

Article

Not peer-reviewed version

Improved Small Object Detection Algorithm CRL-YOLOv5

[Zhiyuan Wang](#), Shujun Men, Yuntian Bai, Yutong Yuan, Jiamin Wang, Kanglei Wang, [Lei Zhang](#)*

Posted Date: 19 August 2024

doi: 10.20944/preprints202408.1218.v1

Keywords: Small Object Detection; Attention Mechanisms; Contextual Information; YOLOv5



Preprints.org is a free multidiscipline platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This is an open access article distributed under the Creative Commons Attribution License which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Article

Improved Small Object Detection Algorithm CRL-YOLOv5

Zhiyuan Wang¹, Shujun Men¹, Yuntian Bai², Yutong Yuan¹, Jiamin Wang¹, Kanglei Wang¹ and Lei Zhang^{1,*}

¹ School of Information Science and Engineering, Yanshan University, Qinhuangdao, 066004, China

² Silesian College of Intelligent Science and Engineering, Yanshan University, Qinhuangdao, 066004, China

* Correspondence: zhangl85@ysu.edu.cn

Abstract: Detecting small objects in images poses significant challenges due to their limited pixel representation and the difficulty in extracting sufficient features, often leading to missed or false detections. To address these challenges and enhance detection accuracy, this paper presents an improved small object detection algorithm, CRL-YOLOv5. The proposed approach integrates the CBAM attention mechanism into the C3 module of the backbone network, which enhances the localization accuracy of small objects. Additionally, the Receptive Field Block (RFB) module is introduced to expand the model's receptive field, thereby fully leveraging contextual information. Furthermore, the network architecture is restructured to include an additional detection layer specifically for small objects, allowing for deeper feature extraction from shallow layers. When tested on the VisDrone2019 small object dataset, CRL-YOLOv5 achieved an mAP₅₀ of 39.2%, representing a 5.4% improvement over the original YOLOv5, effectively boosting the detection precision for small objects in images.

Keywords: small object detection; attention mechanisms; contextual information; YOLOv5

1. Introduction

Small object detection is vital in various applications, including remote sensing image analysis, medical image diagnostics, and intelligent traffic systems. Rapid and accurate detection of small objects, such as unauthorized ships or aircraft in remote sensing images, can significantly enhance national defense security. In the medical field, precise detection of small lesions, like lung nodules, retinal abnormalities, or early-stage tumors, is crucial for early diagnosis and treatment of diseases. Similarly, in intelligent traffic systems, the swift and accurate identification of pedestrians, vehicles, or traffic signs plays a key role in maintaining traffic order and enhancing safety. However, detecting small objects in images is inherently challenging due to factors like limited pixel counts, difficulty in feature extraction, and complex backgrounds. As a result, the task of fast and accurate small object detection has become a focal point of research and innovation in the scientific community in recent years.

Traditional methods for small object detection are often hindered by complex background interference, which adversely affects feature extraction. Deep learning-based object detection algorithms exhibit strong feature extraction capabilities and robustness [1]. These algorithms are broadly categorized into single-stage and two-stage approaches. Two-stage methods, exemplified by R-CNN [2], Fast R-CNN [3], Faster R-CNN [4], and Mask R-CNN [5], typically require initial region proposal generation followed by feature extraction and classification, making them time-consuming compared to single-stage methods. Single-stage detectors like the YOLO [6–8] series and SSD [9] take the entire image as input and directly extract global features, offering higher efficiency and real-time performance, albeit potentially at the cost of lower accuracy and precision. YOLOv5, the fifth iteration of the YOLO series, features an improved network architecture and optimized training strategies, leading to a substantial increase in detection precision. Recent advancements in YOLOv5 adaptations have significantly improved small object detection. Yang et al. introduced KPE-YOLOv5, which incorporates the scSE attention module to enhance the network's focus on small object features,

thereby increasing detection accuracy [10]. Zhang et al. proposed an Adaptive Slicing Approach (ASAH), which optimizes processing speed and detection performance by altering the number of slices rather than their size, effectively reducing redundant computations [11]. Kim et al. developed ECAP-YOLO, which integrates ECA-Net into YOLOv5 to optimize attention on small objects and enhance feature representation [12]. Mahaur et al. improved the Spatial Pyramid Pooling (SPP) module and PANet structure by substituting dilated convolution for pooling operations to further strengthen feature expression [13]. Guo et al. introduced MSFT-YOLO, combining Transformer modules with a multi-scale feature fusion structure (BiFPN), addressing issues of background noise and poor defect detection in industrial settings [14]. Wang et al. presented FE-YOLOv5, which significantly enhances small object detection performance through innovative feature enhancement and spatial awareness technologies [15].

