

Cyber physical systems maintenance with explainable unsupervised machine learning

V. Durga Prasad Jasti¹, Koudegai Ashok², Ramarao Gude³, Prabhakar Kandukuri⁴,
Surendra Nadh Benarji Bejjam⁵, Anusha B.⁶

¹Department of Computer Science and Engineering, Siddhartha Academy of Higher Education, Vijayawada, India

²Department of Computer Science and Engineering (Cyber Security), Vignana Bharathi Institute of Technology, Hyderabad, India

³Department of Electronics and Communication Engineering, G. Pullaiah College of Engineering and Technology, Kurnool, India

⁴Department of Artificial Intelligence and Machine Learning, Chaitanya Bharathi Institute of Technology, Hyderabad, India

⁵Department of Computer Science and Engineering, Koneru Lakshmaiah Education Foundation, Vaddeswaram, India

⁶Department of Electronics and Communication Engineering, Andhra University, Visakhapatnam, India

Article Info

Article history:

Received Feb 23, 2024

Revised Oct 16, 2024

Accepted Dec 13, 2025

Keywords:

Cyber-physical systems

Explainable artificial intelligence

Interpretable machine learning

Self-organizing maps

Unsupervised machine learning

ABSTRACT

As cyber-physical systems (CPS) continue to play a pivotal role in modern technological landscapes, the need for robust and transparent machine learning (ML) models becomes imperative. This research paper explores the integration of explainable artificial intelligence (XAI) principles into unsupervised machine learning (UML) techniques for enhancing the interpretability and understanding of complex relationships within CPS. The key focus areas include the application of self-organizing maps (SOMs) as a representative unsupervised learning algorithm and the incorporation of interpretable ML methodologies. The study delves into the challenges posed by the inherently intricate nature of CPS data, characterized by the fusion of physical processes and digital components. Traditional black-box approaches in unsupervised learning often hinder the comprehension of model-generated insights, making them less suitable for critical CPS applications. In response, this research introduces a novel framework that leverages SOMs, a powerful unsupervised technique, while concurrently ensuring interpretability through XAI techniques. The paper provides a comprehensive overview of existing XAI methods and their adaptation to unsupervised learning paradigms. Special emphasis is placed on developing transparent representations of learned patterns within the CPS domain. The proposed approach aims to enhance model interpretability through the generation of human-understandable visualizations and explanations, bridging the gap between advanced ML models and domain experts.

This is an open access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.



Corresponding Author:

Anusha B.

Department of Electronics and Communication Engineering, Andhra University

Visakhapatnam, Andhra Pradesh, India

Email: anutanhar@gmail.com

1. INTRODUCTION

In the realm of cyber-physical systems (CPS), the integration of machine learning (ML) algorithms is becoming increasingly prevalent. However, the black-box nature of many ML models poses challenges in understanding their decisions, which is critical for the safety and reliability of CPS. This paper explores the application of explainable unsupervised machine learning (EUML) techniques, focusing on explainable artificial intelligence (XAI) principles, self-organizing maps (SOMs), interpretable ML, and unsupervised machine learning (UML) [1]–[3]. The rapid integration of ML techniques into CPS has significantly

advanced the capabilities of these complex systems. However, the opacity of many ML models poses challenges in understanding and trusting their decisions, particularly in safety-critical domains [4], [5]. This paper addresses the need for transparency and interpretability in CPS by focusing on EURL techniques. Key components of this exploration include principles from XAI, SOMs, interpretable machine learning, and UML [6], [7].

As we learn more about EURL, it's important to think about what unsupervised learning means in CPS, where physical and cyber parts are linked and require a higher level of openness. This introduction lays the groundwork for a more in-depth look at EURL. The goal is to close the gap between the fact that uncontrolled learning models aren't always clear and the need for clear, understandable decision-making processes in CPS [8]. The subsequent sections of this paper will unfold the layers of supervised and UML, examine the principles of explainability, and ultimately focus on the intersection of unsupervised learning and interpretability in CPS. The investigation will culminate in a detailed exploration of explainable SOMs as a promising approach to address the challenges posed by traditional black-box models in CPS applications [9], [10]. Through this exploration, we aim to contribute to the evolving landscape of interpretable and trustworthy ML solutions for the intricate domain of CPS [11], [12].

Supervised machine learning (SML) stands as a cornerstone in the field of ML, involving the training of models on labeled datasets to make predictions or classifications. The effectiveness of SML in various domains has been well-established, leading to its widespread adoption. However, the interpretability of these models often diminishes as they grow in complexity, making it challenging to comprehend the decision-making processes. In safety-critical applications, understanding why a model makes a specific prediction is paramount, prompting the exploration of alternative approaches [13], [14]. Unlike SML, UML deals with unlabeled data, seeking to uncover underlying patterns and structures without explicit guidance. Clustering and dimensionality reduction are common tasks in UML, where the absence of labeled examples challenges the interpretability of learned representations. As CPS involves intricate interactions between physical and cyber components, the ability to decipher the latent relationships within data becomes crucial for effective decision-making [15]–[17].

