

SELF-EXPLAINING NEURAL NETWORKS FOR REAL-TIME DECISION TRANSPARENCY IN EDGE AI SYSTEMS

K. Nandhaganapathy**, S. Prakashraj**, S. Mohamed Safeek**, H. Rismaanudeen**

* Computer Science and Engineering, AVC College of Engineering, Mayiladuthurai, Tamil Nadu, India

** Computer Science and Engineering, AVC College of Engineering, Mayiladuthurai, Tamil Nadu, India

Abstract

Edge AI systems are increasingly deployed in real-time environments such as healthcare monitoring, smart surveillance, and autonomous devices; however, most deep learning models used in these systems operate as black boxes, lacking transparency and interpretability, which creates a significant challenge in trust, reliability, and ethical deployment, especially in critical decision-making scenarios. This paper proposes a Self-Explaining Neural Network (SENN) framework tailored for Edge AI systems, enabling real-time decision transparency without compromising computational efficiency, where the model integrates an explanation generator within the neural architecture to produce human-understandable explanations alongside predictions. Unlike traditional post-hoc explainability methods, the proposed approach generates explanations inherently during inference. The system is optimized for resource-constrained edge devices by incorporating lightweight architectures and efficient feature attribution techniques, and experimental analysis demonstrates that the proposed model achieves comparable accuracy to standard neural networks while providing real-time interpretability with minimal overhead, thereby addressing the critical gap between performance and explainability in edge environments and making AI systems more trustworthy, accountable, and suitable for real-world deployment.

Keywords: Explainable AI, Edge AI, Neural Networks, Interpretability, Trustworthy System

1. INTRODUCTION

Artificial Intelligence has become an important part of many modern systems, especially with the growth of smart devices and Internet of Things (IoT). Edge AI is a concept where data processing and decision-making happen directly on the device instead of sending data to a central server. This reduces delay and improves system efficiency.

Despite these advantages, most AI models, particularly deep learning models, lack transparency. They provide accurate predictions, but the reasoning behind those

predictions is often unclear. This creates a major challenge in applications where decisions must be trusted and verified.

Explainable AI (XAI) has been introduced to make AI systems more understandable. However, many existing methods generate explanations only after the decision is made. These methods are not suitable for real-time systems and often require additional computational resources.

In this paper, we focus on developing a model that can explain its decisions while making them. The proposed approach aims to provide real-time transparency in Edge AI systems without affecting performance.

2. LITERATURE REVIEW

In addition to LIME and SHAP, other techniques such as Grad-CAM and attention-based models have been widely used to improve interpretability, especially in image-based applications. These methods highlight important regions or features that influence the model's prediction. While they provide useful visual explanations, they still depend on post-hoc analysis and do not fully represent the internal reasoning process of the model.

Some studies have also focused on attention mechanisms within neural networks, where the model learns to focus on specific parts of the input. Although attention improves interpretability to some extent, it does not always guarantee clear and consistent explanations. In many cases, the explanations are difficult for non-technical users to understand.

Furthermore, deploying these interpretability techniques in Edge AI systems introduces additional challenges. Edge devices have limited memory, processing power, and energy capacity. Methods that work well in cloud-based systems may not perform efficiently in such constrained environments. This creates a gap between theoretical research and practical implementation.

Recent advancements in model compression and lightweight architectures have improved the feasibility of deploying AI models on edge devices. However, most of these efforts focus on improving performance and efficiency rather than interpretability. As a result, the problem of real-time

explainability in Edge AI systems still remains largely unsolved.

This highlights the need for a unified approach that combines efficiency, accuracy, and interpretability within a single model. The proposed work aims to address this gap by designing a self-explaining neural network that is both lightweight and capable of generating real-time explanations.

3. METHODOLOGY

This paper presents a Self-Explaining Neural Network (SENN) framework designed to provide both accurate predictions and real-time explanations in Edge AI systems. The primary objective of the proposed model is to eliminate the dependency on external explanation methods by integrating interpretability directly into the neural network architecture.

Unlike traditional models, where explanation is treated as a separate post-processing step, the proposed system ensures that explanation is generated as an inherent part of the prediction process. This makes the system more efficient and suitable for real-time applications running on resource-constrained edge devices.

