

Paper Type: Original Article

Automated Post-Earthquake Damage Assessment in Reinforced Concrete Structures Using U-Net-Based Convolutional Neural Networks

Shirin Mirabedini^{1,*} , Seyedeh Fatemeh Nourani¹ 

¹ Department of Computer Engineering, Payame Noor University, PO Box 19395-3697 Tehran, I.R of Iran;
sh_mirabedini@pnu.ac.ir; sf.noorani@pnu.ac.ir.

Citation:

Received: 25 December 2025

Revised: 10 March 2025

Accepted: 05 May 2026

Mirabedini, Sh., & Nourani, S. F. (2026). Automated post-earthquake damage assessment in reinforced concrete structures using U-net-based convolutional neural networks. *International journal of researches on civil engineering with artificial intelligence*, 3(2), 144-154.

Abstract

Rapid and accurate assessment of damage to structures following an earthquake is critical for disaster management, emergency response, and rehabilitation planning. Traditional visual inspection methods are time-consuming, costly, and inherently subjective. This research presents an automated approach based on deep learning for identifying and quantifying cracks in Reinforced Concrete (RC) structures after seismic events. The proposed model utilizes a U-Net architecture, a type of Convolutional Neural Network (CNN), for pixel-level crack segmentation. It is trained and evaluated on a dataset of real-world images from the recent Kahramanmaraş earthquake (Türkiye, 2023). Results demonstrate that the model effectively identifies various crack patterns with high accuracy, achieving a strong Intersection over Union (IoU) score of 0.745. This approach offers an efficient, objective, and scalable tool for structural engineers to assess integrity and prioritize repairs. The integration of an Efficient Channel Attention (ECA) mechanism further enhances feature extraction, leading to improved segmentation performance, particularly for subtle and narrow cracks.

Keywords: Deep learning, Convolutional neural networks, U-net, Crack detection, Post-earthquake damage assessment, Reinforced concrete structures, Structural health monitoring, Image segmentation.

1 | Introduction

Earthquakes pose a significant threat to Reinforced Concrete (RC) structures worldwide, often causing severe and widespread damage. The 2023 Kahramanmaraş earthquakes in Türkiye, with magnitudes of 7.8 and 7.5,

 Corresponding Author: sh_mirabedini@pnu.ac.ir

 <https://doi.org/10.48314/ijrceai.v3i2.58>



Licensee System Analytics. This article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0>).

served as a devastating reminder of the urgent need for rapid, accurate, and scalable post-event structural assessments to ensure safety, facilitate timely recovery, and inform reconstruction efforts. Historical seismic events, including Northridge (1994, USA), Kobe (1995, Japan), Kocaeli (1999, Türkiye), Sichuan (2008, China), and Abruzzo (2009, Italy), have consistently demonstrated the immense challenge of assessing damage across hundreds or thousands of structures within a limited timeframe. The sheer scale of destruction often overwhelms available expert resources, highlighting the critical gap between the need for rapid assessment and the capacity of traditional methods [1], [2].

1.1 | Limitations of Traditional Methods

Traditional post-earthquake damage detection relies heavily on visual inspections conducted by trained structural engineers. These inspections are guided by established protocols such as ATC-20, FEMA-273, and NZSEE, which categorize damage based on observed distress indicators like crack width, concrete spalling, rebar buckling, and overall structural integrity. While these frameworks provide a standardized approach, they are fundamentally limited by several factors:

- I. Labor-intensive and time-consuming: Deploying teams of experts to inspect thousands of structures is a logistical nightmare, requiring significant time and resources that are often unavailable in the immediate aftermath of a disaster.
- II. Subjective and inconsistent: Assessment results are inherently dependent on the inspector's experience, judgment, and even physical condition, leading to significant variability and potential for error.
- III. Limited accessibility: Many damaged areas, particularly in high-rise buildings or critical infrastructure like bridges and dams, may be physically inaccessible, preventing thorough inspection.
- IV. Scalability issues: Manual inspection methods simply cannot scale to meet the demands of a major earthquake affecting a large urban area.

Conventional computer vision techniques, such as edge detection (e.g., Canny, Sobel), thresholding, and morphological operations, have also been applied to automate crack detection. However, these methods are highly sensitive to environmental conditions, including lighting, shadows, and background textures. They struggle to differentiate cracks from surface imperfections, stains, or construction joints, leading to high rates of false positives and negatives in real-world scenarios. These limitations underscore the pressing need for more robust, automated, and intelligent solutions [3].