Explainable machine learning (XML) has emerged as a critical field to address the black-box nature of many ML models. While feature importance and model-agnostic methods have been successful in enhancing interpretability, these techniques are primarily designed for supervised learning scenarios. As the integration of ML in CPS intensifies, the need for explainability in unsupervised learning becomes apparent, prompting the exploration of EURL techniques [18]–[20]. With this background, you can better understand the problems that standard supervised and UML models cause, especially when it comes to CPS. In the parts that follow, we'll get into the specifics of EURL by looking at its shortcomings, mapping existing XML terms to unsupervised situations, reviewing recent research, and finally focusing on how it can be used to solve the unique problems of CPS. The idea of adding XAI to the AI work flow is shown in illustration 1. The goal is to use methods that can be explained in different stages of the life cycle of AI.

2. EXPLAINABLE UNSUPERVISED MACHINE LEARNING

EURL encompasses a set of critical requirements that distinguish it from traditional unsupervised learning approaches. The primary desiderata include transparency, interpretability, and the ability to provide insights into the decision-making processes of unsupervised models. In the context of CPS, where the consequences of erroneous decisions can be severe, these desiderata become essential for ensuring the trustworthiness and reliability of the deployed ML models [21].

Developing effective EURL algorithms requires adapting and extending XML concepts to unsupervised contexts. Interpretability, openness, and accountability must be rethought for unsupervised learning. The mapping of XML terms to EURL provides a framework for evaluating and improving UML model interpretability [22]. A comprehensive review of the current state-of-the-art in EURL techniques is presented, highlighting advancements, challenges, and potential applications. This literature review provides insights into the progress made in addressing the interpretability issues of unsupervised learning models, laying the groundwork for the subsequent exploration of EURL in the specific context of CPS [23], [24]. This section focuses on the application of EURL within the realm of CPS. It addresses the unique challenges posed by CPS, such as the dynamic interplay between physical and cyber components, the need for real-time decision-making, and the requirement for transparency in complex, interconnected systems. The discussion includes potential use cases, benefits, and considerations for deploying EURL in CPS applications [25].

The exploration of EURL in this section sets the stage for a more detailed examination of a specific approach – SOMs in the subsequent sections. By laying out the desiderata, mapping existing terms, reviewing literature, and contextualizing EURL within CPS, this paper aims to provide a comprehensive understanding of the potential and challenges associated with interpretable UML in complex, dynamic systems.

2.1. Explainable self-organizing maps

SOMs have emerged as a powerful technique in unsupervised learning, particularly in clustering and dimensionality reduction tasks. Introduced by Kohonen, SOMs map high-dimensional input data onto a lower-dimensional grid of neurons, preserving the topological relationships of the input space. While SOMs exhibit remarkable capabilities in capturing complex structures within data, their interpretability has been limited due to the intrinsic complexity of the learned representations [26], [27].

Figure 1 illustrates the structure of a two-dimensional SOM in both the output space and the input space. Figure 1(a) represents the SOM output space, where neurons are arranged on a fixed two-dimensional grid with predefined topological connections that preserve neighborhood relationships. Figure 1(b) shows the same SOM mapped into the input space after training, where the neurons adapt their positions to fit the distribution and clusters of the input data. The learning behavior of the SOM is governed by key hyperparameters, namely the learning rate and the neighborhood radius of the best matching unit (BMU), both of which are gradually reduced at each epoch to ensure smooth convergence and accurate data representation.

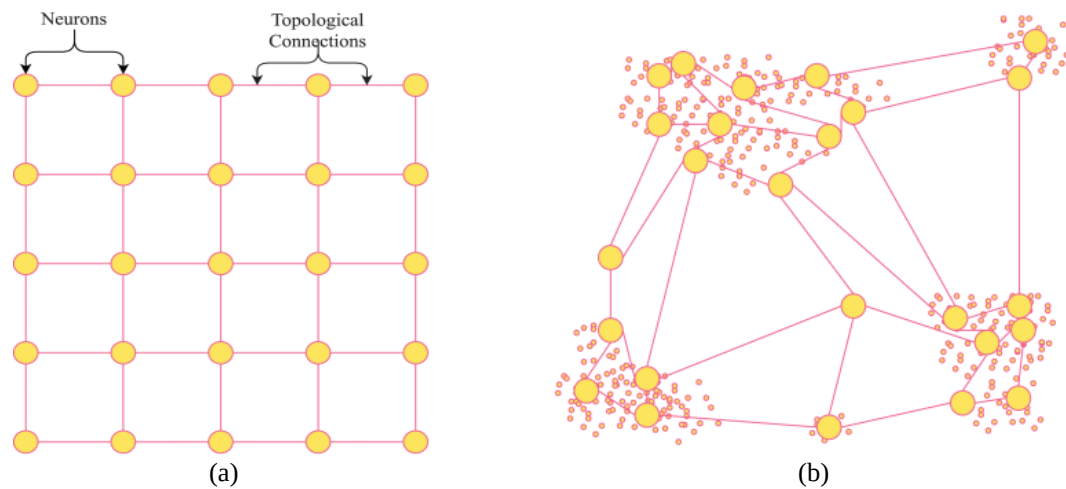


Figure 1. SOMs shown in (a) the output space and (b) the input space changed to fit the 2D spread of the points that were entered