Although the aforementioned methods have improved the accuracy of small object detection to some extent, there remains significant potential for enhancing network feature representation and further leveraging contextual information. This study introduces the following enhancements based on the YOLOv5 model:

1. Integration of CBAM: We incorporate the CBAM attention mechanism [16] into the C3 module of the backbone network to enhance its feature representation capabilities. This modification allows the model to better capture the important parts of an image, facilitating improved accuracy in object detection.

2. Replacement with RFB Module: The SPPF module is replaced by the RFB module [17] to expand the model's receptive field and better utilize contextual information. This change aids in enhancing the model's perception of objects of varying sizes and improves the precision of small object detection.

3. Addition of a Dedicated Small Object Detection Layer: On top of the existing network architecture, a new detection layer focused on small objects is added. This layer makes full use of shallow features for more precise localization of small targets within the image, further enhancing the accuracy of small object detection.

2. Materials and Methods

This research builds upon the YOLOv5 model. Figure 1 illustrates the network structure of YOLOv5, which is divided into three main components: the backbone network, the neck network, and the output section. During object detection, the input image is first resized to a standard dimension of 640×640 pixels. The backbone network then extracts features, generating feature maps of various sizes, which are passed to the neck network for feature fusion. The neck network integrates the structures of FPN [18] and PAN [19], enabling the transfer of semantic information from deeper layers to shallower ones, and the feedback of positional information from shallower to deeper layers. This process results in a feature pyramid that combines both semantic and location information [20]. Finally, the fused feature maps in the neck are processed by convolutional modules to output detections for objects at different scales—large, medium, and small.

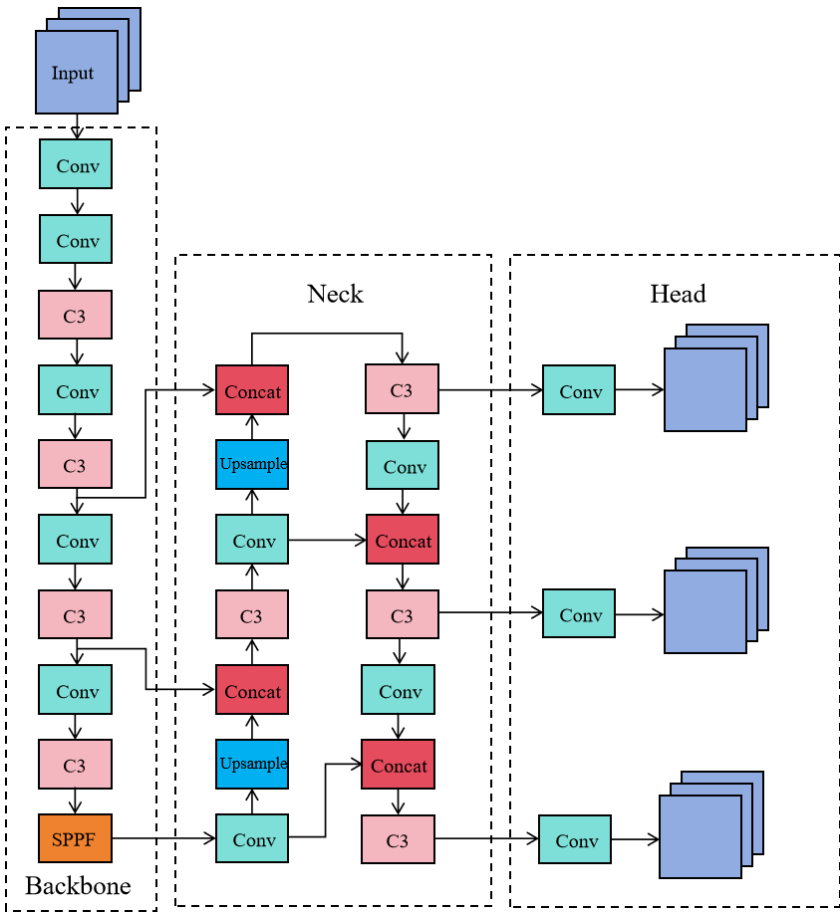


Figure 1. YOLOv5 network structure.