1.2 | Advancements in Deep Learning

Deep learning, particularly Convolutional Neural Networks (CNNs), has revolutionized the field of computer vision and has recently been adopted for structural damage assessment. CNNs are capable of automatically learning hierarchical features from raw image data, eliminating the need for manual feature engineering. This has led to significant breakthroughs in various civil engineering applications, including crack detection, pavement distress assessment, and bridge inspection [4–6].

Research has consistently shown that CNN-based methods significantly outperform traditional image processing techniques in terms of accuracy, robustness, and generalization. Various CNN architectures, such as U-Net, DeepLabV3+, and Fully Convolutional Networks (FCN), have been explored for pixel-level crack segmentation. Among these, the U-Net architecture, originally developed for biomedical image segmentation, has gained particular popularity due to its ability to produce precise segmentation maps by combining high-level semantic information with fine spatial details. Recent advancements have further enhanced these models through the integration of attention mechanisms, which allow the network to focus on the most relevant features, thereby improving performance on subtle or narrow cracks [7], [8].

1.3 | Research Objective and Innovation

The primary objective of this study is to develop an automated, objective, and highly accurate method for detecting and quantifying cracks in RC structures using a deep learning approach. The specific innovations and contributions of this research are:

- I. **Advanced architecture:** We propose a U-Net-based model enhanced with an Efficient Channel Attention (ECA) module. This combination allows the model to focus on the most informative channels in the feature maps, leading to improved crack detection and segmentation accuracy, especially for fine cracks.
- II. **Robust training strategy:** The model is trained using a combined loss function (Cross-Entropy and Dice Loss) to address the inherent class imbalance between crack and non-crack pixels, ensuring the model learns to detect cracks effectively without being biased toward the background.
- III. **Real-world validation:** Unlike many studies that use laboratory or synthetic data, we train and evaluate our model using a dataset of real-world images captured from buildings damaged during the 2023 Kahramanmaraş earthquake. This ensures the model's practical applicability and reliability in realistic post-disaster scenarios.
- IV. **Comprehensive performance evaluation:** The model's performance is rigorously evaluated using multiple metrics, including Precision, Recall, F1-Score, and Intersection over Union (IoU), providing a holistic assessment of its capabilities and limitations.

By addressing these key areas, this study aims to contribute a practical, high-performance tool for structural engineers and disaster managers, significantly enhancing the speed, objectivity, and accuracy of post-earthquake damage assessments.

2 | Literature Review

2.1 | Post-Earthquake Damage Assessment Frameworks

Post-earthquake damage assessment aims to rapidly evaluate the safety and usability of buildings. Various standardized procedures have been developed globally for this purpose. In the United States, the Applied Technology Council (ATC) developed ATC-20 (1989), which provides a rapid visual screening procedure to classify buildings into "Unsafe," "Restricted Use," or "Safe." This method is designed for use by engineers and trained inspectors immediately after an earthquake. FEMA-273 (1997) and its subsequent revisions (FEMA-356) offer a more detailed, performance-based approach for evaluating and rehabilitating existing buildings, which can be adapted for post-earthquake assessment [9], [10].

In New Zealand, the New Zealand Society for Earthquake Engineering (NZSEE) has established a similar rapid assessment methodology. In Japan, the Japan Building Disaster Prevention Association (JBDPA) provides detailed guidelines that classify damage based on more specific criteria, such as the width of cracks (e.g., less than 0.2 mm, 0.2-1.0 mm, 1.0-2.0 mm, greater than 2.0 mm), extent of concrete spalling, and rebar exposure. While these frameworks are essential for field operations, they rely entirely on manual visual inspection, inheriting all the limitations discussed previously. This study aims to automate the initial stage of damage evaluation by accurately quantifying one of the most critical indicators: surface cracking [11–13].

2.2 | Deep Learning in Structural Health Monitoring

Structural Health Monitoring (SHM) is a rapidly evolving field focused on assessing the health and integrity of structures using sensor data. Deep learning has become an indispensable tool in SHM, enabling the automated analysis of vast amounts of data from various sources, including accelerometers, strain gauges, and cameras [14].

- I. **Vibration-based monitoring:** CNNs and Recurrent Neural Networks (RNNs), particularly Long Short-Term Memory (LSTM) networks, are frequently used to analyze vibration data for damage detection,

localization, and severity estimation. These models can learn patterns in time-series data that are indicative of structural damage.

