

Improved YOLOv11-Based Object Detection Algorithm for Autonomous Driving in Low-Light Conditions

XiYe Luo

Tianjin University of Technology and Education, Tianjin, 300222, China

Email: 841593202@qq.com

How to cite this paper: Luo, X. Y. (2026). Improved YOLOv11-Based Object Detection Algorithm for Autonomous Driving in Low-Light Conditions. *Advances in Engineering Research: Possibilities and Challenges*, 3(2), 75-88. ISSN Print: 3079-5192; ISSN Online: 3079-5206. <https://doi.org/10.63313/AERpc.9074>
Published: 2026-02-10

Copyright © 2026 by author(s) and Erytis Publishing Limited.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Abstract

To address the severe challenges of blurred object features and strong background interference in low-light nighttime environments, which often lead to false detections and missed detections in autonomous driving object detection, this paper proposes an improved YOLOv11-based object detection algorithm for autonomous driving in low-light conditions (IEWS-YOLO). First, an inverted residual attention mechanism called the iEMA module is designed and inserted into the backbone network to enhance detection performance in low-light environments. Second, the standard convolutions in the backbone network are replaced with wavelet transform convolution layers (WTConv), improving global feature representation efficiency and noise suppression capability. Finally, the traditional CIoU loss function is replaced with the ShapIoU loss function to optimize bounding box regression accuracy, making predicted box locations more precise and enhancing object localization. Experimental results show that the IEWS-YOLO model achieves significant improvements in various performance metrics compared to the traditional YOLOv11n model. On the ExDark standard low-light dataset, IEWS-YOLO improves mAP@0.5 by 9.6% and mAP@0.5:0.95 by 6.7% compared to YOLOv11n. In conclusion, the IEWS-YOLO model can perform object detection for autonomous driving in low-light environments more accurately.

Keywords

low-light object detection; YOLOv11; autonomous driving; deep learning

1. Introduction

In the new era, with breakthroughs in artificial intelligence technology, intelligent transportation systems and autonomous driving technology have developed rapidly. In real road environments, lighting conditions are diverse and complex. Especially at night or in low-light scenes, vehicle detection of surrounding objects is a crucial part of accident prevention in autonomous driving, and its importance is self-evident.

Statistics show that the probability of traffic accidents at night is much higher than during the day, indicating an urgent need to improve nighttime driving safety. Nighttime environments are characterized by blurred vision, high noise, and difficulty in color recognition. Under low-light conditions such as night, tunnels, and shadow occlusion, images captured by sensors often suffer from insufficient brightness, significant noise, and loss of detail, leading to a sharp decline in the performance of traditional detection models. These factors collectively reduce detection accuracy and significantly increase the risk of traffic violations and accidents.

Object detection, as one of the core tasks in computer vision, aims to accurately locate objects of interest in images or videos and determine their categories. With the rapid development of deep convolutional neural networks, detection frameworks represented by YOLO (You Only Look Once) and Faster R-CNN continue to evolve, achieving significant progress in both detection accuracy and inference speed. The YOLO series of object detection algorithms has achieved remarkable results in daytime scenes, but still exhibits noticeable missed and false detections in low-light environments. In low-light scenes with overlapping objects and uneven illumination, the model's ability to perceive subtle features such as vehicle edges and contours is insufficient, limiting its application in practical systems.

