

Benchmarking Yolov8n, Yolov11n, and Yolov12n for Smart Road Pothole Detection

Ilham Arif Febriansyah^{1,*}, Eko Darmanto¹, Pratomo Setiaji¹

¹ Universitas Muria Kudus, Indonesia

* Corresponding author: Ilham Arif Febriansyah, Universitas Muria Kudus, Indonesia
✉ 202153079@std.umk.ac.id

Copyright: © 2026 by the authors

Received: 30 June 2026 | Revised: 09 July 2026 | Accepted: 10 August 2026 | Published: 24 August 2026

Abstract

Road potholes present a considerable threat to traffic safety, highlighting the need for precise and efficient automatic detection in road-monitoring systems. Although recent research has utilized various YOLO-based models for detecting potholes, there is a scarcity of comparative evaluations of the newest lightweight YOLO architectures under uniform experimental conditions. This study seeks to evaluate the performance of YOLOv8n, YOLOv11n, and YOLOv12n in the context of smart road pothole detection. A quantitative experimental approach was applied using the road pole dataset. Each model underwent training using identical preprocessing, training configurations, and evaluation protocols. Performance metrics included Precision, Recall, mAP50, mAP50–95, inference time, and computational efficiency. The experimental findings indicate that YOLOv11n delivered the best overall results, achieving 90.02% precision, 90.25% mAP50, and 48.35% mAP50–95, along with the quickest inference time of 3.13 ms and the highest end-to-end processing speed. Although YOLOv12n recorded the highest recall at 86.11%, its overall detection performance and computational efficiency were inferior to those of YOLOv11n. These results suggest that YOLOv11n offers the most balanced compromise between detection accuracy and computational efficiency, providing empirical support for choosing lightweight object detection models for intelligent road-monitoring systems.

Keywords: computer vision; deep learning; object detection; road pothole detection; yolo

To cite this article: Febriansyah, I. A., Darmanto, E., & Setiaji, P. (2026). Benchmarking Yolov8n, Yolov11n, And Yolov12n For Smart Road Pothole Detection. *Edumatic: Jurnal Pendidikan Informatika*, 10(2), 400–409. <https://doi.org/10.29408/edumatic.v10i2.35908>

INTRODUCTION

Road transportation plays a crucial role in intelligent transportation systems (ITS) and the development of smart cities, where dependable infrastructure monitoring is vital for maintaining traffic safety, enhancing mobility efficiency, and sustainably managing urban areas. The latest progress in computer vision and deep learning has facilitated the creation of automated visual inspection systems that aid infrastructure maintenance by providing precise, real-time object recognition (Ranyal et al., 2022). As a result, smart road monitoring has emerged as a prominent research field, offering a scalable alternative to traditional inspection techniques, while also lowering operational expenses and enhancing decision-making for transportation authorities (Ettalibi et al., 2024; Rasee et al., 2025).

Among the various forms of road infrastructure damage, potholes are particularly significant owing to their role in heightening accident risks, diminishing driving comfort, and hastening vehicle wear and tear. Consequently, the precise detection of potholes is essential for



proactive road maintenance and effective infrastructure management (Giordani et al., 2026; Safyari et al., 2024). However, the evaluation of road conditions largely relies on manual visual inspections, which are laborious, time intensive, expensive, and challenging to implement uniformly across large road networks (Zanevych et al., 2024). These challenges have led to the increasing use of automated computer vision methods that can identify road damage using digital images or video footage.

Numerous deep learning methods have been explored for detecting road damage, such as image classification, semantic segmentation, and object detection. Object detection stands out for road monitoring, as it combines object localization and classification in one framework, offering the possibility of real-time processing (Gorro et al., 2024; Setiaji et al., 2025). This feature is crucial for intelligent transportation systems, where detection mechanisms might need to handle continuous images or video streams with limited computational power. In this scenario, lightweight YOLO architectures have gained significant interest because they aim to preserve detection accuracy while minimizing computational demands, thus enabling their use on edge devices with constrained resources (Lee & Hwang, 2022; Lin & Pan, 2025; Sahputra & Muhammad, 2026).