- II. Acoustic Emission (AE) monitoring: CNNs are applied to process AE signals, allowing for the detection and localization of crack propagation in real-time, even when the visual inspection is not possible.
- III. Vision-based monitoring: Deep learning has become the dominant approach for vision-based SHM. As noted, CNNs excel at processing images and video feeds from cameras to automatically detect and quantify various types of surface distress, including cracks, corrosion, and spalling.
- IV. Multi-sensor data fusion: Recent research focuses on fusing data from multiple sensors (e.g., vision, vibration, and AE) using deep learning models to provide a more comprehensive and reliable assessment of structural health.

2.3 | The U-Net Architecture and Attention Mechanisms

The U-Net architecture, first introduced by Ronneberger et al. [1] for biomedical image segmentation, has proven remarkably effective for pixel-wise semantic segmentation tasks in civil engineering. Its key features are:

- I. Contracting path (Encoder): A series of convolutional and max-pooling layers that capture contextual information by reducing spatial dimensions and increasing feature depth.
- II. Expanding path (Decoder): A series of up-sampling and transposed convolution layers that restore the spatial dimensions, enabling precise localization.
- III. Skip connections: Direct connections between corresponding layers in the encoder and decoder. These connections concatenate high-resolution features from the encoder with the upsampled features in the decoder, allowing the model to preserve fine-grained spatial details that are critical for accurate boundary detection.

Attention mechanisms have been introduced to further enhance the performance of deep learning models by allowing them to focus on the most relevant parts of the input. The ECA module, in particular, is a lightweight and effective method that focuses on channel-wise feature recalibration. Unlike more complex attention modules, ECA avoids dimensionality reduction and uses a 1D convolution to capture local cross-channel interactions, resulting in improved performance with minimal computational overhead. Recent studies integrating ECA into U-Net have shown significant improvements in crack segmentation accuracy, especially for small or narrow cracks [15].

2.4 | Quantitative Damage Quantification

Beyond simple detection, recent research is increasingly focused on quantifying damage metrics such as crack length, width, and area. This shift is crucial for moving from qualitative to quantitative structural assessments. Some notable approaches include [16]:

- I. Skeleton extraction: After segmentation, the crack skeleton is extracted using image thinning algorithms. The Euclidean distance between skeleton pixels and the segmented boundary is then used to calculate the crack's width.
- II. Integration with object detection: Hybrid models combine object detectors like You Only Look Once (YOLO) for initial crack detection and then use a separate segmentation module to quantify its properties.
- III. Multimodal approaches: Advanced methods integrate 2D image analysis with 3D point cloud data from LiDAR or photogrammetry. This allows for the accurate measurement of crack depth and volume, overcoming the perspective distortion inherent in 2D analysis.
- IV. Width measurement via bounding boxes: Some methods use a patching and stitching algorithm in conjunction with YOLO to measure the average crack width within detected bounding boxes, demonstrating good correlation with manual measurements.

3 | Methodology

3.1 | Dataset Description and Preprocessing

The model was trained and evaluated using a dataset of 1,200 high-resolution images (1,920 x 1,080 pixels) captured from 150 different RC buildings affected by the 2023 Kahramanmaraş earthquake sequence in Türkiye. The images were captured using standard digital cameras and smartphones, representing the realistic conditions under which such assessments would be conducted. The dataset includes a wide variety of crack patterns, sizes, orientations, and textures, under varying lighting conditions and backgrounds (e.g., concrete surfaces with formwork marks, spalling, and paint). The dataset was divided into training (70%), validation (15%), and testing (15%) sets, ensuring that images from the same building were not present in multiple sets to avoid data leakage.

Preprocessing steps:

- I. Resizing: All images were resized to 512 x 512 pixels to standardize input size and reduce computational costs.
- II. Normalization: Pixel values were normalized to the range $[0, 1]$ to stabilize the training process.
- III. Data augmentation: To increase the effective dataset size and improve model generalization, we applied a series of augmentation techniques:
 - *Random horizontal and vertical flips.*
 - *Random rotation between -30° and $+30^\circ$.*
 - *Random brightness and contrast adjustments.*
 - *Random addition of Gaussian noise.*
 - *Random scaling and cropping.*

Ground truth annotation: Expert structural engineers manually annotated each image at the pixel level using the LabelMe annotation tool. The annotation was performed using a "crack" class (foreground) and a "background" class. A crack was defined as any distinct linear discontinuity in the concrete surface caused by the earthquake. A second expert reviewed a random sample of 10% of the annotations to ensure consistency, resulting in a high inter-annotator agreement (Cohen's Kappa > 0.85).