Currently, object detection techniques for low-light scenes have been widely applied. In the improved YOLO11-based algorithm for small object detection in degraded environments, Peng Zhexuan et al. constructed a lightweight module named C3K2-H by integrating the C3K2 module with the heterogeneous convolution HetConv module, which reduces computational load while maintaining detection accuracy. They also fused the SNet network with the SPPF module to form an SPPFSE module, introducing an attention mechanism to enhance the capture of important object features. In the frequency-domain enhancement-based end-to-end low-light image object detection method, Li Yang et al. achieved end-to-end training for image enhancement and object detection through a joint loss function. This loss incorporates a magnitude difference loss based on Mahalanobis distance to precisely adjust image magnitude and replaces the original regression loss of YOLOv8n with MPDIoU loss to improve detection accuracy. Based on a Retinex decomposition network, Feng et al. proposed an efficient semantic-guided machine vision module (EMV), which dynamically adapts to object detection tasks through a lightweight network for feature decomposition. In the improvement based on the YOLOv7 architecture, Zhao et al. introduced a hybrid convolution module into the backbone structure and feature fusion module of YOLOv7, integrating grouped convolution to enhance the expressive capability of edge regions, thereby significantly improving detection performance in low-light environments. In the research on vehicle recognition in weak-light environments, Zhang Junyi et al., based on YOLOv7, introduced the SimAM attention mechanism to design feature

extraction and feature fusion modules, improving the vehicle detection capability of the network in weak-light conditions.

Although the above methods have improved autonomous driving detection performance in low-light environments to some extent, mainstream low-light autonomous driving detection still faces issues such as poor detection capability, leading to false and missed detections. Therefore, this paper proposes an improved YOLOv11-based object detection algorithm for autonomous driving in low-light conditions (IEWS-YOLO). The specific improvements of this paper are as follows:

- (1) An inverted residual attention mechanism (iEMA) module is designed to enhance the feature quality processed by the neck and detection head, optimize the robustness of overall feature extraction for multi-scale objects in complex and changing scenes, and further improve detection performance in low-light environments.
- (2) The original Conv layers are replaced with WTConv (Wavelet Transform Convolution layers) to address the over-parameterization problem encountered by convolutional neural networks (CNNs) when achieving large receptive fields. This enables CNNs to more effectively capture local and global features, providing a more efficient, robust, and easily integrated convolutional layer solution.
- (3) The original CIoU loss function is replaced with ShapIoU to optimize bounding box regression accuracy, making predicted box locations more accurate and enhancing object localization.

2. YOLOv11 Algorithm

YOLOv11 is a new network open-sourced by Ultralytics in 2024. It improves upon YOLOv8 and supports various tasks such as object detection, instance segmentation, and pose estimation. It inherits the core idea of single-stage, end-to-end detection and incorporates multiple optimizations in network architecture, training strategies, and loss functions. YOLOv11 mainly consists of five models with different parameter sizes: n, s, m, l, and x. YOLOv11n, as the lightest network architecture in the YOLOv11 series, has the smallest number of parameters and floating-point operations, the fastest inference speed, and maintains good detection accuracy. YOLOv11 primarily comprises three parts: the Backbone network, the Neck network, and the Head network.

- (1) Backbone: This is the backbone network of YOLOv11, responsible for extracting multi-level feature maps from input images. YOLOv11 uses the C2f2 module to replace the C2f module from YOLOv8, reducing redundant computations. YOLOv11 also supports quick adaptation to different computational devices through scaling factors.
- (2) Neck: This is the neck network of the YOLOv11 model. It continues the PAN-FPN architecture from YOLOv8 and introduces an Adaptive Feature Fusion (AFF) module to replace traditional concatenation and weighting operations. The AFF module can

automatically assign weights based on the significance of feature maps, enhancing detail features for small objects and semantic features for large objects.

(3) Head: This is the detection head of the YOLOv11 model, also the final layer. YOLOv11 uses a Decoupled Head to replace the Coupled Head from YOLOv8 and introduces an Anchor-Free and Anchor-Based dual-mode switching mechanism to improve bounding box localization accuracy.

3. IEWS-YOLO

This paper improves and optimizes based on the YOLOv11n baseline model, aiming to enhance detection accuracy while reducing model complexity in low-light environments. Figure 1 shows the network structure of IEWS-YOLO. An inverted residual attention mechanism (iEMA) module is designed to improve detection performance in low-light autonomous driving. WTConv is adopted to address the over-parameterization problem in CNNs when achieving large receptive fields, providing a more efficient, robust, and easily integrated convolutional layer. Finally, the Shape-IoU loss function is introduced to enhance object localization and address insufficient regression accuracy.