Recent research has highlighted the success of YOLO-based methods in detecting potholes and pavement damage (Li et al., 2025). However, the outcomes reported are often not directly comparable owing to the use of varied datasets, annotation methods, image preprocessing techniques, training setups, and evaluation protocols across studies. Consequently, discrepancies in reported performance may arise not only from architectural features but also from differing experimental conditions. This concern is especially significant when assessing lightweight models, where even minor architectural modifications can influence the trade-off between detection accuracy, computational demands, and inference speed. Additionally, earlier studies indicated that model performance can differ significantly in real-world scenarios, particularly under challenging environmental and visual conditions (Jakubec et al., 2023; Reddy et al., 2026).

Evaluating the true impact of architectural changes within the YOLO family presents a significant challenge owing to methodological variations. Although the latest YOLO versions incorporate modifications aimed at enhancing feature representation, detection accuracy, and computational efficiency, it is uncertain whether these modifications consistently lead to better pothole detection outcomes when models are tested under identical experimental conditions. Performance discrepancies may be influenced by differences in data composition, annotation methods, training settings, or evaluation standards without a controlled benchmarking framework. Therefore, to determine the influence of the model architecture on the detection performance more accurately, a standardized comparison using the same dataset, preprocessing approach, training setup, and evaluation protocol is essential.

This study performed a controlled cross-generational comparison of YOLOv8n, YOLOv11n, and YOLOv12n for detecting road potholes, utilizing a consistent experimental setup. The three models were assessed using the same dataset, annotation method, training setup, and evaluation protocol to reduce methodological differences and allow for a clearer comparison of their architectural features. The assessment examines not only detection performance metrics, such as Precision, Recall, mAP50, and mAP50–95, but also computational aspects, such as parameter count, GFLOPs, training duration, inference time, and processing speed. This thorough evaluation aims to ascertain whether the advancements in newer YOLO versions are linked to significant improvements in computational efficiency and practical real-time performance.

The significance of this study lies in its experimentally controlled evaluation of how lightweight YOLO architectures have evolved to detect potholes. Instead of proposing another modified detector, the research aims to create a fair empirical foundation to assess if newer

YOLO versions consistently outperform older versions when key experimental conditions remain unchanged. The findings shed light on the balance between accuracy and efficiency among YOLOv8n, YOLOv11n, and YOLOv12n, providing practical advice for choosing lightweight detection models suitable for real-time road condition monitoring and intelligent transportation systems.

METHOD

In this study, an experimental benchmarking method was utilized to assess the performance of three lightweight YOLO architectures for detecting road potholes under the same experimental conditions. To ensure an unbiased comparison, the benchmarking framework used a consistent dataset, annotation process, training setup, and evaluation protocol. The research process included acquiring the dataset, extracting frames, cleaning data, annotating images, splitting the dataset, training models, and evaluating performance. Figure 1 illustrates the workflow.

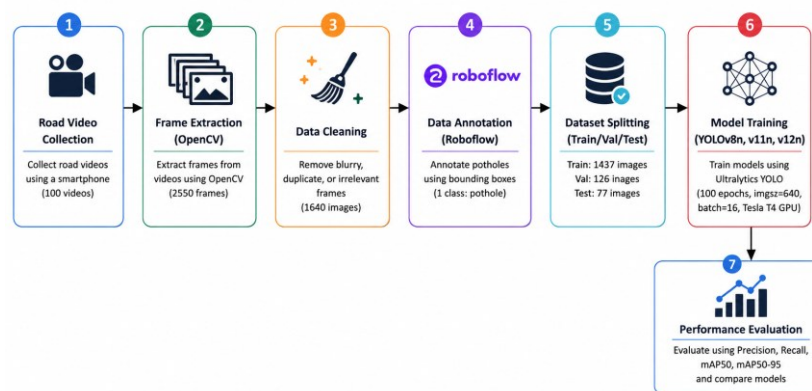


Figure 1. Pothole detection research workflow

The data were gathered in Karanganyar District, Demak Regency, employing a smartphone camera with 1080p resolution under diverse road and lighting scenarios. Out of 100 videos, 2,550 frames were extracted through OpenCV and then refined to 1,640 unique, high-quality images, which were annotated in YOLO format with a single pothole category. The dataset was split into 1,437 images for training, 126 for validation, and 77 for testing.