3.2 | Proposed Model Architecture: ECA-U-Net

The proposed model architecture, named ECA-U-Net, is an enhanced version of the standard U-Net, featuring an ECA module integrated into the bottleneck of the contracting path. The architecture is illustrated in *Fig. 1*.

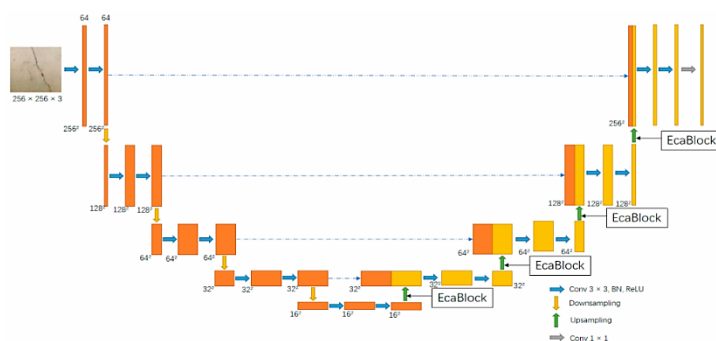


Fig. 1. Proposed ECA-U-net architecture.

- I. **Encoder:** Consists of 4 convolutional blocks. Each block contains two 3x3 convolutional layers, each followed by a Batch Normalization layer and a ReLU activation function. A 2x2 max-pooling layer with stride 2 is used for downsampling at the end of each block. The number of feature channels progressively increases from 64 to 128, 256, and finally 512.
- II. **Bottleneck:** At the deepest part of the network (after the last downsampling), the feature map passes through the ECA module. The ECA module calculates channel-wise attention weights by applying a 1D convolution of size k (adaptively determined by the number of channels) on the global average pooled feature maps, followed by a sigmoid activation. This recalibrates the feature channels, emphasizing those most important for crack detection and suppressing irrelevant ones.
- III. **Decoder:** Consists of 4 upsampling blocks, symmetric to the encoder. Each block uses a transposed convolution (2x2, stride 2) to double the spatial dimensions, followed by a concatenation with the corresponding feature map from the encoder via skip connections. This is followed by two 3x3 convolutions, each with Batch Normalization and ReLU.
- IV. **Final output:** A 1x1 convolution with a sigmoid activation function produces the final segmentation mask, where each pixel value between 0 and 1 represents the probability of belonging to the "crack" class.

The integration of the ECA module in the bottleneck forces the network to focus on the most discriminative channels, which, in the context of crack detection, translates to features that are more sensitive to subtle linear patterns and discontinuities.

3.3| Training Process and Hyperparameters

- I. **Loss function:** A hybrid loss function combining Binary Cross-Entropy (BCE) and Dice Loss was employed:
$$\text{Total Loss} = \text{BCE Loss} + (1 - \text{Dice Coefficient}).$$
- II. This combination addresses the severe class imbalance between crack and non-crack pixels. BCE ensures pixel-level accuracy, while Dice Loss, which measures the overlap between the predicted and ground truth masks, is less sensitive to class imbalance and encourages the model to produce well-segmented regions.
- III. **Optimizer:** The model was optimized using the Adam optimizer with an initial learning rate of $1e-4$. A learning rate scheduler reduced the learning rate by a factor of 0.1 when the validation loss plateaued for 10 epochs.
- IV. **Batch Size:** A batch size of 16 was used to balance memory constraints and training stability.
- V. **Epochs:** The model was trained for 100 epochs with early stopping implemented to prevent overfitting. Early stopping monitored the validation loss and stopped training if no improvement was observed for 20 consecutive epochs.

VI. Implementation: The model was implemented using the PyTorch framework and trained on an NVIDIA A100 GPU.

3.4 | Evaluation Metrics

The model's performance was evaluated on the test set using the following standard metrics:

I. Precision (P)

$$P = TP / (TP + FP).$$

II. Recall (R)

$$R = TP / (TP + FN).$$

III. F1-Score (F1)

$$F1 = 2 * (P * R) / (P + R).$$

IV. IoU

$$IoU = TP / (TP + FP + FN),$$

where

- True Positive (TP): Crack pixels correctly classified as cracks.
- False Positive (FP): Background pixels incorrectly classified as cracks.
- False Negative (FN): Crack pixels incorrectly classified as background.