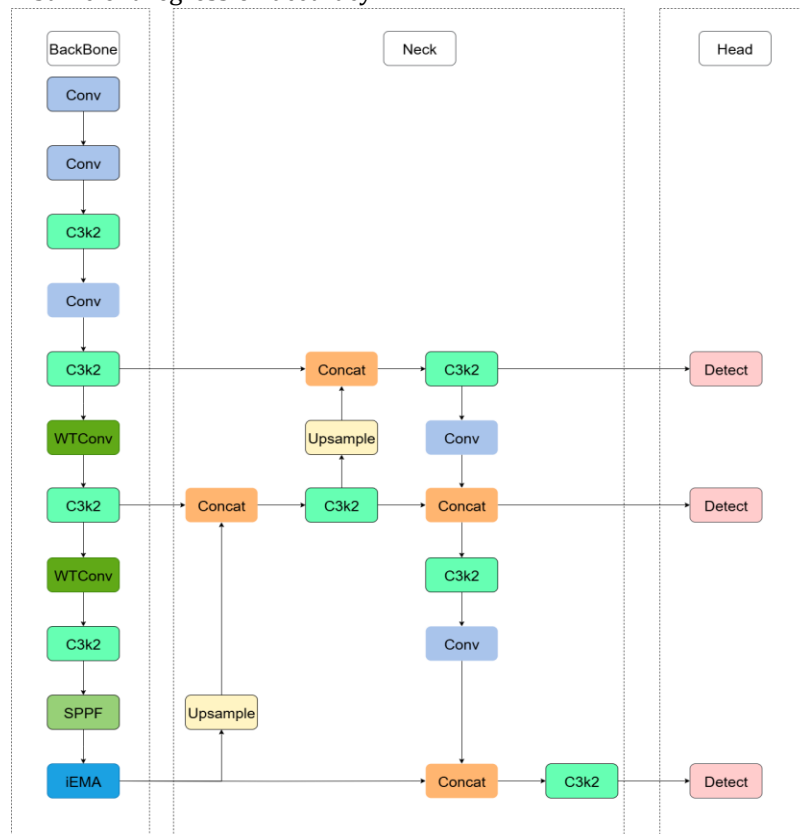


Figure 1. IEWS-YOLO Network Structure

3.1. iEMA Module

The iEMA module is a lightweight multi-scale attention module designed for object detection tasks. It primarily enhances feature extraction and representation capabilities for multi-scale objects in low-light scenes and is commonly used in the backbone network or neck feature fusion layers of detection algorithms like the YOLO series.

The iEMA module features a multi-branch parallel structure and a dual attention mechanism encompassing both channel and spatial dimensions. One branch extracts and enhances local detail features, while another captures global contextual information. The outputs of the two branches are fused and then undergo weight calibration through cross-space learning, enabling the model to adaptively strengthen feature channels and spatial regions rich in information while suppressing responses from noise and irrelevant backgrounds. In low-light autonomous driving object detection tasks, this module can overcome image quality degradation caused by insufficient lighting. Compared to traditional attention mechanisms, it is less susceptible to interference from degradation factors during feature extraction, resulting in better model performance. Traditional attention modules like the ECA module use single-scale channel attention, which cannot adapt to the feature requirements of coexisting large and small objects in object detection. Modules like CBAM suffer from computational redundancy, increasing model computation time.

The iEMA module optimizes traditional attention mechanisms in the following ways: First, it decomposes input feature maps into multi-scale feature maps through parallel convolution branches, capturing global semantic features for large objects and local detail features for small objects, avoiding information redundancy in multi-branch structures. Second, in channel attention design, it performs adaptive pooling operations combining average and max pooling on the multi-scale decomposed feature maps, significantly reducing parameter size and computation time. Finally, iEMA simplifies feature interaction and weight generation methods, ensuring lightweight computation for spatial attention while enhancing the attention mechanism.