The YOLOv8n, YOLOv11n, and YOLOv12n models were trained with Ultralytics on a Tesla T4 GPU through Google Colab, maintaining consistent conditions (input size of 640×640 , batch size of 16, and 100 epochs) to allow for an equitable comparison. The models' detection capabilities were measured using Precision, Recall, mAP@50, and mAP@50–95 metrics, while their computational efficiency was evaluated based on training duration, inference duration, FPS, parameter count, and GFLOPs.

The evaluation process utilized four main object detection metrics: Precision, Recall, mAP@50, and mAP@50–95. Precision measures the reliability of positive predictions by comparing correctly detected objects with all positive predictions produced by the model. Recall evaluates the model's ability to identify actual pothole objects by measuring the proportion of correctly detected objects against all existing objects. mAP@50 assesses detection performance using an IoU threshold of 0.50, while mAP@50–95 provides a stricter evaluation by calculating the average precision across IoU thresholds ranging from 0.50 to 0.95. These metrics provide a comprehensive assessment of both detection accuracy and bounding box localization capability.

$$Precision = \frac{TP}{TP + FP} \quad (1)$$

Precision in Equation (1) represents the accuracy of the model's positive predictions, where True Positive (TP) indicates correctly detected potholes and False Positive (FP) represents incorrect detections. A higher Precision value indicates that the model produces fewer false detections. Recall in Equation (2) describes the model's capability to detect all existing potholes, where False Negative (FN) represents objects that are present but not detected by the model. A higher Recall value indicates better detection coverage.

$$Recall = \frac{TP}{TP + FN} \quad (2)$$

mAP@50 in Equation (3) measures the detection performance based on an IoU threshold of 0.50, evaluating both object recognition and bounding box localization accuracy. Meanwhile, mAP@50–95 in Equation (4) evaluates model performance across multiple IoU thresholds from 0.50 to 0.95, providing a more comprehensive measurement of localization consistency and detection robustness.

$$mAP50 = \frac{1}{N} \sum_{i=1}^N AP_i \quad (3)$$

$$mAP50-90 = \frac{1}{10} \sum_{iou=0.50}^{0.95} AP \quad (4)$$

RESULTS AND DISCUSSION

Results

The Road Pothole Dataset utilized in this study includes road images marked with pothole bounding boxes formatted for YOLO. Figure 2 shows sample annotated images employed for benchmarking purposes. This dataset offers a uniform foundation for assessing the detection performance of YOLOv8n, YOLOv11n, and YOLOv12n under identical experimental conditions.

Figure 2 displays sample annotated images from the Road Pothole Dataset, which were utilized during the benchmarking phase. Potholes are marked with bounding boxes in YOLO format, offering uniform ground-truth annotations for both model training and performance assessment. This annotation method guarantees uniform object localization throughout the dataset, allowing for an equitable comparison among YOLOv8n, YOLOv11n, and YOLOv12n under the same experimental settings.