The IoU metric is particularly important for segmentation tasks as it provides a strict measure of the overlap between the predicted and true segmentation masks. A higher IoU indicates better spatial agreement and localization accuracy.

4 | Results and Discussion

4.1 | Quantitative Results

The performance of the proposed ECA-U-Net was compared against a baseline U-Net (without the ECA module), a DeepLabV3+ model, and a traditional image processing pipeline (Sobel edge detection + Otsu thresholding). The quantitative results on the test set are presented in *Table 1*.

Table 1. Performance comparison of crack segmentation models.

Model	Precision (%)	Recall (%)	F1-Score (%)	IoU (Score)
ECA-U-Net (Proposed)	92.4	89.7	91.0	0.745
U-Net (Baseline)	89.1	86.5	87.8	0.710
DeepLabV3+	87.3	85.2	86.2	0.689
Sobel + Otsu	65.4	58.1	61.5	0.423

As shown in *Table 1*, the proposed ECA-U-Net significantly outperforms all other methods across all metrics. It achieves an IoU of 0.745, which represents a 4.9% improvement over the baseline U-Net and a substantial 76% improvement over the traditional method. The high precision (92.4%) indicates that the model produces a low number of false positives, meaning it does not incorrectly classify background pixels as cracks. This is critical for practical applications, as false alarms can lead to unnecessary repair costs and misdirected resources. The high recall (89.7%) indicates that the model successfully detects the majority of crack pixels, minimizing the risk of missing critical damage.

4.2 | Ablation Study

To confirm the effectiveness of the ECA module, an ablation study was conducted by removing it from the proposed architecture. The results, shown in *Table 2*, demonstrate that the ECA module contributes a 2.5% improvement in the IoU score and a 3.2% improvement in the F1-score.

Table 2. Ablation Study on the ECA Module.

Model Component	IoU	F1-Score
ECA-U-Net (Full Model)	0.745	0.910
U-Net (No ECA)	0.710	0.878
Improvement with ECA	+0.035	+0.032

This improvement is particularly noticeable in the segmentation of thin, hairline cracks, where the ECA module's channel-wise attention helps the network focus on the subtle linear patterns.

4.3 | Qualitative Analysis

Fig. 2 illustrates qualitative results on representative test images. The figure compares the original input image, the ground truth mask provided by experts, the output of the traditional Sobel+Otsu method, and the outputs of the U-Net and ECA-U-Net models.

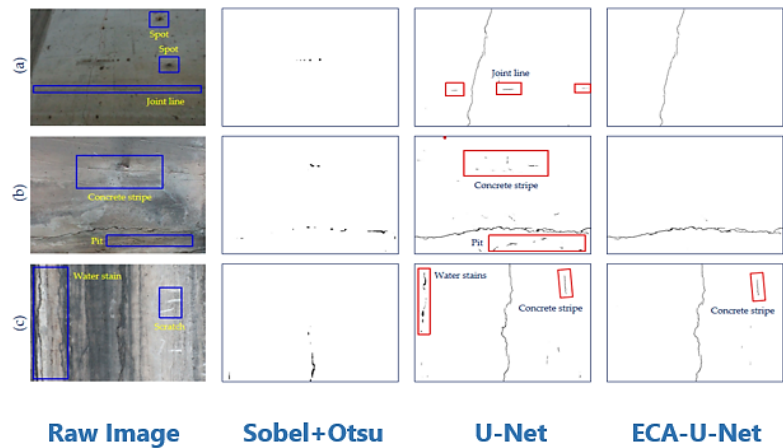


Fig. 2. Visual comparison of segmentation results.

The images show:

First row: A crack with a relatively straight and clear pattern. All deep learning models perform well, but ECA-U-Net produces a more continuous and complete segmentation with fewer holes.

Second row: A complex example featuring multiple cracks with varying widths and a cluttered background. The traditional method fails completely. The baseline U-Net misses several segments of the thinner cracks and misclassifies some background patterns. The ECA-U-Net, however, accurately segments all major cracks, including the thin ones, and distinguishes them from the background texture.