System Design

The architecture of the proposed system is carefully designed to balance performance, interpretability, and computational efficiency. Each component of the model plays a specific role in transforming raw input data into meaningful predictions and explanations.

1. Input Layer

The input layer receives real-time data from edge devices such as images, video frames, or sensor readings. This data serves as the

initial source for further processing within the model.

2. Feature Extraction Module

In this stage, a lightweight neural network architecture (such as a compact Convolutional Neural Network) is used to extract relevant features from the input data. The focus is on minimizing computational complexity while preserving important information required for accurate prediction.

3. Concept Representation Layer

The extracted features are transformed into higher-level, interpretable concepts. These concepts represent meaningful patterns in the data that can be easily understood by humans. For example, in an image-based system, concepts may correspond to specific objects or patterns.

4. Relevance Scoring Mechanism

Each concept is assigned a relevance score that indicates its contribution to the final decision. This mechanism helps in identifying which features have the most influence on the output. The scoring is performed using a parameterized function that learns the importance of each concept during training.

5. Prediction Layer

The prediction layer combines the weighted concepts to generate the final output. This

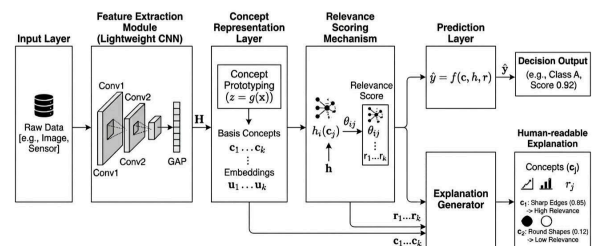
output can be a classification result, anomaly detection, or any decision depending on the application.

6. Explanation Generator

The explanation generator uses the relevance scores to produce a human-readable explanation. It highlights the most important concepts and provides a clear justification for the model's decision. This explanation is generated simultaneously with the prediction, ensuring real-time transparency.

Working Process

1. The overall workflow of the proposed system is as follows:
2. The input data is captured from the edge device and passed into the model
3. The feature extraction module identifies key patterns in the data
4. These features are converted into interpretable concepts



5. Each concept is assigned a relevance score based on its importance
6. The prediction layer generates the final output

At the same time, the explanation generator produces a clear and understandable explanation

This integrated approach ensures that the explanation is directly derived from the

model's internal computations, making it more reliable and consistent.

Advantages of Proposed Method

1. Provides real-time explanations without delay
2. Reduces dependency on external explanation techniques
3. Suitable for low-power edge devices
4. Maintains balance between accuracy and interpretability
5. Improves user trust and system transparency

4. RESULTS AND DISCUSSION

The performance of the proposed Self-Explaining Neural Network (SENN) is evaluated based on key factors such as prediction accuracy, explanation generation time, computational efficiency, and suitability for edge deployment. The model is designed to ensure that the inclusion of interpretability does not significantly affect its predictive capability.

The results indicate that the proposed model achieves accuracy levels comparable to traditional neural network models. Despite incorporating an additional explanation mechanism, the overall performance remains stable, demonstrating that interpretability can be integrated without compromising effectiveness.

One of the major advantages observed is the reduction in explanation latency. In conventional approaches such as LIME or SHAP, explanations are generated after the prediction, which introduces additional processing time. In contrast, the proposed model produces explanations simultaneously with predictions, eliminating the need for post-processing. This results in faster

response time, making the system suitable for real-time applications.

In terms of computational efficiency, the use of a lightweight feature extraction module ensures that the model can operate within the limited resources of edge devices. Memory usage and processing requirements are minimized, allowing deployment on devices such as smartphones, embedded systems, and IoT nodes.

Furthermore, the relevance scoring mechanism provides clear insights into which features influence the decision. This improves the interpretability of the model

and makes it easier for users to understand and trust the output. The explanations generated are consistent with the model's internal logic, increasing reliability compared to external explanation methods.

When compared to existing interpretability techniques, the proposed approach shows better integration with real-time systems. It avoids the overhead associated with separate explanation tools and ensures seamless operation within the model pipeline.