Table 1. Performance comparison of yolov8n, yolov11n, and yolov12n

Model	Precision (%)	Recall (%)	mAP50 (%)	mAP50–95 (%)
YOLOv8n	86,245	80,952	84,871	42,778
YOLOv11n	90,020	85,903	90,248	48,348
YOLOv12n	85,770	86,109	89,390	48,133

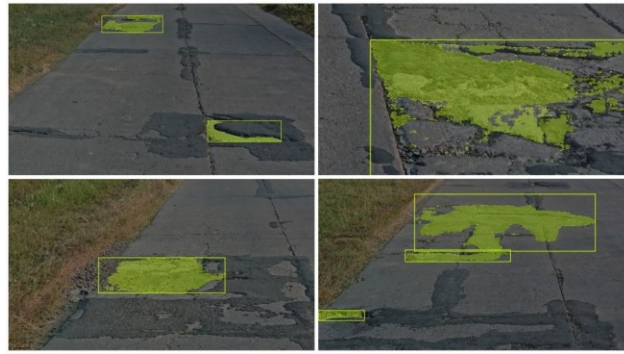


Figure 2. Example of annotation results of pothole dataset using roboflow

Table 1 summarizes the benchmarking results for the three lightweight YOLO models. Among them, YOLOv11n recorded the highest precision at 90.02%, mAP50 at 90.25%, and mAP50–95 at 48.35%. In contrast, YOLOv12n excelled in recall, achieving 86.11%. These findings suggest that architectural advancements in YOLOv11n enhance both detection accuracy and localization performance, whereas YOLOv12n appears to focus more on object coverage than on prediction precision. In summary, YOLOv11n offered the most effective balance between the detection capability and localization quality when tested under the same experimental conditions.

Table 2. Comparison of computational efficiencies of yolov8n, yolov11n, and yolov12n

Model	Layers	Parameters	GFLOPs	Training Time (hours)	Inference Time (ms)	FPS Inference	End-to-End FPS
YOLOv8n	73	3,005,843	8.1	1,069	3.23	317.32	95.30
YOLOv11n	101	2,582,347	6.3	1,098	3.13	320.12	97.24
YOLOv12n	159	2,556,923	6.3	1,135	7.17	141.76	63.68

Table 2 presents a comparison of the computational efficiency of the three models, focusing on model complexity, training cost, and inference performance. Despite YOLOv12n having the least number of parameters (2.56 M), it exhibited the longest inference duration (7.17 ms), suggesting that a smaller model size does not always lead to quicker execution. Conversely, YOLOv11n not only achieves the shortest inference time, but also the highest inference FPS and end-to-end FPS, while maintaining a relatively low computational complexity. These results indicate that YOLOv11n provides a more advantageous balance between computational efficiency and detection performance, making it more appropriate for real-time applications.

Figures 3–5 illustrate the learning curves for YOLOv8n, YOLOv11n, and YOLOv12n, respectively, across 100 training epochs. Throughout the training process, each model exhibits a steady decline in both training and validation losses, while Precision, Recall, mAP@50, and mAP@50–95 progressively increase until they stabilize as the training concludes. The lack of significant fluctuations suggests stable convergence under identical training conditions, indicating that all three models were sufficiently trained for subsequent performance evaluations.

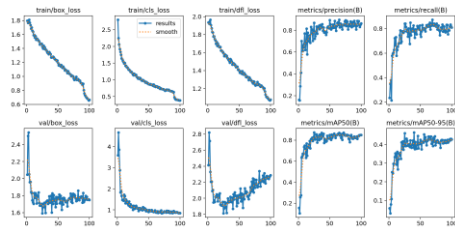


Figure 3. Learning curve for the yolov8n model.