To address the interpretability challenges of traditional SOMs, modifications and extensions have been proposed, giving rise to explainable SOMs. This section delves into the enhancements made to SOMs, focusing on how these modifications render the model more interpretable. Techniques such as neuron importance scoring, feature attribution, and visualization of learned representations are explored to facilitate a deeper understanding of the decision-making processes within the SOMs framework [28]. Explainability in SOMs is important in CPS, where understanding complex dataset correlations is critical. Explainable SOMs promise to provide a robust solution for unsupervised learning in CPS applications by combining SOMs' topological structure capture with explainability's transparency.

The subsequent sections will further explore the experimental setup and results of explainable SOMs, evaluating their model fidelity, local and global interpretability, and usability within CPS. This empirical analysis aims to validate the effectiveness of explainable SOMs in addressing the specific challenges posed by CPS and assess their potential for real-world applications in complex, dynamic environments.

3. EXPERIMENT SETUP AND RESULTS

3.1. Model fidelity

In the experimental setup, the fidelity of explainable SOMs is rigorously assessed. Comparative analyses are conducted against traditional SOMs, evaluating the ability of explainable SOMs to accurately represent the intricate relationships within the given CPS dataset. Metrics such as clustering accuracy, preservation of topological structures, and reconstruction errors are employed to quantify the fidelity of the models.

3.2. Local interpretability

The local interpretability of explainable SOMs is examined to understand how well the model provides insights into individual data points. Neuron importance scoring and feature attribution techniques are applied to identify the key factors influencing the decisions made by the model on a per-instance basis. This analysis aims to highlight the granularity of interpretability achieved by explainable SOMs in the context of CPS.

3.3. Global interpretability

A broader perspective is taken to evaluate the global interpretability of explainable SOMs. By examining the learned representations at a system-wide level, the model's ability to uncover overarching patterns, anomalies, and relationships within the CPS dataset is assessed. Visualization techniques, such as heatmaps and cluster summaries, are employed to facilitate a comprehensive understanding of the global interpretability achieved by explainable SOMs.

3.4. Usability within cyber-physical systems

Real-world experiments are conducted to evaluate the usability of explainable SOMs within CPS. The models are deployed in CPS environments, and their performance is assessed in scenarios that mimic the dynamic, interconnected nature of these systems. This analysis includes considerations for real-time decision-making, adaptability to changing conditions, and the overall impact on system reliability and safety.

The variation in cluster quality matrices utilized in this investigation for the bank marketing data set is illustrated in Figure 2. It illustrates how the number of clusters and SOM dimensions influence the evolution of these matrices. Cluster quality metrics are computed for both the SOM neuron weights (blue) and the training dataset (orange) for a given SOM size. The objective of this analysis is to determine the ideal SOM dimension and cluster count. When the trained SOM neurons accurately represent the entire dataset, it can be expected that cluster analysis of the SOM weights would reveal comparable patterns to that of the entire dataset. Figure 2 illustrates that as the number of clusters increases, they adhere to the same patterns. The Davies-boulding index value and the Silhouette coefficient indicate that three to five clusters are optimal for the examined SOM dimensions (8,16,2,40).