To enhance the network's feature extraction capabilities, the CBAM attention mechanism has been integrated into the C3 module of YOLOv5's backbone network, optimizing the distribution of feature weights. Additionally, the SPPF module in the backbone has been replaced with an RFB module, which broadens the model's receptive field and improves the integration of multi-scale features, thereby enhancing detection precision. Furthermore, a new detection layer specifically designed for small objects has been added to the network structure, maximizing the use of shallow features. These modifications collectively enable the enhanced YOLOv5 to achieve higher detection accuracy. The updated network structure is illustrated in Figure 2.

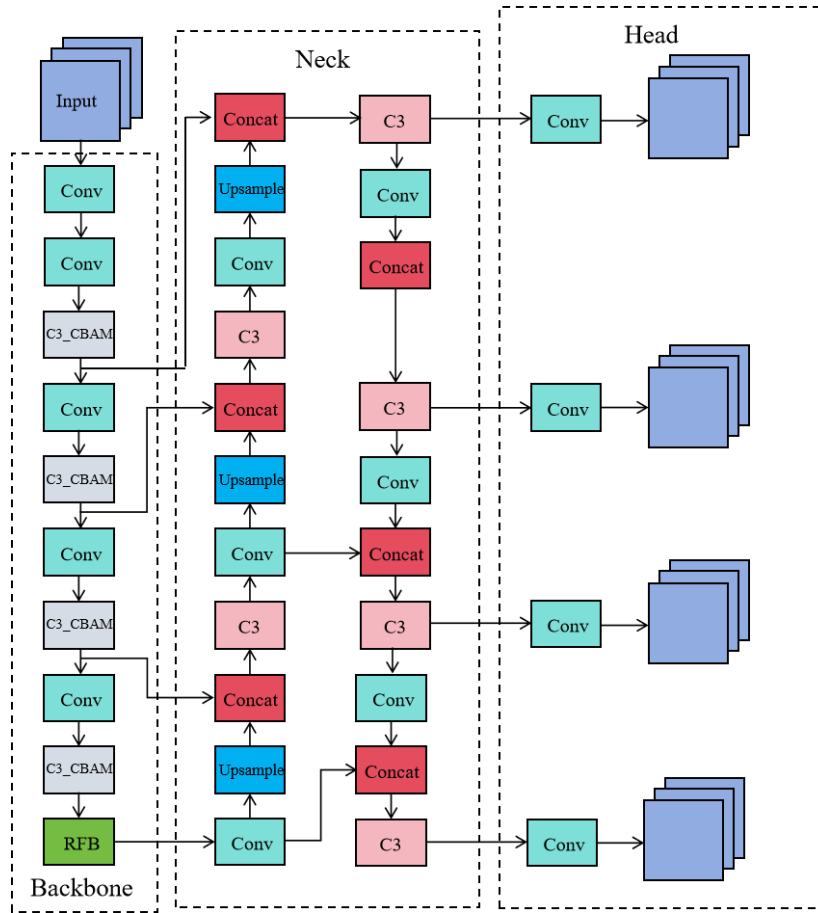


Figure 2. Improved YOLOv5 network structure.

2.1. C3_CBAM Module

The Convolutional Block Attention Module (CBAM) consists of two key components: the Channel Attention Module (CAM) and the Spatial Attention Module (SAM). CAM enhances the model's ability to distinguish important features by adjusting the significance of different channels. In contrast, SAM improves the model's capacity to capture positional information by extracting critical details from the spatial dimension.

Figure 3 is the principle of the Channel Attention Module (CAM). Initially, the input feature map F undergoes both max pooling and average pooling operations, yielding the max-pooled feature map $F_{\max}^c$ and the average-pooled feature map F_{avg}^c , respectively. These feature maps are then fed into a fully connected layer responsible for learning the attention weights for each channel. This process allows the model to adaptively recognize and emphasize the channels that are more important for the current task. Subsequently, the global maximum and average feature vectors are combined to form the final attention weight vector. A sigmoid activation function is employed to ensure that the channel attention weights are scaled between 0 and 1. These weights are then applied to each channel of the original feature map [21]. The resulting attention-weighted channel feature map, $Ac(F)$, is produced by multiplying the derived attention weights with each channel of the original feature map, emphasizing channels that are beneficial for the current task while suppressing irrelevant ones. The processed channel feature map $Ac(F)$ can be expressed as: σ denotes the sigmoid activation function, W_1 and W_2 respectively represent the weight matrices of the first and second hidden layers in the fully connected layer.