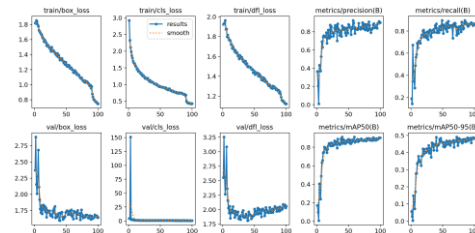


Figure 4. Learning curve for the yolov11n model

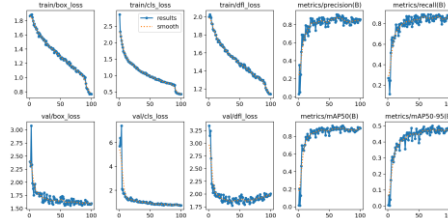


Figure 5. Learning curve for the yolov12n model

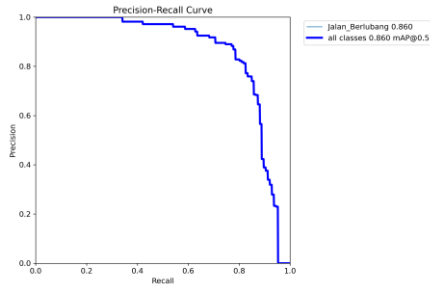


Figure 6. Precision–recall curve of yolov8n model

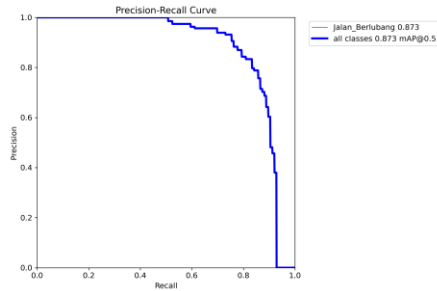


Figure 7. Precision–recall curve of yolov11n model

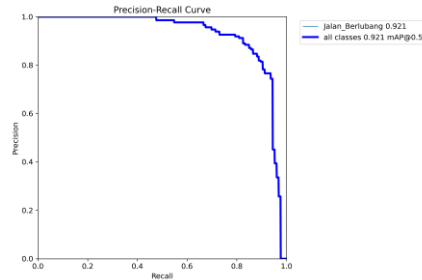


Figure 8. Precision–recall curve yolov12n model

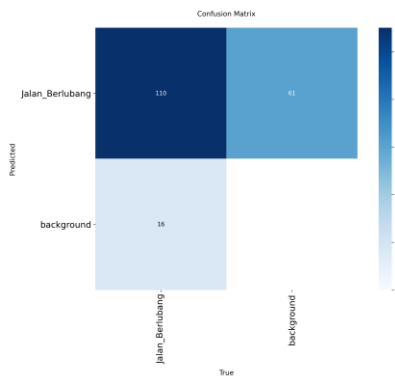


Figure 9. Confusion matrix model yolov8n

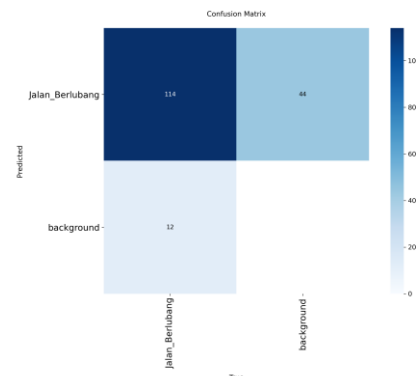


Figure 10. Confusion matrix model yolov11n

Figures 6, 7, and 8 display the precision–recall curves for YOLOv8n, YOLOv11n, and YOLOv12n when applied to the test dataset. These models consistently exhibit high precision over a broad spectrum of recall values, suggesting that they generalize well to new pothole images. The Average Precision (AP) values for these models were 0.860 for YOLOv8n, 0.873 for YOLOv11n, and 0.921 for YOLOv12n. Despite YOLOv12n achieving the highest AP in the representative curve, Table 1's repeated benchmarking results consistently show that YOLOv11n delivers a better average detection performance across various evaluation runs.