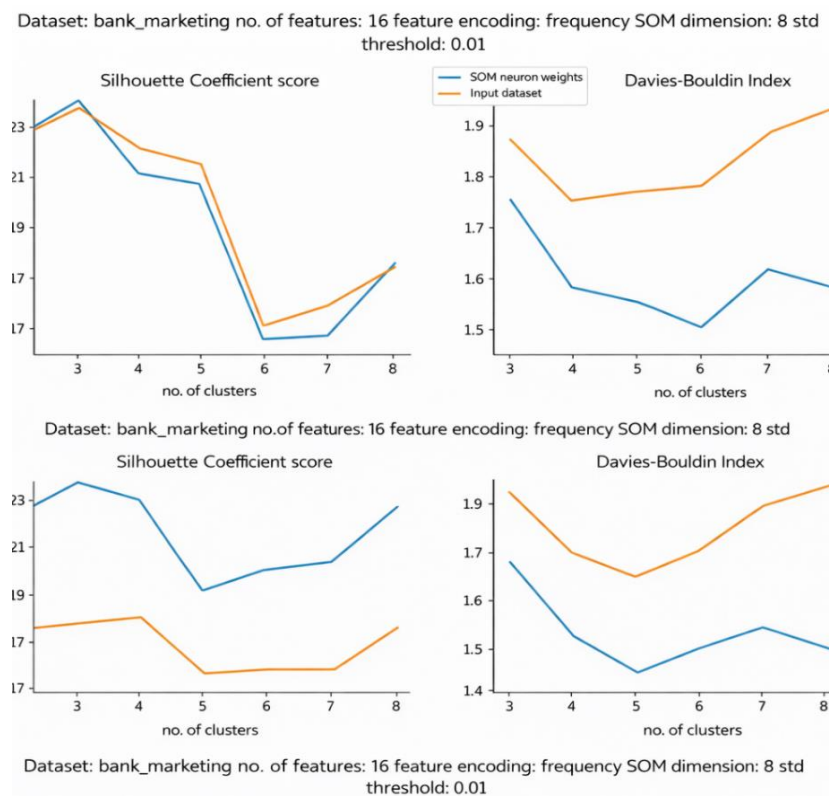


Figure 2. A way to judge the quality of K clusters using the Silhouette coefficient and the Davies-bouldin index for various SOM map sizes

We sorted features by standard deviation and adjusted random or inconsequential features to test our hypothesis. That is, we modified $p\%$ for the most important features (those with the lowest standard deviation values), randomly selected features, and least important features, all given $p\%$ cardinality. We checked each data record in the test set to see if modifying the feature value changed its cluster label in each of the three situations above. Two possibilities were studied for each scenario: i) the proportion of test data records where another cluster may replace the cluster label, and ii) the proportion where all other clusters can be substituted. For some data points, perturbing feature values with close cluster criteria may not be enough to remove them from the initial cluster. We made sure feature values from at least one extra cluster might affect a data point's cluster categorization. Consider a four-cluster scenario with data point j in cluster 2. We attempt changing its cluster designation from 2 to another and replacing its feature values with the averages of clusters 1, 2, and 4. Reducing cardinality $n\%$ and growing swapped percentages are expected. The swap % for all datasets is shown in Figure 3. Except for the second KDD dataset scenario (which asks, "What is the percentage of test data records where all other clusters can swap the cluster label?"), blue bars represent relevant features and brown bars random features. This applies to all datasets. The KDD dataset's extremely unbalanced classes and significant training-testing gap may explain this poor performance. However, KDD works as expected in the first scenario (how many test data records contain a cluster label that might be changed?). These practical results validated our prediction by showing that the suggested SOM technique selected data record cluster labels based on critical criteria.

An additional experiment was conducted to verify the proportion of the chosen K characteristics that were included in the most crucial feature list of a BMU (Algorithm IV). The feature-wise l_1 distance between each data record and its BMU was computed for every data record in the test set. This was followed by an arrangement of the features according to the increasing l_1 distances. We postulated that the most prominent features of a data point would be those that are geographically nearest to it, and that these features would be part of the BMU's prioritised feature lists. After the feature distances are sorted in increasing order, one of three strategies—i) closest, ii) random, or iii) furthest—is used to choose K features. Furthermore, we determined what proportion of K characteristics make it onto the BMU's key feature list. Figure 4 shows the results, with the X-axis showing the total number of features (K) and the Y-axis showing the proportion of those characteristics (%) that were considered relevant (feature list). The colour blue denotes characteristics that are close by, whereas yellow denotes features that are random and green denotes features that are far away. While the yellow bar displays the second-highest percentage, the blue bar shows the highest percentage for all K characteristics. This suggests that each BMU's determined essential feature lists contain the nearby features.

Figure 5 shows the global interpretability of the 'flag' feature in the KDD dataset as it varies between clusters. The SOM neurons were grouped into three groups, and the U-matrix showed the distances and separation between the clusters. Critical feature value ranges were shown against cluster assignments. Additionally, u-maps were employed to verify the dispersion of clusters. The 'flag' characteristic of the KDD dataset is illustrated in Figure 5. The first picture is the raw data, and it depicts the SOM neurons' cluster separation. Across three clusters (component plans), the 'flag' feature's value is depicted in the second image of the first row. The feature value of "flag" varies between the three groups. The three clusters are clearly delineated in the third image of the initial raw data set, where a region of lighter colour signifies the distance between neurons. The greater the area, the more distinct the clusters are. A fine-grained representation of the scale of feature values across clusters is shown in the second row of Figure 5. One thing to keep in mind is that even within the same cluster, there can be neurons with varying feature values for the same feature. If a domain expert wants to know how a specific feature acts inside a cluster, they need this data. Cluster 0 has a larger feature value for the 'flag' feature (0.7-1.0), cluster 1 displays an intermediate range (0.45-0.52), and cluster 3 displays a very low range (0.15). In addition, it displays the likelihood of a particular feature value being present within a cluster. Take cluster 0 as an example; the 90ature value varies from cluster to cluster.