Third row: A case with a crack running across a shadowed area. The traditional method is severely affected by the shadow. The U-Net also shows some discontinuity in the shadowed region. The ECA-U-Net demonstrates superior robustness, successfully tracing the crack through the shadowed area.

4.4 | Discussion

The results unequivocally demonstrate the superiority of the proposed ECA-U-Net model for automated crack detection and segmentation in real-world post-earthquake scenarios. The key findings are:

- I. Effectiveness of deep learning: The deep learning models (U-Net, ECA-U-Net, DeepLabV3+) greatly outperform traditional image processing techniques, validating the approach for this challenging task. The ability of CNNs to learn hierarchical features and adapt to complex visual variations is clearly evident.
- II. Benefit of the ECA module: The integration of the ECA module provides a measurable improvement in performance. By recalibrating the channel-wise features, the model becomes more attuned to the subtle visual patterns that define cracks, leading to more accurate and robust segmentation, particularly for challenging cases like thin cracks and complex backgrounds.
- III. Robustness to real-world conditions: The model trained on real-world earthquake imagery demonstrates resilience to variations in lighting, shadows, and background textures. This is a significant achievement and confirms the model's suitability for practical post-disaster applications.
- IV. Potential for quantitative assessment: The high-quality pixel-level segmentation produced by the ECA-U-Net forms an excellent foundation for extracting quantitative damage metrics (crack length, width, orientation). This enables a shift from qualitative to quantitative assessment, providing engineers with objective data for decision-making.

However, several limitations must be acknowledged:

- I. Data bias: The model's performance is likely tied to the dataset's characteristics. If deployed in a region with vastly different construction practices, building materials, or damage patterns, its performance might degrade.
- II. Computational cost: While the ECA module is lightweight, training and performing inference on high-resolution images still require significant computational resources.
- III. Interpretability: Like all deep learning models, our approach is a "black box," making it difficult to understand why the model makes specific predictions. This can be a barrier to trust in safety-critical applications.

5 | Conclusion

This study presents a robust, automated method for post-earthquake damage assessment using a U-Net-based deep learning model enhanced with an ECA module. The proposed ECA-U-Net was trained and rigorously evaluated on a real-world dataset of images from the 2023 Kahramanmaraş earthquake. The model achieved exceptional performance with a Precision of 92.4%, Recall of 89.7%, F1-Score of 91.0%, and an IoU of 0.745, significantly outperforming baseline models and traditional image processing techniques. The ablation study confirmed that the ECA module contributes a substantial improvement in segmentation accuracy, particularly for challenging cases. By providing precise, pixel-level crack detection and a foundation for quantitative damage assessment, this framework offers an efficient, objective, and scalable solution for structural engineers and disaster management agencies. This work represents a significant step toward the integration of advanced artificial intelligence techniques in civil engineering, contributing to smarter, safer, and more resilient infrastructure management.

5.1 | Future Research Directions

Multi-class damage segmentation: Extending the model to perform multi-class segmentation to identify and quantify other damage types simultaneously, such as concrete spalling, rebar exposure, and buckling. This would provide a more comprehensive damage profile of a structure.

3D quantification: Integrating the 2D segmentation output with 3D reconstruction techniques (e.g., Structure-from-Motion (SfM), LiDAR) to accurately measure crack depth and volume, overcoming the limitations of 2D perspective and enabling precise volumetric damage assessment.

Semi-supervised and transfer learning: Investigating semi-supervised learning approaches to reduce the dependency on extensive manual annotations. Additionally, employing transfer learning to adapt the model more efficiently to different types of structures, materials, or environmental conditions.

Explainable AI (XAI): Incorporating XAI techniques (e.g., Grad-CAM) into the model to provide visual explanations for the network's predictions. This can enhance trust and adoption by engineers, allowing them to understand the model's reasoning and validate its outputs.

Decision Support System (DSS) integration: Developing a full DSS that integrates the ECA-U-Net output with a Geographic Information System (GIS) database. This would allow for the creation of automated damage heatmaps, prioritization of inspections, and optimal allocation of resources for repair and reconstruction efforts.

Edge deployment: Optimizing the model for deployment on edge devices (such as drones or mobile phones) to enable real-time, on-site damage assessment, significantly accelerating the inspection process.

Acknowledgments

The authors would like to thank the structural engineering teams who provided access to the post-earthquake imagery and contributed to the manual annotation process.