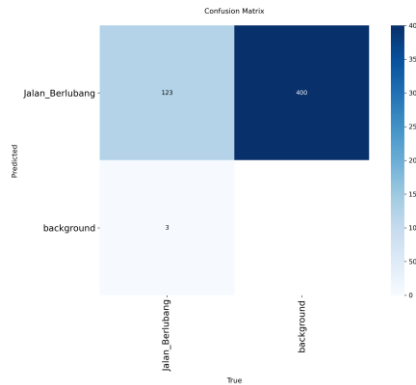


Figure 11. Confusion matrix model yolov12n

Figures 9, 10, and 11 display the confusion matrices for YOLOv8n, YOLOv11n, and YOLOv12n, respectively, on the test dataset. Most predictions aligned with the main diagonal, signifying that all three models differentiated most potholes from background areas. However, each model still produces false positives and negatives, indicating that road textures and visually similar pavement conditions continue to pose difficulties for object recognition. Among the models, YOLOv11n shows the fewest classification errors, aligning with its higher precision and mAP values shown in Table 1. These findings further confirm YOLOv11n's robustness.



Figure 12. Example of YOLOv11n detection results on test data.

Figure 12 shows typical detections using YOLOv11n on the test set. Although the majority of potholes were precisely identified, there were instances of missed and incorrect detections, particularly in areas with low contrast and uneven pavement texture. These examples enhance the quantitative findings by demonstrating common detection mistakes and verifying the model's capability to identify potholes under diverse road conditions.

Among the lightweight models assessed, YOLOv11n demonstrated the best performance. It excelled in blending detection accuracy, localization capability, and inference efficiency, while keeping computational demands relatively low. This conclusion is

consistently supported by quantitative metrics, such as precision–recall curves, confusion matrices, training patterns, and visual outcomes, which collectively suggest that YOLOv11n provides an advantageous balance between accuracy and efficiency for real-time pothole detection in the tested scenarios.

Discussion

The benchmark reveals that YOLOv11n achieves the best-rounded performance among the three lightweight architectures. This advantage stems from the combination of detection accuracy, localization quality, and inference efficiency under the same experimental conditions. This observation implies that newer YOLO architectures do not always outperform each other in every aspect, highlighting the necessity of evaluating architectural modifications through controlled benchmarking. Consequently, the results offer a reliable foundation for comparing the practical distinctions between lightweight YOLO generations (Zeng & Zhong, 2024).

YOLOv11n's performance enhancements are linked to advancements in feature extraction and aggregation, which aid in representing potholes of diverse sizes, shapes, and visual traits. By effectively aggregating the features, the local surface details are maintained while integrating the contextual information necessary for precise localization. Consequently, more efficient feature processing can enhance detection capabilities without a corresponding increase in computational demands. Similar advantages of lightweight architectural optimization have been observed in studies on pavement defect detection (Li et al., 2024; Lin et al., 2026).

YOLOv12n demonstrated a higher Recall but a lower Precision compared to YOLOv11n, suggesting it is more sensitive but also prone to more false positive detections. In road inspection, achieving high recall is crucial because overlooking potholes can delay repairs and pose safety hazards. However, too many false positives can undermine the trustworthiness of automated inspection outcomes. This tendency may be attributed to the attention-focused architecture of YOLOv12, which improves feature representation and global contextual understanding. Therefore, the ideal balance between Recall and Precision depends on whether the application emphasizes comprehensive pothole detection or dependable automated identification (Tian et al., 2025).

The results align with previous research, indicating that minor architectural changes can enhance the detection of pavement defects. Improvements in feature extraction, multiscale processing, and lightweight components have been noted in RDD-YOLO, YOLO11-WLBS, and YOLOv8-PD (Li et al., 2024; Lin et al., 2026; Zeng & Zhong, 2024). Nonetheless, caution is advised when making direct numerical comparisons, because the studies employed varied datasets, annotation distributions, training setups, and evaluation methods. The detection performance may also be affected by differences in pothole size, contrast, shape, and texture of the surrounding pavement. Consequently, the observed advantage of YOLOv11n is specific to the controlled conditions of this benchmark and does not represent a general ranking of YOLO architectures (Li et al., 2026).

These findings suggest that the number of parameters is not the sole factor affecting the practical inference efficiency. Other elements, such as the computational graph, memory access patterns, feature aggregation, and parallel execution capabilities, also play significant roles in determining the runtime. As a result, a model with a smaller parameter count may not always achieve reduced latency. This study offers empirical support for the notion that architectural efficiency should be assessed in conjunction with model size and detection performance. By employing consistent experimental protocols across different YOLO generations, this study clarifies the connection between architectural design and computational efficiency.