$$Ac(F) = \sigma(W_2(W_1(F_{\max}^c)) + W_2(W_1(F_{\text{avg}}^c))) \quad (1)$$

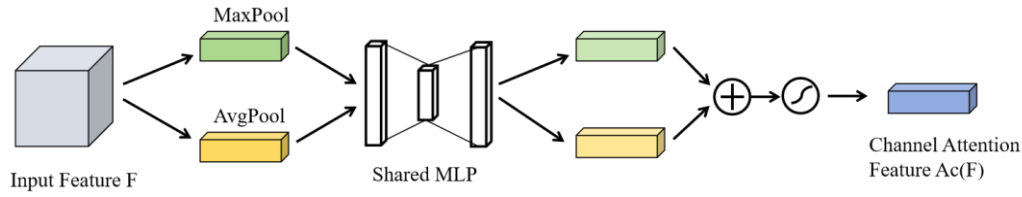


Figure 3. Schematic diagram of CAM principles.

Figure 4 is the operational principle of the Spatial Attention Module (SAM): Initially, the channel-optimized feature map F_c undergoes max pooling and average pooling along the channel dimension, followed by concatenation of the two pooled feature maps. This forms a composite feature map that encapsulates contextual information at different scales. The composite feature map is then further processed through a convolutional layer to produce the spatial feature map $As(F)$. Similar to the Channel Attention Module, the creation of the spatial feature map also utilizes a sigmoid activation function to ensure that attention weights are scaled between 0 and 1. This processing highlights significant areas within the image and reduces the impact of less important regions. The spatial feature map $As(F)$ can be expressed as: σ represents the sigmoid activation function and $*$ denotes the convolution operation.

$$As(F) = \sigma(*([MaxPool(F); AvgPool(F)])) = \sigma(*([F_{max}; F_{avg}])) \quad (2)$$

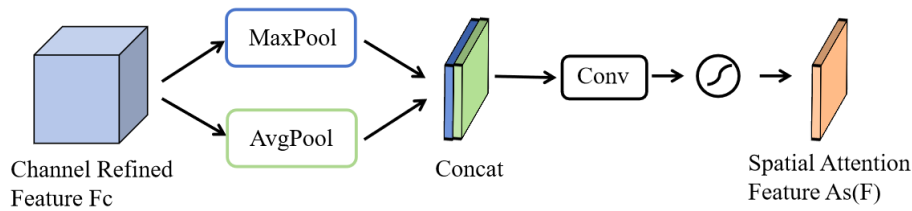


Figure 4. Schematic diagram of SAM principles.

CBAM combines the output features from CAM and SAM by weighting them to produce the final attention-enhanced features. These features are then used as input for subsequent network layers, preserving key information while suppressing noise and irrelevant details. The operational principle of CBAM is illustrated in Figure 5 and can be mathematically represented as: F_c represents the feature map optimized by the Channel Attention Module, and F_s represents the feature map optimized by the Spatial Attention Module.

$$F_c = Ac(F) \otimes F \quad (3)$$

$$F_s = As(F_c) \otimes F_c \quad (4)$$

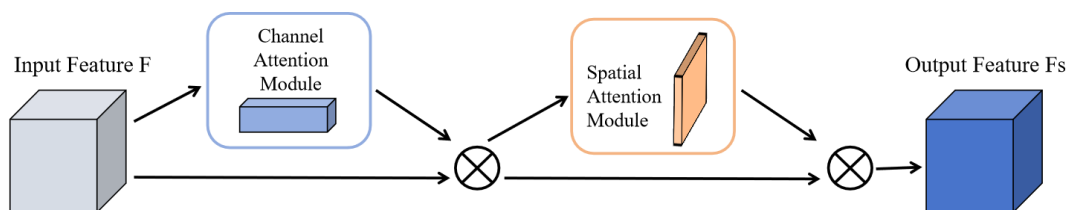


Figure 5. Schematic diagram of CBAM module principles.