Initial experiment integrity test (Figure 4). We tried several values for p , active features, randomly picked features, and least important features. We calculated the percentage of data points where the cluster label changed after changing p out of all features. We examined two scenarios: the proportion of test data records where another cluster's label could replace the cluster label (left) and the proportion where no cluster's label could replace it (right).

The results obtained from the experimental setup are critically discussed, emphasizing the strengths and potential limitations of explainable SOMs in the context of CPS. Considerations for scalability, computational efficiency, and generalizability are addressed. Additionally, insights into the practical implications of deploying explainable SOMs in real-world CPS applications are discussed, providing a comprehensive understanding of their effectiveness in enhancing model transparency and interpretability. This section shows how explainable SOMs can solve UML problems in CPS. The next section will draw inferences from the data and suggest further study and development in this fast expanding sector.

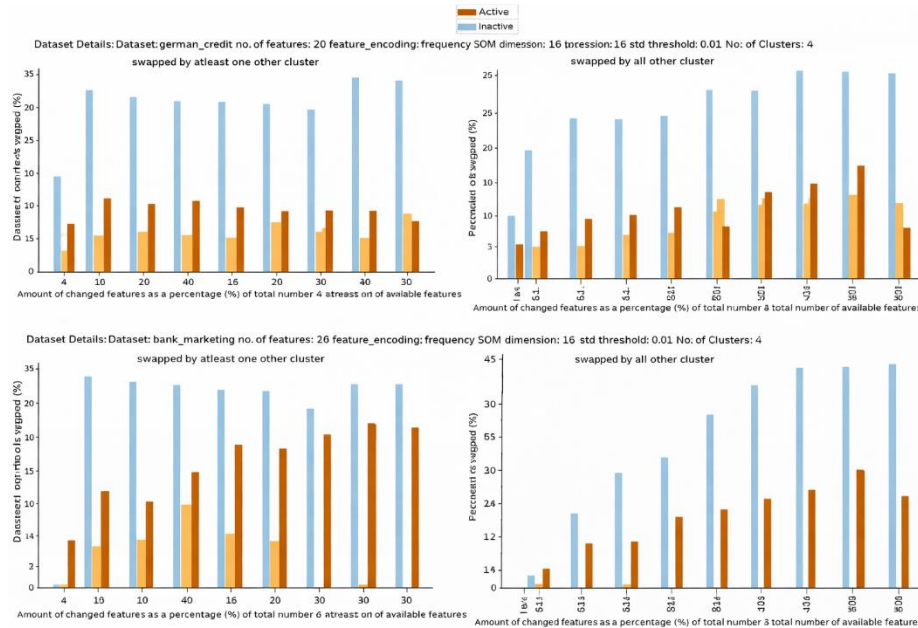


Figure 3. Clustering performance comparison using SOM and input dataset

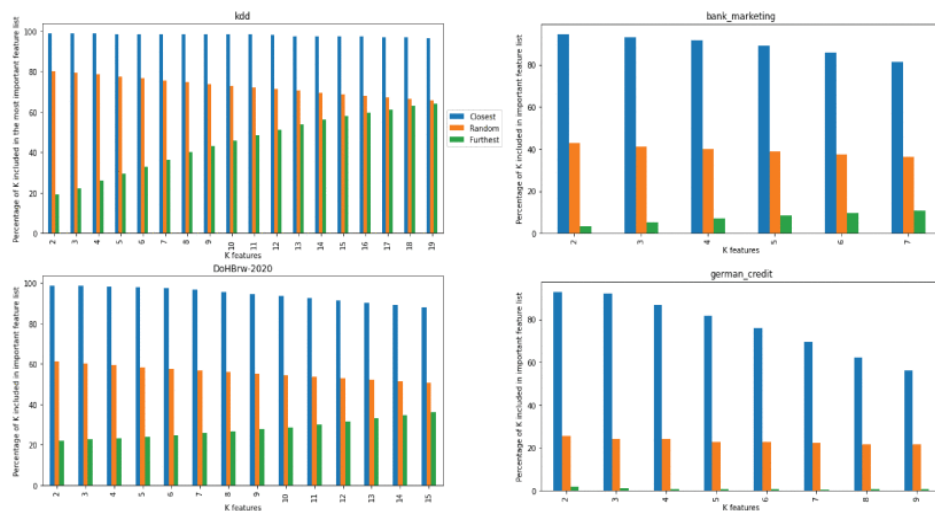


Figure 4. The percentage of nearest K features in the BMU's most significant feature list