In the 1×1 convolutional path of the iEMA module, one-dimensional global average pooling is first used to compute attention in the horizontal and vertical directions. The formula for one-dimensional global average pooling μ is as follows (Equation 1):

$$\mu = \frac{1}{N} \sum_{i=1}^N x_i \quad (1)$$

where N is the length of the spatial dimension (height H or width W), and x_i is the i -th element of the feature map.

In the YOLOv11n model, the iEMA module is added to the backbone network to enhance the feature quality processed by the neck and detection head, optimize the robustness of overall feature extraction for multi-scale objects in complex and changing scenes, and further improve detection performance in low-light

autonomous driving. This enables YOLOv11n to more efficiently integrate feature maps of various depths, enhancing the model's detection accuracy for multi-scale objects in low-light scenes.

3.2. WTConv

In autonomous driving systems under low-light conditions, the accuracy of visual perception is directly related to vehicle safety. Low-light conditions lead to blurred images, reduced contrast, and increased noise. Traditional convolution operations, when extracting image features, cannot simultaneously enhance local details and suppress noise, resulting in a significant decline in detection performance. The WTConv module is a core innovative unit designed for feature degradation in low-light images. Integrating it into the YOLOv11n architecture enables CNNs to more effectively capture local and global features, providing a more efficient, robust, and easily integrated convolutional layer solution.

The WTConv module creates a parallel dual-branch fusion architecture. The first branch is the wavelet transform feature enhancement path. It first performs a discrete wavelet transform on the input feature map, decomposing it into low-frequency and high-frequency sub-bands, achieving preliminary separation of image and noise in the frequency domain. Then, it uses independent convolutional layers to process different frequency sub-bands separately, strengthening edge high-frequency signals while suppressing harmful components caused by lighting noise. The second branch is the lightweight global context path, which employs an optimized Transformer encoder or self-attention module to model dependencies among all pixel positions in the feature map, capturing semantic information about the shapes of fragmented objects.

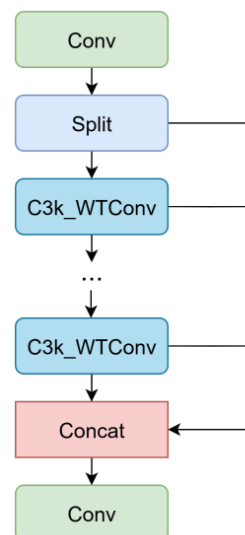


Figure 2. WTConv Module Structure

In the YOLOv11n model, the WTConv module replaces the deep feature extraction

units in the backbone network. The original C2f module is removed and replaced with the C3k2-WTConv module. Its detailed structure is shown in Figure 2. Ultimately, the average detection accuracy for key targets such as vehicles, pedestrians, and traffic signs in low-light autonomous driving detection is significantly improved, showing good robustness especially in extreme situations like dim small objects and strong glare interference.

3.3. ShapeIoU

ShapeIoU is a loss function designed to improve bounding box localization accuracy in object detection. In low-light autonomous driving scenes, image quality degrades due to environmental factors, causing blurred object edges and distorted contours. The traditional CIoU loss function tends to produce deviations during optimization. It cannot fully capture the geometric shapes and scale characteristics of surrounding objects in autonomous driving, leading to decreased localization accuracy for key targets like vehicles and pedestrians, and issues such as box drift or size inaccuracy. The ShapeIoU loss function treats the shape and scale information of target objects as independent optimization factors. Based on calculating the overlap between predicted and ground truth boxes, it innovatively introduces two key cost terms. The shape cost term guides the network to learn to match the actual contour proportions of targets by measuring the difference in aspect ratios between the two boxes. The scale cost term ensures the network produces size-adaptive prediction boxes for targets of different sizes by comparing the similarity of box areas.

In the YOLOv11n model, replacing the original CIoU loss function with ShapeIoU significantly improves bounding box regression accuracy on the low-light dataset. ShapeIoU guides the network to learn the position, shape, and scale of bounding boxes during training, enhancing the model's boundary perception ability for blurred objects in low-light conditions.