To augment the feature extraction capabilities of the network, this study integrates the CBAM attention mechanism with the C3 module of the YOLOv5 backbone, forming the feature enhancement module C3_CBAM. The structure of the C3 module, as shown in Figure 6, consists of two residual blocks, each extracting rich features through multi-layer convolution. In this configuration, the Shortcut option is set to true by default, allowing the input feature map to be added to the output feature maps processed by two convolutional layers, thereby forming residual connections. Furthermore, the Add operation performs a simple pixel-level addition to increase the information content. The Concat operation, on the other hand, concatenates the output feature maps of the residual blocks with the feature maps from the skip connections along the channel dimension, thus creating new feature maps with richer content.

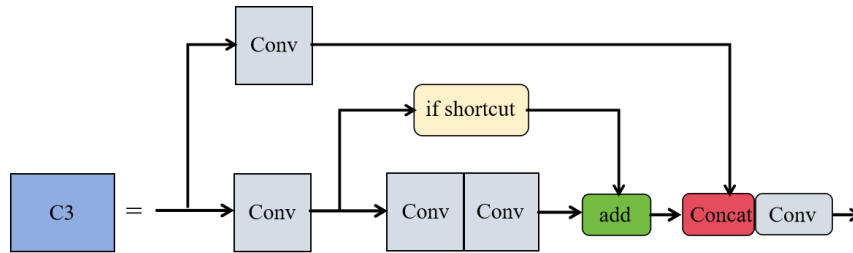


Figure 6. C3 module structure diagram.

The CBAM attention mechanism is integrated into the residual blocks of the C3 module in YOLOv5, forming a new feature extraction module called C3_CBAM, as illustrated in Figure 7. The incorporation of CBAM is intended to improve the model's ability to recognize features of small objects, particularly by enhancing the feature channels that are crucial for small object identification. This integration not only increases the model's sensitivity to small object features but also sharpens its focus on critical spatial information while reducing the impact of irrelevant background details. As a result, this enhancement leads to more precise and reliable localization and identification of small objects.

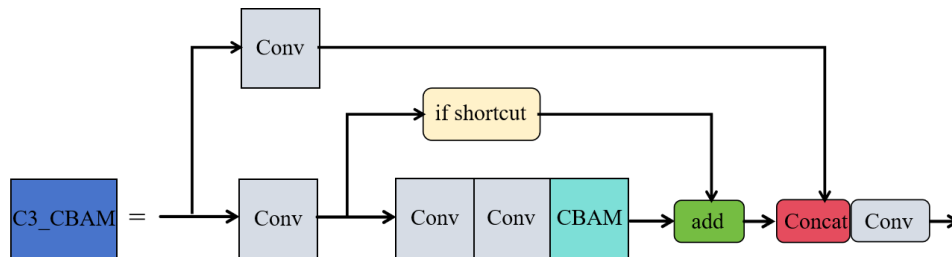


Figure 7. C3_CBAM structure diagram.

2.2. RFB Module

The SPPF (Spatial Pyramid Pooling-Fast) module in the YOLOv5 network structure helps to expand the model's receptive field and gather global information, but it performs moderately in handling precise multi-scale object detection tasks. Inspired by the human visual receptive field, the design of the RFB (Receptive Field Block) module is apt for processing complex and varied visual information. Replacing the SPPF module with the RFB module further enlarges the model's receptive field and enhances the acquisition of comprehensive contextual information. The structure of the RFB module, as shown in Figure 8, utilizes a combination of multi-branch convolution and dilated convolution. The features processed by convolution are concatenated and then fused through a 1×1 convolution, producing the final output feature map. By employing convolutional kernels of various sizes and dilated convolutions, the RFB module effectively captures multi-scale features, significantly boosting the model's capability to detect objects with substantial differences in shape and size. The use of dilated convolution not only enlarges the receptive field without additional computational cost

but also maintains the resolution of the feature map while improving the efficiency of processing large area contextual information. This design enhances the model's discriminative power and robustness, making it more efficient and accurate in facing diverse detection tasks.