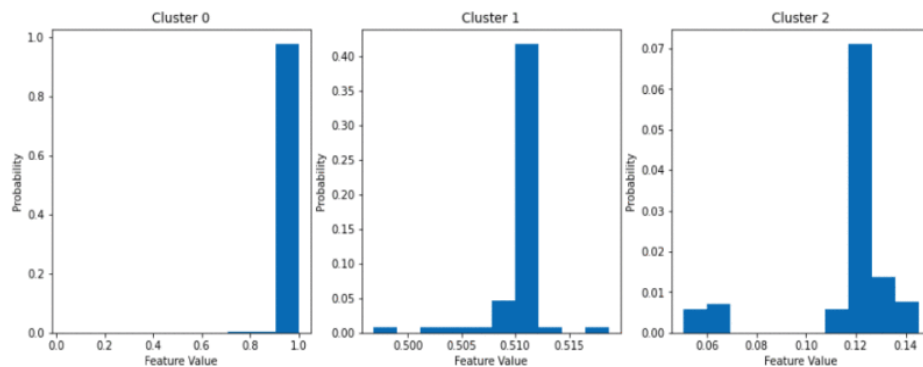


Figure 5. Feature behavior for the 'flag' feature of the KDD data set across clusters (SOM neurons were clustered into three categories; the distances between clusters and the degree of separation between clusters were represented by a U-matrix); the 'flag' feature value varies across clusters

4. CONCLUSION

The exploration of EUML, delved into its desiderata, the adaptation of existing XML terms to unsupervised scenarios, a review of current literature, and the application of EUML principles within the complex landscape of CPS. The introduction of explainable SOMs as a promising EUML technique addressed the need for interpretability in unsupervised learning models. The subsequent section detailed the experiment setup and results, critically examining the model fidelity, local and global interpretability, and the usability of explainable SOMs within CPS environments. By conducting real-world experiments and analyzing performance metrics, this section provided empirical evidence supporting the effectiveness of explainable SOMs in enhancing transparency and interpretability in CPS applications. The discussion brought together theoretical insights and empirical findings, highlighting the strengths, limitations, and practical implications of explainable SOMs in CPS. Considerations for scalability, computational efficiency, and real-time decision-making were addressed, providing a holistic view of the potential impact of EUML on the field. As we look to the future, continued research in EUML, especially within the context of CPS, holds great promise. Advancements in interpretable unsupervised learning techniques can contribute significantly to the ongoing development of safe, reliable, and transparent ML solutions for complex, dynamic systems. This paper serves as a stepping stone, encouraging further exploration and innovation in the intersection of EUML and CPS.

ACKNOWLEDGMENTS

The authors would like to express their sincere gratitude to their respective institutions for providing the necessary facilities, infrastructure, and academic support to carry out this research work. The authors also thank colleagues and reviewers for their valuable suggestions and constructive feedback, which helped improve the quality of this paper.

FUNDING INFORMATION

This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

AUTHOR CONTRIBUTIONS STATEMENT

This journal uses the Contributor Roles Taxonomy (CRediT) to recognize individual author contributions, reduce authorship disputes, and facilitate collaboration.

Name of Author	C	M	So	Va	Fo	I	R	D	O	E	Vi	Su	P	Fu
V. Durga Prasad Jasti	✓	✓	✓	✓	✓	✓		✓	✓	✓			✓	
Koudegai Ashok		✓				✓		✓	✓	✓	✓	✓		
Ramarao Gude	✓		✓	✓			✓			✓	✓		✓	✓
Prabhakar Kandukuri	✓	✓	✓	✓	✓	✓		✓	✓	✓			✓	
Bejjam S. N. Benarji		✓				✓		✓	✓	✓	✓	✓		
Anusha B.	✓		✓	✓			✓		✓	✓	✓		✓	✓

C : **C**onceptualization

M : **M**ethodology

So : **S**oftware

Va : **V**alidation

Fo : **F**ormal analysis

I : **I**nvestigation

R : **R**esources

D : **D**ata Curation

O : Writing - **O**riginal Draft

E : Writing - Review & **E**ditng

Vi : **V**isualization

Su : **S**upervision

P : **P**roject administration

Fu : **F**unding acquisition

CONFLICT OF INTEREST STATEMENT

The authors declare that there is no conflict of interest regarding the publication of this paper.

DATA AVAILABILITY

The data used to support the findings of this study are available from the corresponding author upon reasonable request.

REFERENCES

- [1] M. I. Malik, A. Ibrahim, P. Hannay, and L. F. Sikos, "Developing resilient cyber-physical systems: a review of state-of-the-art malware detection approaches, gaps, and future directions," *Computers*, vol. 12, no. 4, p. 79, Apr. 2023, doi: 10.3390/computers12040079.

