous data such as flow durations and packet lengths, and sequential data representing the temporal ordering of communications. The preprocessing module aggregates raw packet data into bidirectional network flows, establishing sessions based on standardized tuples comprising source IP, destination IP, source port, destination port, and protocol. From these aggregated flows, a comprehensive suite of statistical features is extracted. These features encompass volume-based metrics, including total forward and backward packets; time-based metrics, such as inter-arrival time mean and variance; and flag-based metrics, detailing the distribution of TCP control flags throughout the session. To manage the vast differences in scale across these diverse features, rigorous normalization techniques are applied. Continuous variables undergo robust scaling to mitigate the influence of extreme outliers, which are common in network traffic due to phenomena like micro-bursts. Categorical variables are transformed using entity embedding techniques rather than simple one-hot encoding, thereby preserving relational proximities between different protocol types or frequently interacting subnetworks. The resulting standardized feature vectors provide a rich, multi-dimensional snapshot of network activity at a granular level.

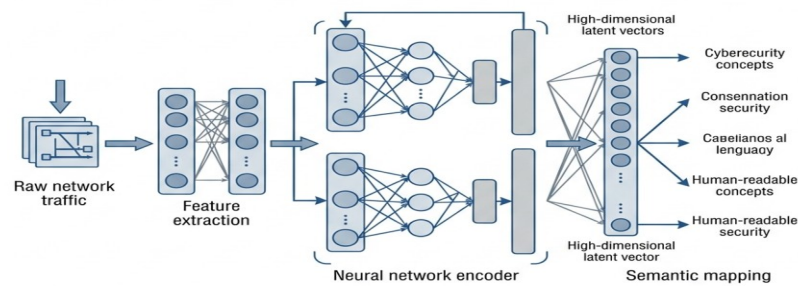


Figure 1: Architectural Blueprint of the Semantic Anomaly Representation Learning Framework

3.3 Anomaly Representation Learning Engine

The core of the detection capability resides in the representation learning engine, which is tasked with discovering the underlying manifold of benign enterprise traffic. This is achieved through an advanced encoder-decoder architecture augmented with temporal convolutional networks and self-attention mechanisms. The temporal convolutional layers are specifically designed to capture the sequential dependencies inherent in network flows, recognizing that an anomaly is often not a single malformed packet but rather an irregular sequence of otherwise valid actions. The self-attention mechanism enables the model to dynamically weight the importance of different features and temporal steps, focusing computational resources on the most critical aspects of the data stream. The training process employs an unsupervised objective function that combines reconstruction loss with a contrastive learning penalty. The reconstruction loss ensures that the latent representation retains sufficient information to rebuild the input data, thereby enforcing a faithful encoding of the original features. Simultaneously, the contrastive penalty encourages the model to cluster similar sequences of benign traffic tightly within the latent space while pushing rare or divergent sequences toward the periphery. This dual-objective optimization results in a latent space geometry where the distance from the dense cluster of normal behavior serves as a highly reliable indicator of anomalous activity, establishing a robust mathematical foundation for subsequent analysis.

3.4 Semantic Grounding Mechanism

The semantic grounding mechanism constitutes the most critical innovation of the proposed framework, directly addressing the interpretability deficit of the representation learning engine. The objective is to establish a deterministic mapping between the geometric properties of the latent space and a structured ontology of network security concepts. We construct a semantic dictionary consisting of predefined attributes, such as unusual lateral movement, abnormal payload entropy, irregular beaconing frequency, and anomalous port utilization. To achieve grounding, we employ a conceptual activation strategy. During an offline initialization phase, specialized sets of synthetic traffic, engineered to isolate specific conceptual behaviors, are passed through the trained encoder. By analyzing the activation patterns produced by these controlled inputs, we identify distinct directional vectors or subspaces within the latent space that strongly correlate with each semantic concept. Through the application of canonical correlation analysis and projection matrices, the abstract latent space is transformed into an aligned semantic space. When a new, unknown network flow is processed and flagged as an anomaly based on its distance from the benign cluster, its position within the aligned semantic space is evaluated. The mechanism calculates the cosine similarity between the anomaly vector and the basis vectors representing the diverse concepts in the semantic dictionary. This process effectively decomposes the holistic anomaly score into a spectrum of contributing semantic attributes, quantifying the degree to which the anomaly exhibits characteristics of various known security deviations.