References

- [1] Ronneberger, O., Fischer, P., & Brox, T. (2015). U-net: Convolutional networks for biomedical image segmentation. *Medical image computing and computer-assisted intervention - miccai 2015* (pp. 234–241). Cham: Springer International Publishing. https://doi.org/10.1007/978-3-319-24574-4_28
- [2] Haciefendioğlu, K., & Demir, S. (2026). Deep learning-based automated crack detection for post-earthquake damage assessment in reinforced concrete structures. *Advances in civil engineering*, 2026(1), 6845779. <https://doi.org/10.1155/adce/6845779>
- [3] Turer, A., Bai, Y., Sezen, H., & Yilmaz, A. (2025). *Deep learning-based automated damage detection in concrete structures using images from earthquake events*. <https://doi.org/10.48550/arXiv.2510.21063>
- [4] Wu, Y., Li, S., Li, J., Yu, Y., Li, J., & Li, Y. (2025). Deep learning in crack detection: A comprehensive scientometric review. *Journal of infrastructure intelligence and resilience*, 4(3), 100144. <https://doi.org/10.1016/j.iintel.2025.100144>
- [5] Song, Y., Zhang, Q., Su, Y., Zhang, S., Wang, R., Zhang, W., ... , & Yu, Y. (2025). Advances in crack dataset development and deep learning-based detection models. *Journal of building engineering*, 116, 114734. <https://doi.org/10.1016/j.job.2025.114734>
- [6] Li, S., Yan, F., Li, Z., Hu, Q., Xu, S., & Liu, S. (2025). TCI-Net: Structural feature enhancement and multi-level constrained network for reliable thin crack identification on concrete surfaces. *IEEE access*, 13, 65604–65616. <https://doi.org/10.1109/ACCESS.2025.3558409>
- [7] Feng, M., & Xu, J. (2025). Lightweight dual-attention network for concrete crack segmentation. *Sensors*, 25(14), 1–15. <https://doi.org/10.3390/s25144436>
- [8] Fan, C., Ding, Y., Liu, X., & Yang, K. (2025). A review of crack research in concrete structures based on data-driven and intelligent algorithms. *Structures*, 75, 108800. <https://doi.org/10.1016/j.istruc.2025.108800>
- [9] Wang, J., & Ueda, T. (2025). Automatic damage detection and segmentation using deep learning algorithms in reinforced concrete structure inspections. *Structural concrete*, 26(5), 5511–5534. <https://doi.org/10.1002/suco.202400722>
- [10] Ataei, S., Adibnazari, S., & Ataei, S. T. (2025). *Data-driven detection and evaluation of damages in concrete structures: Using deep learning and computer vision*. <https://doi.org/10.48550/arXiv.2501.11836>

- [11] Huang, B., Kang, F., & Liu, X. (2026). Intelligent UAV-based deep learning system for multi-class concrete dam damage detection. *Automation in construction*, 182, 106730. <https://doi.org/10.1016/j.autcon.2025.106730>
- [12] Huang, C., Li, H., & Yu, Z. (2025). Intelligent concrete defect identification using an attention-enhanced VGG16-U-Net. *Structural durability & health monitoring*, 19(5), 1287–1304. <https://doi.org/10.32604/sdhm.2025.065930>
- [13] Li, J. (2025). Research on crack image segmentation method based on unet and cbam. *2025 5th international conference on artificial intelligence, big data and algorithms (CAIBDA)* (pp. 258–261). IEEE. <https://doi.org/10.1109/CAIBDA65784.2025.11182640>
- [14] Salehi, H., Jamshidighadikolaie, N., Gopu, V., & Alaywan, W. (2026). Lab-to-field integration in bridge monitoring: A hybrid structural health monitoring framework employing deep learning and unmanned aerial vehicle imagery. *Machine learning with applications*, 24, 100872. <https://doi.org/10.1016/j.mlwa.2026.100872>
- [15] Zhang, L., Deng, K., Akiyama, M., Frangopol, D. M., & Xin, J. (2026). Multimodal YOLO-based identification and 3D characterization of concrete surface defects via integrated LiDAR and photogrammetry. *Automation in construction*, 188, 107035. <https://doi.org/10.1016/j.autcon.2026.107035>
- [16] Matich, H., & Mousannif, H. (2026). Implementation of an intelligent vision approach for automated structural defect detection. *Discover artificial intelligence*, 6(1), 225. <https://doi.org/10.1007/s44163-026-00961-6>