The results hold significant practical importance for intelligent transportation systems (ITS), especially in the realm of automated road condition monitoring. YOLOv11n's balance makes it ideal for intelligent road inspections, enabling the processing of images or video streams to detect potholes and assist in prioritizing maintenance tasks. Its computational efficiency is also advantageous for edge AI applications that require local and near-real-time processing (Abdelwahed et al., 2025). Additionally, integration with cameras mounted on vehicles could enhance the monitoring of smart city infrastructure and Advanced Driver Assistance Systems (ADAS), where both reliable detection and minimal latency are crucial.

This study was limited by its reliance on a single dataset and controlled experimental settings, which may not adequately capture the variety found in real-world road environments. Future investigations should perform cross-dataset validation under different lighting, weather, road surface, and camera conditions. Additionally, examining newer versions of YOLO, such as YOLOv26, alongside transformer-based detectors, can provide insights into the progression of model architectures. Further exploration is required in areas such as model compression, knowledge distillation, quantization, domain adaptation, and explainable object detection to enhance the deployment efficiency, generalization, and interpretability of the model.

CONCLUSION

This study introduced a cross-generational benchmarking framework for YOLOv8n, YOLOv11n, and YOLOv12n, using the same datasets, training setups, and evaluation methods for pothole detection. By standardizing these elements, the framework enables an impartial assessment of how YOLO's architectural advancements affect performance. Findings showed that YOLOv11n achieved the best balance of accuracy, localization ability, and computational efficiency under tested conditions. The framework offers a foundation to differentiate performance variations due to architectural changes from those arising from experimental differences. The results will aid model selection for real-time road monitoring and intelligent transportation systems. Nonetheless, the conclusions were confined to the dataset and experimental conditions used in this study. Future studies should broaden evaluation through cross-dataset validation, varying illumination, weather, and road conditions, and implementation on edge-computing devices to evaluate robustness and practical deployment performance.

REFERENCES

- Abdelwahed, S. H., Sharobim, B. K., Wasfey, B., & Said, L. A. (2025). Advancements in real - time road damage detection : a comprehensive survey of methodologies and datasets. *Journal of Real-Time Image Processing*, 22(4), 1–13. <https://doi.org/10.1007/s11554-025-01683-1>
- Ettalibi, A., Elouadi, A., & Mansour, A. (2024). AI and Computer Vision-based Real-time Quality Control: A Review of Industrial Applications. *Procedia Computer Science*, 231(2023), 212–220. <https://doi.org/10.1016/j.procs.2023.12.195>
- Giordani, E., Arcioni, L., Gil-martín, M., Foresti, G. L., & Marini, M. R. (2026). Real-world road damage dataset with potholes , cracks , and maintenance holes. *Scientific Reports*, 16(15318), 1–22. <https://doi.org/10.1038/s41598-026-46679-4>
- Gorro, K., Ranolo, E., Roble, L., & Santillan, R. N. (2024). *Road Pothole Detection Using YOLOv8 with Image Augmentation*. 12(4), 417–426. <https://doi.org/10.18178/joig.12.4.417-426>
- Jakubec, M., Lieskovska, E., Bučko, B., & Záborská, K. (2023). Comparison of CNN-based models for pothole detection in real-world adverse conditions: Overview and evaluation. *Applied Sciences*, 13(9), 5810. <https://doi.org/10.3390/app13095810>
- Lee, J., & Hwang, K. (2022). YOLO with adaptive frame control for real - time object detection