The Shape-IoU calculation formulas are as follows:

The distance loss, representing the center distance difference between the predicted box and the ground truth box, is shown in Equation (2):

$$L_{\text{distance}} = hh \times (x_c - x_c^{\text{gt}})/c^2 + ww \times (y_c - y_c^{\text{gt}})/c^2 \quad (2)$$

where c is the diagonal length of the minimum enclosing box, and the weight coefficients ww and hh are given by Equation (3):

$$\begin{cases} hh = \frac{2 \times (w^{\text{gt}})^{\text{scale}}}{(w^{\text{gt}})^{\text{scale}} + (h^{\text{gt}})^{\text{scale}}} \\ ww = \frac{2 \times (h^{\text{gt}})^{\text{scale}}}{(w^{\text{gt}})^{\text{scale}} + (h^{\text{gt}})^{\text{scale}}} \end{cases} \quad (3)$$

The shape loss, representing the size difference between the predicted box and the

ground truth box, is shown in Equation (4):

$$\Omega^{shape} = \sum_{t=(w,h)} (1 - e^{w_t})^\theta \quad (4)$$

where, as shown in Equation (5):

$$\begin{cases} \omega_h = ww \times \frac{|h - h^{gt}|}{ww(h, h^{gt})} \\ \omega_w = hh \times \frac{|w - w^{gt}|}{max(w, w^{gt})} \end{cases} \quad (5)$$

The bounding box regression loss formula is given by Equation (6):

$$L_{shape-IoU} = 1 - IoU + L_{distance} + 0.5 \times \Omega^{shape} \quad (6)$$

4. Dataset and Evaluation Metrics

4.1. Experimental Setup and Dataset

The experiments in this paper were conducted on a Windows 11 Pro, 64-bit operating system with 128GB RAM, an Intel Xeon Gold 6226R CPU, and an NVIDIA GeForce RTX 3080 GPU. Python programming language (version 3.8.19) was used for training. The batch size was set to 4, and the number of epochs was set to 300.

To better study the effectiveness of YOLOv11 for object detection in low-light autonomous driving environments, this paper uses the ExDark low-light dataset for validation. The ExDark dataset is specifically designed for computer vision research in low-light conditions. It was released in 2019 by a team from the Faculty of Computer Science and Information Technology, University of Malaya. It contains a total of 7,363 high-quality images captured under low-light conditions, with finely annotated bounding boxes for over 100,000 objects, including 12 categories such as bicycles, cars, pedestrians, and animals. The dataset was divided into a test set (738 images), a training set (5,884 images), and a validation set (741 images). Sample images from the dataset are shown in Figure 3.



Figure 3. Sample Dataset Images

4.2. Evaluation Metrics

This paper uses Average Precision (AP), Precision (P), Recall (R), and mean Average Precision (mAP) as evaluation metrics for object detection. The corresponding calculation formulas are shown in Equations (7) to (10):

$$AP = \int_0^1 PRdR \quad (7)$$

$$P = \frac{TP}{TP + FP} \quad (8)$$

$$R = \frac{TP}{TP + FN} \quad (9)$$

$$mAP = \frac{1}{n} \sum_{i=1}^N AP_i \quad (10)$$

where: TP represents the number of samples where the target is correctly predicted as positive, FP represents the number of samples where the target is incorrectly predicted as positive, FN represents the number of samples where the target is incorrectly predicted as negative; N is the total number of categories.

5. Experimental Results and Analysis

5.1. Ablation Experiment Results and Analysis

To verify the effectiveness of the improvement methods in the IEWS-YOLO model, a series of ablation experiments were designed to evaluate different improvement methods, providing references for future algorithm optimization. The specific ablation experiment results are shown in Table 1. Figure 4 shows the

Precision-Recall curve for the baseline YOLOv11n model, and Figure 5 shows the Precision-Recall curve for the improved IEWS-YOLO model.