- [2] R. Jabeen, Y. Singh, and Z. A. Sheikh, "Machine learning for security of cyber-physical systems and security of machine learning: attacks, defences, and current approaches," in *The International Conference on Recent Innovations in Computing*, 2022, pp. 813–841, doi: 10.1007/978-981-99-0601-7_62.
- [3] F. S. Mozaffari, H. Karimipour, and R. M. Parizi, "Learning based anomaly detection in critical cyber-physical systems," in *Security of Cyber-Physical Systems*, Cham: Springer International Publishing, 2020, pp. 107–130.
- [4] M. K. Hasan, R. A. Abdulkadir, S. Islam, T. R. Gadekallu, and N. Safie, "A review on machine learning techniques for secured cyber-physical systems in smart grid networks," *Energy Reports*, vol. 11, pp. 1268–1290, 2024, doi: 10.1016/j.egy.2023.12.040.
- [5] R. Salih Ahmed, E. Sayed Ali Ahmed, and R. A. Saeed, "Machine learning in cyber-physical systems in industry 4.0," in *Artificial intelligence paradigms for smart cyber-physical systems*, 2021, pp. 20–41.
- [6] C. S. Wickramasinghe, K. Amarasinghe, D. L. Marino, C. Rieger, and M. Manic, "Explainable unsupervised machine learning for cyber-physical systems," *IEEE Access*, vol. 9, pp. 131824–131843, 2021, doi: 10.1109/ACCESS.2021.3112397.
- [7] V. Božić, "Explainable artificial intelligence (XAI): enhancing transparency and trust in AI systems," *Preprint*, 2023.
- [8] A. Giannaros *et al.*, "Autonomous vehicles: sophisticated attacks, safety issues, challenges, open topics, blockchain, and future directions," *Journal of Cybersecurity and Privacy*, vol. 3, no. 3, pp. 493–543, 2023, doi: 10.3390/jcp3030025.
- [9] N. Burkart and M. F. Huber, "A survey on the explainability of supervised machine learning," *Journal of Artificial Intelligence Research*, vol. 70, pp. 245–317, 2021.
- [10] Mamta, N. Garla, I. U. Haq, and H. Dhiman, "Revolutionizing Biomedical engineering with quantum computing and AI," in *Quantum Innovations at the Nexus of Biomedical Intelligence*, 2023, pp. 206–222.
- [11] S. Ali *et al.*, "Explainable artificial intelligence (XAI): what we know and what is left to attain trustworthy artificial intelligence," *Information Fusion*, vol. 99, 2023, doi: 10.1016/j.inffus.2023.101805.
- [12] V. F. Santos, C. Albuquerque, D. Passos, S. E. Quincozes, and D. Mossé, "Assessing machine learning techniques for intrusion detection in cyber-physical systems," *Energies*, vol. 16, no. 16, p. 6058, Aug. 2023, doi: 10.3390/en16166058.
- [13] M. M. Taye, "Understanding of machine learning with deep learning: architectures, workflow, applications and future directions," *Computers*, vol. 12, no. 5, p. 91, Apr. 2023, doi: 10.3390/computers12050091.
- [14] I. H. Sarker, "Machine learning: algorithms, real-world applications and research directions," *SN Computer Science*, vol. 2, no. 3, p. 160, May 2021, doi: 10.1007/s42979-021-00592-x.
- [15] S. Derkarabetian, S. Castillo, P. K. Koo, S. Ovchinnikov, and M. Hedin, "A demonstration of unsupervised machine learning in species delimitation," *Molecular Phylogenetics and Evolution*, vol. 139, p. 106562, Oct. 2019, doi: 10.1016/j.ympev.2019.106562.
- [16] K. Karadžuzović-Hadžiabdić and A. Peters, "Artificial intelligence in clinical decision-making for diagnosis of cardiovascular disease using epigenetics mechanisms," in *Epigenetics in Cardiovascular Disease*, Elsevier, 2021, pp. 327–345.
- [17] K. A. Abbas *et al.*, "Unsupervised machine learning technique for classifying production zones in unconventional reservoirs," *International Journal of Intelligent Networks*, vol. 4, pp. 29–37, 2023, doi: 10.1016/j.ijin.2022.11.007.
- [18] P. Linardatos, V. Papastefanopoulos, and S. Kotsiantis, "Explainable AI: a review of machine learning interpretability methods," *Entropy*, vol. 23, no. 1, pp. 1–45, 2021, doi: 10.3390/e23010018.
- [19] R. Zhou and T. Hu, "Evolutionary approaches to explainable machine learning," in *Handbook of evolutionary machine learning*, 2024, pp. 487–506.
- [20] K. Rasheed, A. Qayyum, M. Ghaly, A. Al-Fuqaha, A. Razi, and J. Qadir, "Explainable, trustworthy, and ethical machine learning for healthcare: a survey," *Computers in Biology and Medicine*, vol. 149, p. 106043, Oct. 2022, doi: 10.1016/j.compbiomed.2022.106043.
- [21] P. Thisovithan, H. Aththanayake, D. P. P. Meddage, I. U. Ekanayake, and U. Rathnayake, "A novel explainable AI-based approach to estimate the natural period of vibration of masonry infill reinforced concrete frame structures using different machine learning techniques," *Results in Engineering*, vol. 19, p. 101388, Sep. 2023, doi: 10.1016/j.rineng.2023.101388.
- [22] I. Horváth and J. Tavčar, "Designing cyber-physical systems for runtime self-adaptation: knowing more about what we miss...," *Journal of Integrated Design and Process Science*, vol. 25, no. 2, pp. 1–26, 2021, doi: 10.3233/JID210030.
- [23] V. Belle and I. Papantonis, "Principles and practice of explainable machine learning," *Frontiers in Big Data*, vol. 4, Jul. 2021, doi: 10.3389/fdata.2021.688969.
- [24] M. Elahi, S. O. Afolaranmi, J. L. M. Lastra, and J. A. P. Garcia, "A comprehensive literature review of the applications of AI techniques through the lifecycle of industrial equipment," *Discover Artificial Intelligence*, vol. 3, no. 1, p. 43, Dec. 2023, doi: 10.1007/s44163-023-00089-x.
- [25] P. Bindra, M. Kshirsagar, C. Ryan, G. Vaidya, K. K. Gupt, and V. Kshirsagar, "Insights into the advancements of artificial intelligence and machine learning, the present state of art, and future prospects: seven decades of digital revolution," *Smart Innovation, Systems and Technologies*, vol. 225, pp. 609–621, 2021, doi: 10.1007/978-981-16-0878-0_59.
- [26] B. A. Yilma, "Personalisation in cyber-physical-social systems," Université de Lorraine, 2021.
- [27] D. Miljkovic, "Brief review of self-organizing maps," in *2017 40th International Convention on Information and Communication Technology, Electronics and Microelectronics (MIPRO)*, May 2017, pp. 1061–1066, doi: 10.23919/MIPRO.2017.7973581.
- [28] A. Jamil, A. A. Hameed, and Z. Orman, "A faster dynamic convergence approach for self-organizing maps," *Complex and Intelligent Systems*, vol. 9, no. 1, pp. 677–696, 2023, doi: 10.1007/s40747-022-00826-2.