- p applications.
- Multimedia Tools and Applications*
- , 36375–36396.
- <https://doi.org/10.1007/s11042-021-11480-0>
- Li, B., Wang, K., Liu, Z., Huang, M., & Zhang, S. (2025). DMC-YOLOv11: an improved pavement damage detection model based on YOLOv11. *Systems Science & Control Engineering An Open Access Journal*, 2583. <https://doi.org/10.1080/21642583.2025.2579987>
- Li, X., Zhang, N., Pan, Y., Lv, Y., Xu, X., & Wang, Z. (2026). A lightweight and attention-enhanced framework for robust pavement defect detection. *Engineering Applications of Artificial Intelligence*, 165(PB), 113545. <https://doi.org/10.1016/j.engappai.2025.113545>
- Li, Y., Yin, C., Lei, Y., Zhang, J., & Yan, Y. (2024). RDD-YOLO: road damage detection algorithm based on improved you only look once version 8. *Applied Sciences*, 14(8), 3360. <https://doi.org/10.3390/app14083360>
- Lin, J., Wang, P., Ruan, Y., & Sun, Y. (2026). YOLO11-WLBS: an efficient model for pavement defect detection. *Scientific Reports*, 16(5284), 1–16. <https://doi.org/10.1038/s41598-026-35743-8>
- Lin, Z., & Pan, W. (2026). YOLO-ROC: a high-precision and ultra-lightweight model for real-time road damage detection. *Measurement Science and Technology*, 37(3), 035405. <https://doi.org/10.48550/arXiv.2507.23225>
- Ranyal, E., Sadhu, A., & Jain, K. (2022). Road condition monitoring using smart sensing and artificial intelligence: A review. *Sensors*, 22(8), 3044. <https://doi.org/10.3390/s22083044>
- Rasee, M. A., Ung, L. L., Tan, G. J., Tan, C. W., Buking, R., Yahya, N., & Ismail, H. (2025). AI-driven Vision-based Pothole Detection for Improved Road Safety. 33(3), 1535–1562. <https://doi.org/10.47836/pjst.33.3.20>
- Reddy, S. S., Janarthanan, M., Khan, I. U., & Amrutha, K. (2026). A Deep Learning Framework for Real-Time Pothole Detection from Combined Drone Imagery and Custom Dataset Using Enhanced YOLOv8 and Custom Feature Extraction. *Mathematics*, 1–32. <https://doi.org/10.3390/math14050898>
- Safyari, Y., Mahdianpari, M., & Shiri, H. (2024). A review of vision-based pothole detection methods using computer vision and machine learning. *Sensors*, 24(17), 5652. <https://doi.org/10.3390/s24175652>
- Sahputra, A., & Muhammad, A. H. (2026). Identity-Aware Lightweight MobileNetV2 with Distillation and Optuna for Face Spoofing Detection. *Edumatic: Jurnal Pendidikan Informatika*, 10(2), 320–329. <https://doi.org/10.29408/edumatic.v10i2.34863>
- Setiaji, P., Triyanto, W. A., & Nurhaliza, M. (2025). Real-Time Traffic Density and Anomaly Monitoring Using YOLOv8, OpenCV and Pattern Recognition for Smart City Applications in Demak. *Jurnal Teknik Informatika (JUTIF)*, 6(4), 1769–1782. <https://doi.org/10.52436/1.jutif.2025.6.4.4867>
- Tian, Y., Ye, Q., & Doermann, D. (2026). Yolov12: Attention-centric real-time object detectors. *Advances in neural information processing systems*, 38, 78433–78457. <https://doi.org/10.52202/085713-2627>
- Zanevych, Y., Yovbak, V., Basystiuk, O., Shakhovska, N., Fedushko, S., & Argyroudis, S. (2024). Evaluation of pothole detection performance using deep learning models under low-light conditions. *Sustainability*, 16(24), 10964. <https://doi.org/10.3390/su162410964>
- Zeng, J., & Zhong, H. (2024). YOLOv8 - PD: an improved road damage detection algorithm based on YOLOv8n model. *Scientific Reports*, 14(12052), 1–14. <https://doi.org/10.1038/s41598-024-62933-z>