Table 1. Ablation Experiment Results Comparison

Model	iEMA	WTConv	Shape-IoU	mAP@0.5/%	mAP@0.5:0.95/%	Model Size/MB	FPS/(f/s)
Yolov11n	×	×	×	53.4	33.4	6.6	42.0
Yolov11n-IE	√	×	×	55.8	35.5	5.7	37.8
Yolov11n-W	×	√	×	57.9	36.2	5.9	40.3
Yolov11n-S	×	×	√	55.1	33.7	6.3	38.6
Yolov11n-IEW	√	√	×	59.4	36.3	5.6	35.7
Yolov11n-IES	√	×	√	54.3	34.2	6.0	50.8
Yolov11n-WS	×	√	√	61.8	39.0	5.9	39.6
IEWS-YOLO	√	√	√	63.0	40.1	5.5	41.3

As can be seen from Table 1, the baseline YOLOv11n model achieves an mAP@0.5 of 53.4% on the dataset. When only the iEMA module is added, the YOLOv11n-IE model improves mAP@0.5 by 2.4% compared to the baseline, indicating that iEMA enhances the extraction capability of blurred features in low-light conditions through multi-scale attention. After adding WTConv, the YOLOv11n-W model improves mAP@0.5 by 4.5% compared to the baseline. When replacing the loss function with Shape-IoU, the YOLOv11n-S model improves mAP@0.5 by 1.7% compared to the baseline, compensating for the shape modeling deficiency of the traditional CIoU loss. YOLOv11n-IEW, YOLOv11n-IES, and YOLOv11n-WS are models with combined improvements. Finally, the detection model integrates all three innovations. Both YOLOv11n-IEW and YOLOv11n-WS show improvements in mAP@0.5 compared to YOLOv11n-IE, YOLOv11n-W, and YOLOv11n-S. Although the YOLOv11n-IES model does not perform outstandingly in mAP@0.5, its FPS improves. The final IEWS-YOLO model achieves the best mAP@0.5 of 63.0%, an improvement of 9.6% over the baseline model. The model size is compressed to 5.5 MB, and FPS reaches 41.3 f/s, validating its effectiveness for object detection in low-light autonomous driving environments.