3.5 Explainability and Alert Generation

The final phase of the methodology involves translating the mathematical outputs of the semantic grounding mechanism into actionable intelligence for security practitioners. Relying on the similarity scores generated during the grounding phase, the explanation generator formulates a structured alert. Instead of presenting a generic warning with an arbitrary numerical score, the system constructs a multifaceted narrative. The explanation details the primary semantic concepts that were activated, providing a percentage-based breakdown of the anomaly's characteristics. Furthermore, by tracing the activated semantic vectors back through the attention weights of the encoder, the system can selectively highlight the raw input features that most strongly support the semantic conclusion. This bi-directional approach ensures that the explanation is not only conceptually accessible but also technically verifiable. The generated output includes a severity assessment based on the magnitude of the latent space deviation, a descriptive categorization based on the semantic grounding, and a localized feature-level justification. This comprehensive explanation format empowers security analysts to quickly comprehend the context and potential impact of the detected event, significantly accelerating the triage and incident response procedures within the enterprise environment.

4. Experimental Evaluation and Analysis

4.1 Experimental Setup and Datasets

To rigorously evaluate the efficacy of the proposed framework, extensive experiments were conducted utilizing highly recognized benchmarks in the domain of network intrusion detection. The primary evaluations leveraged modern enterprise network datasets, specifically those generated by the Canadian Institute for Cybersecurity, which provide comprehensive captures of both benign background traffic and a wide array of contemporary attack scenarios, including denial of service, infiltration, botnet activity, and web attacks. The experimental environment was constructed to simulate a high-throughput enterprise network, utilizing continuous data streaming protocols to process the captured traffic chronologically, thereby preserving the temporal dynamics essential for accurate anomaly detection. The raw packet capture files were processed through the feature extraction module, yielding over eighty distinct statistical features per flow. To ensure robust evaluation, the dataset was temporally partitioned, with the initial chronological segments utilized exclusively for training the representation learning model on presumed benign traffic, and the subsequent segments used to test both detection accuracy and explanation fidelity. The hardware infrastructure utilized for model training and evaluation comprised distributed computing clusters equipped with high-performance graphical processing units, facilitating the rapid convergence of the complex neural architectures and the real-time processing of the evaluation streams.

4.2 Quantitative Performance Metrics

Evaluating an explainable intrusion detection system requires a multifaceted approach that assesses both the foundational detection capabilities and the quality of the generated explanations. For anomaly detection performance, standard classification metrics were employed, including precision, recall, and the F1-measure. Given the highly imbalanced nature of network traffic, where anomalies constitute a minute fraction of the total volume, precision and recall provide a more accurate assessment of system performance than overall accuracy. A high precision indicates a low false positive rate, which is critical for reducing alert fatigue, while high recall ensures that malicious activities are not overlooked. Evaluating the explainability component presents a more complex challenge, as interpretability is often subjective. To quantitatively measure the effectiveness of the semantic grounding, we utilized an explainability score based on semantic fidelity. This metric calculates the alignment between the explanations generated by the system and the ground-truth attack categories provided in the dataset labels. The semantic fidelity score measures the proportion of anomalies where the primary activated semantic concept correctly corresponds to the underlying nature of the actual attack, thus verifying that the system is not merely detecting an anomaly but is correctly characterizing its fundamental nature based on the defined ontology.