BIOGRAPHIES OF AUTHORS



V. Durga Prasad Jasti    received his M.Tech (CSE) from Acharya Nagarjuna University, Guntur. He is currently pursuing Ph.D. from Acharya Nagarjuna University, Guntur and working as an Assistant Professor in Siddhartha Academy of Higher Education, deemed to be University, in the Department of Computer Science and Engineering, Vijayawada, Andhra Pradesh. His research interests include deep learning and image processing. He can be contacted at email: prasadjasti2018@gmail.com.



Koudegai Ashok    is currently working as associate professor in Vignana Bharathi Institute of Technology affiliated to JNTUH, Hyderabad. He has 21 years of teaching experience in engineering education. He received his B.Tech in 1999 from RGM CET, Nandyala and M.Tech in 2006 from JNTUA Anantapur and currently pursuing Ph.D. from K L University, Vijayawada. His research interests include cyber security, computer networks, network security, and operating systems. He can be contacted at email: koudegai.ashok@vbithyd.ac.in.



Ramarao Gude    is currently associate professor of Electronics and Communication Engineering of G. Pullaiah College of Engineering and Technology, Kurnool, he has about 18 years of experience in teaching and research. He has published several research papers in journals of both international and national repute. He holds M.Tech degree from JNTUCEA, Ananthapuramu in the field of electronics and communication engineering. He can be contacted at email: ramaraog19@gmail.com.



Prabhakar Kandukuri    received his Ph.D. in computer science and engineering from JNT University Ananthapur, India. M. Tech. Degree in computer science and engineering from JNTUA College of Engineering, B. Tech. Degree in computer science and engineering from Acharya Nagarjuna University, and he received his Diploma from the SBTET, Hyderabad. He is a professor of Artificial Intelligence and Machine Learning Department, Chaitanya Bharathi Institute of Technology, Hyderabad, India. His main research interest includes software engineering, machine learning. He can be contacted at email: prabhakarcs@gmail.com.



Mr. Surendra Nadh Benarji Bejjam    is currently working as assistant professor in KL University affiliated to Koneru Lakshmaiah Education Foundation, Vaddeswaram, Guntur. He has 8-years of teaching experience in engineering education. He received B.Tech in 2012 from Acharya Nagarjuna University (KLCE), M.Tech in 2015 from Pondicherry Central University and part-time research scholar (Ph.D.) from Pondicherry Central University. His Research interests include cyber security, machine learning, deep learning, networks, artificial intelligence, and cloud security. He can be contacted at email: benarji@kluniversity.in.



Anusha B.    is currently in the Department of Electronics and Communication Engineering, Andhra University, Andhra Pradesh, India. Research interests are artificial intelligence in ECE, embedded systems, nanoelectronics, internet of things (IoT) security, robotics and automation. She can be contacted at email: anutanhar@gmail.com.