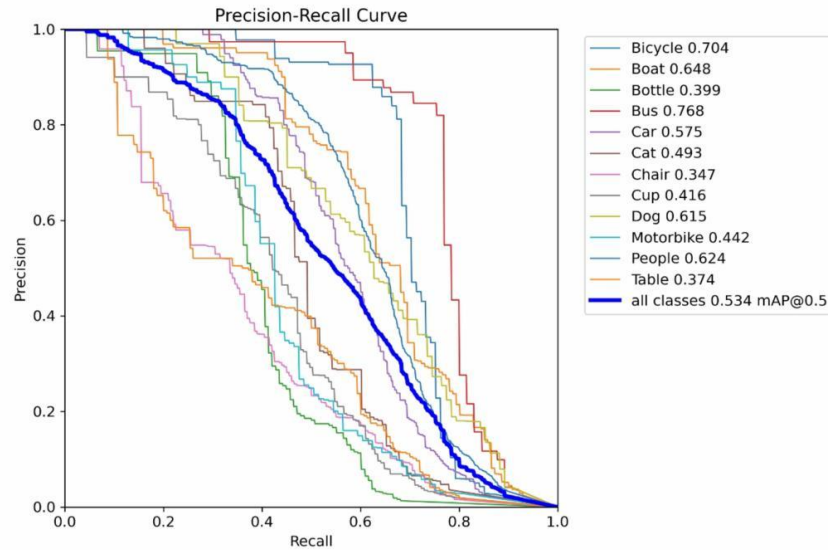


Figure 4. Baseline YOLOv11n Model

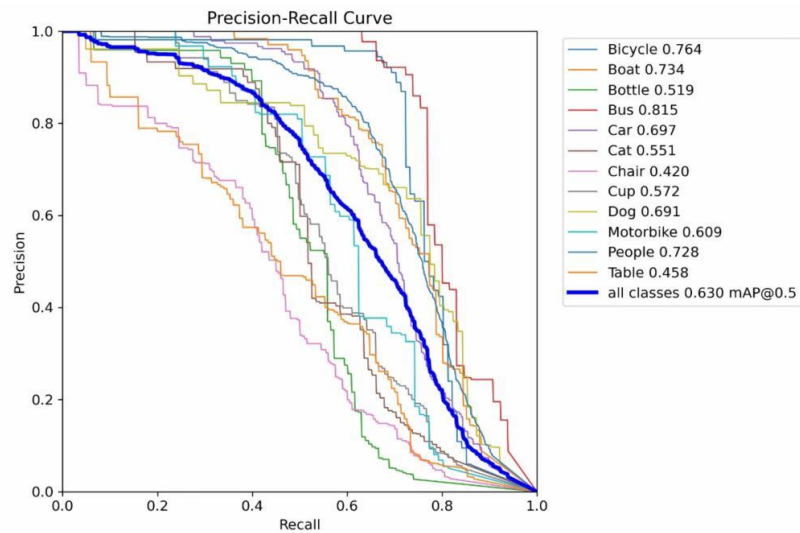


Figure 5. IEWS-YOLO Model

5.2. Different Model Comparison Experiment Results and Analysis

To further verify the effectiveness of the proposed IEWS-YOLO model for object detection in low-light autonomous driving environments, this paper compares it with four different models—YOLOv5n, YOLOv8n, YOLOv10n, and YOLOv11n—trained on the ExDark dataset. The experimental results are shown in Table 2.

Table 2. Different Model Comparison Experiment Results

Model	mAP @0.5/%	mAP @0.5:0.95/%	Model Size/MB	FPS/(f/s)
-------	---------------	--------------------	------------------	-----------

YOLOv5n	41.9	27.8	7.3	35.5
YOLOv8n	52.6	32.3	8.0	37.3
YOLOv10n	44.5	30.7	6.5	33.6
YOLOv11n	53.4	33.4	6.6	42.0
IEWS-YOLO	63.0	40.1	5.5	41.3

From Table 2, it can be seen that on the ExDark dataset, the IEWS-YOLO model achieves the highest mAP@0.5 and mAP@0.5:0.95 among all compared models, at 63.0% and 40.1% respectively. The improved IEWS-YOLO model shows improvements in recall, precision, and mean average precision compared to the original YOLOv11n model in the object detection task. The IEWS-YOLO model exhibits more stable precision changes, whereas the traditional YOLOv11n model shows a more obvious declining trend in precision. This indicates that the improved IEWS-YOLO model has more stable performance.

5.3. Visualization Comparison Experiment

Considering the complex and variable traffic situations in low-light scenes, which involve issues like insufficient brightness, significant noise, and loss of detail, a visualization comparison experiment was conducted on a constructed low-light scene object detection set to more intuitively verify the adaptability of the improved model. The aim is to compare the performance differences in object detection between the traditional model and the improved IEWS-YOLO model.

Representative scenes were selected from the validation set of the ExDark dataset, and the detection results of the original YOLOv11n model and the proposed IEWS-YOLO model were visually compared side by side. The comparison results are shown in Figures 6, 7, and 8. The left side of each figure shows the object detection results from the original YOLOv11n model, and the right side shows the results from the IEWS-YOLO model. Figure 6 shows a comparison of detecting bicycles and cars parked on a road. Figure 7 shows a comparison of detecting parked bicycles and pedestrians on a street where people are resting while riding bicycles. Figure 8 shows a comparison of detecting a kitten and a cup in a room.

From Figure 6, it can be seen that the original model has low object detection accuracy in low-light scenes. Figures 7 and 8 show that the original model suffers from false detections and missed detections in low-light scenes. The improved model reduces missed detections and improves detection accuracy, demonstrating better applicability for object detection in low-light scenes.