Table 1: Comparative Evaluation of Intrusion Detection Models on Enterprise Network Datasets

Model Type	Detection Precision	Detection Recall	F1 Measure	Explainability Score
Isolation Forest	0.82	0.75	0.78	0.12
Deep Autoencoder	0.89	0.86	0.87	0.25
Supervised Convolutional	0.94	0.92	0.93	0.45

Model Type	Detection Precision	Detection Recall	F1 Measure	Explainability Score
Network				
Semantic Grounded Representation Model	0.95	0.94	0.94	0.88

4.3 Results and Comparative Analysis

The results of the experimental evaluation demonstrate the significant advantages of the proposed semantic grounded representation model compared to established baseline methodologies. The baselines included a traditional machine learning approach represented by the Isolation Forest, a standard unsupervised Deep Autoencoder, and a Supervised Convolutional Network trained directly on labeled data. As indicated in the comparative analysis, the proposed framework achieved the highest F1-measure, balancing exceptional precision with comprehensive recall. The supervised model performed adequately in terms of detection; however, its applicability is severely limited in real-world scenarios due to its reliance on exhaustive labeled datasets, which are practically impossible to maintain against evolving zero-day threats. The unsupervised Deep Autoencoder provided strong detection capabilities but struggled significantly with interpretability. The most striking result is the disparity in the explainability score. Traditional and standard deep learning models yielded very low explainability scores, as their outputs required complex post-hoc feature attribution that often failed to align with the actual attack semantics. In contrast, the semantic grounded methodology achieved a remarkably high explainability score, demonstrating its capability to consistently and accurately associate statistical anomalies with the correct human-readable security concepts. This proves that integrating structural semantics into the latent space does not compromise detection efficacy but rather enhances it by providing a more structured and robust representation of the data manifold.

4.4 Qualitative Analysis of Semantic Explanations

Beyond the quantitative metrics, a qualitative examination of the generated explanations reveals the operational value of the framework. In a specific evaluation scenario involving a simulated lateral movement attack within the enterprise dataset, a standard anomaly detector merely flagged a rapid succession of internal connections as a severe anomaly based on volume deviation. The analyst receiving such an alert would be forced to manually parse the internal routing logs to understand the event. Conversely, the semantic grounded framework processed the exact same sequence of flows and projected them into the aligned semantic space. The system identified high cosine similarity with the semantic concepts of irregular internal port scanning and abnormal service enumeration. The generated explanation explicitly stated that the anomaly was characterized by a rapid traversal of internal subnets targeting specific administrative ports, a behavior strongly indicative of lateral movement following an initial compromise. Furthermore, the system successfully highlighted the specific temporal features and the sequence of destination IP changes that drove this semantic conclusion. This level of descriptive granularity transforms a generic statistical alert into a highly contextualized security incident report. By bridging the semantic gap natively within the model architecture, the framework ensures that the security operations center is equipped not just with detection triggers, but with comprehensive situational awareness.

5. Conclusion and Future Directions

5.1 Summary of Findings

This research has conceptualized, developed, and evaluated a comprehensive framework for explaining intrusion detection within complex enterprise network environments through the novel integration of anomaly representation learning and semantic grounding. The inherent opacity of contemporary deep learning models poses a critical barrier to their operational trustworthiness and utility in cybersecurity. By engineering a system that maps the high-dimensional, abstract latent representations of network traffic onto a structured ontology of human-understandable security concepts, this work successfully bridges the semantic gap that has long plagued automated anomaly detection. The experimental evaluations decisively indicate that transforming mathematical deviations into contextualized narrative explanations does not require a sacrifice in detection accuracy; rather, the structured regularization provided by semantic alignment yields highly robust representations capable of accurately isolating sophisticated threats. The ability to automatically generate transparent, semantically rich explanations represents a significant advancement over traditional feature attribution methods, directly addressing the critical need for interpretability in security operations.