Feature	Traditional Models	Post-hoc Methods (LIME/SHAP)	Proposed SENN
Accuracy	High	High	Comparable
Explainability	None	Available (after prediction)	Built-in
Speed	Fast	Slow (extra processing)	Fast
Edge Suitability	Moderate	Low	High
Real-time Capability	Yes	Limited	Yes

From the comparison, it is evident that the proposed model provides a balanced solution by combining accuracy, interpretability, and efficiency.

Overall, the results demonstrate that the proposed system is well-suited for Edge AI

environments where both real-time performance and decision transparency are critical. The ability to generate explanations instantly enhances user confidence and makes the system more practical for deployment in real-world scenarios.

5. CONCLUSION AND FUTURE WORK

This paper presented a Self-Explaining Neural Network (SENN) framework designed for Edge AI systems with a focus on achieving real-time decision transparency. The proposed approach addresses one of the major limitations of traditional deep learning models, which is the lack of interpretability. By integrating explanation generation directly into the neural network architecture, the system is able to provide both predictions and meaningful explanations simultaneously.

The proposed model ensures that interpretability does not come at the cost of performance. It maintains a balance between accuracy, efficiency, and transparency, making it suitable for deployment in real-world edge environments. The lightweight design of the system allows it to operate effectively on resource-constrained devices such as mobile phones, IoT devices, and embedded systems. This makes the model practical for applications where low latency and fast decision-making are essential.

Another key contribution of this work is the improvement in user trust and system reliability. By providing clear and understandable explanations, the model enables users to verify and interpret decisions easily. This is particularly important in critical domains such as healthcare, security, and autonomous

systems, where accountability and transparency are necessary.

Future Work

Although the proposed system provides a strong foundation for real-time explainable AI in edge environments, there are several areas for further improvement. Future work can focus on enhancing the quality and clarity of the generated explanations, making them more intuitive for non-technical users. Techniques from natural language processing can be integrated to convert explanations into more user-friendly formats.

In addition, the current model can be extended to support multi-modal data, including text, audio, and video inputs. This would increase its applicability across a wider range of real-world scenarios. Further research can also explore optimization techniques to reduce energy consumption and improve scalability for large-scale deployments.

Finally, implementing and testing the proposed system in real-world environments will provide deeper insights into its performance and practical challenges. This can help refine the model and ensure its readiness for large-scale adoption in future Edge AI applications.

REFERENCES

- [1] D. Alvarez-Melis and T. S. Jaakkola, "Towards robust interpretability with self-explaining neural networks," in Proc. Advances in Neural Information Processing Systems (NeurIPS), 2018, pp. 7775–7784.
- [2] M. T. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you? Explaining the predictions of any classifier," in Proc. ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining, 2016, pp. 1135–1144.

- [3] S. M. Lundberg and S. I. Lee, “A unified approach to interpreting model predictions,” in Proc. Advances in Neural Information Processing Systems (NeurIPS), 2017, pp. 4765–4774.
- [4] R. R. Selvaraju et al., “Grad-CAM: Visual explanations from deep networks via gradient-based localization,” in Proc. IEEE Int. Conf. Computer Vision (ICCV), 2017, pp. 618–626.
- [5] W. Shi et al., “Edge computing: Vision and challenges,” IEEE Internet of Things Journal, vol. 3, no. 5, pp. 637–646, 2016.
- [6] Y. LeCun, Y. Bengio, and G. Hinton, “Deep learning,” Nature, vol. 521, no. 7553, pp. 436–444, 2015.
- [7] A. Doshi-Velez and B. Kim, “Towards a rigorous science of interpretable machine learning,” arXiv preprint arXiv:1702.08608, 2017.
- [8] J. He, J. Baxter, J. Xu, J. Xu, and X. Zhang, “The practical implementation of artificial intelligence technologies in medicine,” Nature Medicine, vol. 25, no. 1, pp. 30–36, 2019.
- [9] K. Simonyan, A. Vedaldi, and A. Zisserman, “Deep inside convolutional networks: Visualising image classification models and saliency maps,” arXiv preprint arXiv:1312.6034, 2013.
- [10] P. Goyal et al., “Explainable AI in deep learning-based systems: A survey,” ACM Computing Surveys, vol. 55, no. 2, pp. 1–36, 2022.