Figure 6. Comparison of YOLOv11n and IEWS-YOLO Detection Results



Figure 7. Comparison of YOLOv11n and IEWS-YOLO Detection Results



Figure 8. Comparison of YOLOv11n and IEWS-YOLO Detection Results

6. Conclusion

Aiming at the problems existing in object detection methods for autonomous driving in low-light environments, this paper proposes an improved YOLOv11-based object detection algorithm for autonomous driving in low-light conditions (IEWS-YOLO). Based on the YOLOv11 framework, this algorithm first designs an inverted residual attention mechanism (iEMA) module to enhance the feature quality processed by the neck and detection head, further improving detection performance in low-light autonomous driving. Second, it replaces the standard convolutions in the backbone network with wavelet transform convolution layers (WTConv), improving global feature representation efficiency and noise suppression capability. Finally, it adopts the ShapeIoU loss function to replace the traditional CIoU, optimizing bounding box regression accuracy, making predicted box locations more accurate, and enhancing object localization. Compared to the traditional YOLOv11 model, the IEWS-YOLO model shows significant improvements in various performance aspects. On the ExDark standard low-light dataset, the accuracy, recall, and mean average precision for object detection indicate that the IEWS-YOLO model achieves significant improvements in multiple key metrics. mAP@0.5 improves by 9.6%, reaching 63.0%; mAP@0.5:0.95 improves by 6.7%, reaching 40.1%. The improved model can more accurately perform object detection for autonomous driving in low-light environments.

References

- [1] Zhang W ,Xu H ,Zhu X , et al. RFSC-net: Re-parameterization forward semantic compensation network in low-light environments[J]. Image and Vision Computing,2024,151105271-105271.DOI:10.1016/J.IMAVIS.2024.105271.

- [2] He, Lh., Zhou, Yz., Liu, L. et al. Research on object detection and recognition in remote sensing images based on YOLOv11. *Sci Rep* 15, 14032 (2025).
- [3] Lin, J., Wang, P., Ruan, Y. et al. YOLO11-WLBS: an efficient model for pavement defect detection. *Sci Rep* (2026).
<https://doi.org/10.1038/s41598-026-35743-8>
- [4] Ren, B., Xu, Z., Zhao, J. et al. Complex dark environment-oriented object detection method based on YOLO-AS. *Sci Rep* 15, 21873 (2025).
<https://doi.org/10.1038/s41598-025-07348-0>
- [5] Yang, S. et al. LightingNet: An integrated learning method for low-light image enhancement. *IEEE Trans. Comput. Imaging* 9, 29 – 42 (2023).
- [6] Li, X., Pu, X., Ling, W. et al. YOLO-SAM an end-to-end framework for efficient real time object detection and segmentation. *Sci Rep* 15, 40854 (2025).
<https://doi.org/10.1038/s41598-025-24576-6>
- [7] Han Y, Qi K, Zheng J, et al. Lightweight Cattle Facial Recognition Method Based on Improved YOLOv11. *Smart Agriculture*, 2025, 7(3): 173-184.
<https://doi.org/10.12133/j.smartag.SA202502010>
- [8] Xu, X. et al. Exploring image enhancement for salient object detection in low light images. *ACM Trans. Multimed. Comput. Commun. Appl.* 17, 1 – 19 (2021).
- [9] Wang J ,Yang P ,Liu Y , et al. Research on Improved YOLOv5 for Low-Light Environment Object Detection[J]. *Electronics*,2023,12(14):DOI:10.3390/ELECTRONICS12143089.
- [10] Wang Yin, Wang Lide, Qiu Ji. Real-time Enhancement Algorithm for Orbital Dark Light Environment Based on DenseNet Structure[J]. *Journal of Southwest Jiaotong University*, 2022, 57(6): 1349-1357. doi: 10.3969/j.issn.0258-2724.20210199
- [11] Xue R ,Duan J ,Du Z .MPE-DETR: A multiscale pyramid enhancement network for object detection in low-light images[J].*Image and Vision Computing*,2024,150105202-105202.DOI:10.1016/J.IMAVIS.2024.105202.