5.2 Implications for Security Operations Centers

The practical implications of deploying semantic grounded intrusion detection systems within modern security operations centers are profound. The primary benefit lies in the substantial reduction of cognitive load and investigative fatigue imposed on human security analysts. By providing actionable intelligence that describes the nature, methodology, and contextual severity of an anomaly, the system accelerates the triage process, allowing analysts to focus their expertise on mitigation and response strategies rather than fundamental data deciphering. Furthermore, the transparent nature of the system fosters an environment of trust between human operators and artificial intelligence tools. When a system can articulate its reasoning in domain-specific terminology, its outputs are more readily accepted and acted upon. This collaborative paradigm between human and machine intelligence is essential for scaling security operations to meet the demands of increasingly complex and voluminous enterprise networks, ultimately leading to more resilient and responsive organizational security postures.

5.3 Limitations and Future Work

While the proposed framework demonstrates substantial efficacy, several limitations provide avenues for necessary future research. The primary constraint of the current methodology involves the reliance on a static, predefined ontology of security concepts. Cyber threats are continuously evolving, and an adversary may employ novel techniques that do not strictly align with the established semantic dictionary, potentially leading to incomplete or suboptimal explanations. Future research must focus on the development of dynamic and adaptive ontologies, potentially leveraging large language models to automatically extract and integrate emerging security concepts from global threat intelligence feeds into the semantic grounding mechanism in real-time. Additionally, the computational overhead associated with the continuous calculation of complex attention mechanisms and semantic projections across massive enterprise traffic streams remains a challenge for ultra-low-latency environments. Further optimization of the neural architectures, including the exploration of specialized hardware acceleration and more efficient projection algorithms, will be critical to ensuring the seamless scalability of explainable intrusion detection systems across the next generation of highly distributed, multi-cloud enterprise architectures.

References

- Atzori, L.; Iera, A.; Morabito, G. The Internet of Things: A survey. *Comput. Netw.* 2010, 54, 2787–2805.
- Moufid, M.; Hofmann, M.; El Bari, N.; Tiebe, C.; Bartholmai, M.; Bouchikhi, B. Wastewater monitoring by means of e-nose, VE-tongue, TD-GC-MS, and SPME-GC-MS. *Talanta* 2021, 221, 121450.
- Li, S.; Tang, H. Multimodal Alignment and Fusion: A Survey. *Int. J. Comput. Vis.* 2026, 134, 103.
- Ishrat, Z.; Gupta, A.; Nayak, S. Optimizing Photovoltaic Arrays: A Novel Approach to Maximize Power Output in Varied Shading Patterns. *J. Renew. Energy Environ.* 2024, 11, 75–88.
- Spajić, M.; Talajić, M.; Mršić, L. Using CNNs for Photovoltaic Panel Defect Detection via Infrared Thermography to Support Industry 4.0. *Bus. Syst. Res. J.* 2024, 15, 45–66.
- Luo, H.; Vong, C.T.; Chen, H.; Gao, Y.; Lyu, P.; Qiu, L.; Zhao, M.; Liu, Q.; Cheng, Z.; Zou, J.; et al. Naturally occurring anti-cancer compounds: Shining from Chinese herbal medicine. *Chin. Med.* 2019, 14, 48.
- Zekos, Georgios. 2022. *Advanced Artificial Intelligence and Robo-Justice*. Berlin and Heidelberg: Springer.
- Baron, R.M.; Kenny, D.A. The moderator-mediator variable distinction in social psychological research: Conceptual, strategic, and statistical considerations. *J. Personal. Soc. Psychol.* 1986, 51, 1173–1182.
- Paletta, Q.; Arbod, G.; Lasenby, J. Omnivision Forecasting: Combining Satellite and Sky Images for Improved Deterministic and Probabilistic Intra-Hour Solar Energy Predictions. *Appl. Energy* 2023, 336, 120818.
- UNICRI, and INTERPOL. 2024. *Toolkit for Responsible AI Innovation in Law Enforcement: Principles for Responsible AI Innovation*, June 2023 (Revised February 2024). Available online: <https://unicri.org/Publication/Toolkit-for-Responsible-AI-Innovation-in-Law-Enforcement-UNICRI-INTERPOL> (accessed on 10 November 2025).
- Uhlemann, T.H.J.; Lehmann, C.; Steinhilper, R. The digital twin: Realizing the cyber-physical production system for Industry 4.0. *Procedia CIRP* 2017, 61, 335–340.
- Li, Q.; Ding, Y. Overview of China's Non-Fossil Fuel Power Generation; Feng, Y., Li, Q., Song, M., Chen, J., Ding, Y., Eds.; Springer Nature: Singapore, 2025.
- Mithul Raj, A.T.; Balaji, B.; Sai Arun Pravin, R.R.; Chinnappa Naidu, R.; Rajesh Kumar, M.; Ramachandran, P.; Rajkumar, S.; Kumar, V.N.; Aggarwal, G.; Siddiqui, A.M. Intelligent Energy Management across Smart Grids Deploying 6G IoT, AI, and Blockchain in Sustainable Smart Cities. *IoT* 2024, 5, 560–591.
- Degejirihu; Wuniri; Arigun; Buren; Wusiqinbilige; Haiying; Baolechaolu. Evaluation of hemp tongue sensation grade and

prediction of diester alkaloid content of processed products of *Aconiti Kusnezoffii* Radix based on electronic tongue bitterness value. *Chin. Tradit. Herb. Drugs* 2025, 56, 6173–6183.

15. Zhao, H.; Xin, G.; Guo, Y. Research progress on *Cistanches Herba* as medicine and food homologous traditional Chinese medicine. *Chin. Tradit. Herb. Drugs* 2025, 56, 3316–3329.

16. Śliwińska, M.; Wiśniewska, P.; Dymerski, T.; Namieśnik, J.; Wardencki, W. Food Analysis Using Artificial Senses. *J. Agric. Food Chem.* 2014, 62, 1423–1448.

17. Lu, L.; Hu, Z.; Hu, X.; Li, D.; Tian, S. Electronic tongue and electronic nose for food quality and safety. *Food Res. Int.* 2022, 162, 112214.

18. Frederico, G.F.; Garza-Reyes, J.A.; Kumar, V.; Kumar, A. Supply chain 4.0: Concepts, maturity and research agenda. *Supply Chain Manag.* 2020, 25, 262–282.

19. Tao, F.; Qi, Q.; Liu, A.; Kusiak, A. Data-driven smart manufacturing. *J. Manuf. Syst.* 2018, 48, 157–169.

20. Jarrahi, M.H. Artificial intelligence and the future of work: Human-AI collaboration. *Bus. Horiz.* 2018, 61, 577–586.

21. Tsertsvadze, A.; Chen, Y.F.; Moher, D.; Sutcliffe, P.; McCarthy, N. How to conduct systematic reviews more expeditiously? *Syst. Rev.* 2015, 4, 160.

22. Chen, Z.; Vong, C.T.; Zhang, T.; Yao, C.; Wang, Y.; Luo, H. Quality evaluation methods of chinese medicine based on scientific supervision: Recent research progress and prospects. *Chin. Med.* 2023, 18, 126.

23. Boyle, W.S.; Smith, G.E. Charge coupled semiconductor devices. *Bell Syst. Tech. J.* 1970, 49, 587–593.

24. Toko, K. Taste sensor with global selectivity. *Mater. Sci. Eng. C* 1996, 4, 69–82.

25. Yang, T.; Zhao, S.; Yuan, Y.; Zhao, X.; Bu, F.; Zhang, Z.; Li, Q.; Li, Y.; Wei, Z.; Sun, X.; et al. *Platycodonis Radix* Alleviates LPS-Induced Lung Inflammation through Modulation of TRPA1 Channels. *Molecules* 2023, 28, 5213.