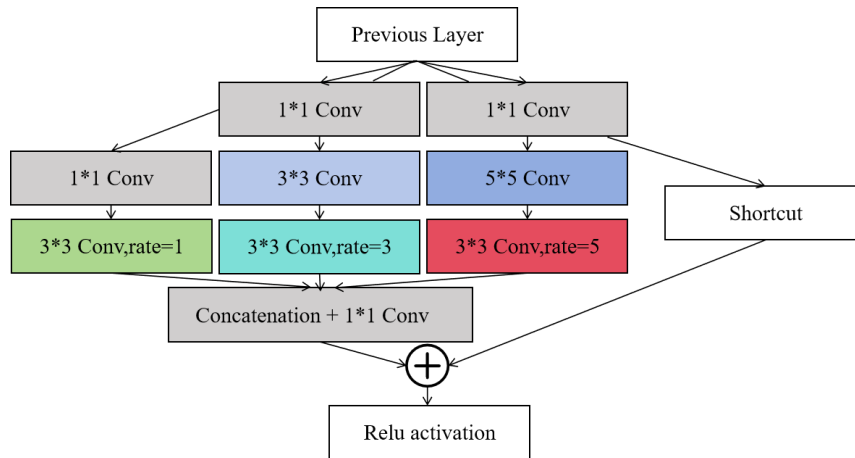


Figure 8. RFB module structure diagram.

2.3. Adding a Dedicated Small Object Detection Layer

The YOLOv5 network incorporates an FPN (Feature Pyramid Network) and PAN (Path Aggregation Network) structure, where the FPN is responsible for conveying deep semantic information to the shallower layers, and the PAN transfers shallow positional information to the deeper layers. This architecture effectively merges features from different levels, enhancing the network's ability to detect objects. The network's three output heads are designed to detect small, medium, and large-scale objects respectively. However, as the resolution of feature maps decreases with network depth, fine details of small objects are often lost, which can negatively impact their accurate localization and identification. To address this issue, an additional layer dedicated to small object detection is introduced, processing higher resolution feature maps. This added layer enables the model to detect objects across various scales, enhancing its multi-scale detection capability and allowing it to handle targets of different sizes simultaneously. As a result, the network achieves more accurate recognition and localization in complex scenes. The input to this new small object detection layer is derived from shallower feature layers of the backbone network, and it outputs smaller-sized detection heads. This design captures more detail, further improving the performance of small object detection.

3. Experiments and Analysis

3.1. Experimental Data and Environment

The dataset selected for the experiments in this study is VisDrone2019, which is shared by the AISKEYE team at Tianjin University. It consists of 10,209 images, with 6,471 images used for training, 3,190 for testing, and 548 for validation. The dataset includes labels for ten categories: pedestrians, people, bicycles, cars, vans, trucks, tricycles, awning-tricycles, buses, and motorcycles. Figure 9 shows the distribution of all label sizes within the training set, with the x-axis representing label width and the y-axis representing label height. As observed, the labels are densely packed in the lower left corner, indicating a prevalence of small objects, making this dataset well-suited for this study.

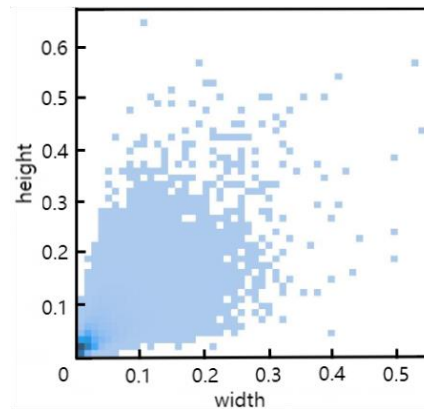


Figure 9. Distribution of all label sizes in the training set.

The experiments were conducted on a server running Ubuntu 22.04 operating system, utilizing Python 3.10 as the programming language. The deep learning framework used was PyTorch 2.1.0, with CUDA 12.1 as the parallel computing platform. The hardware included an RTX 4090 GPU. All experiments were performed with consistent hyperparameters across tests. Specifically, the training was carried out for 150 epochs with a batch size of 16, an initial learning rate of 0.01, and an image resolution of 640×640 pixels.

3.2. Evaluation Metrics

The performance of the model in this experiment was evaluated using several metrics: Precision (P), Recall (R), mAP50, mAP50-95, the number of parameters, GFLOPs, and FPS. Precision (P) measures the accuracy of the model, defined as the ratio of correctly predicted positive observations to the total predicted as positive. Recall (R) assesses the model's ability to capture all relevant cases, defined as the ratio of correctly identified positive cases to all actual positive cases. AP represents the area under the Precision-Recall (PR) curve, and the mean Average Precision (mAP) is calculated by averaging the AP across all categories. mAP50 refers to the mean precision at an Intersection over Union (IoU) threshold of 0.5, while mAP50-95 refers to the mean precision across an IoU threshold ranging from 0.5 to 0.95. The number of parameters indicates the total parameters used by the model. GFLOPs measure the billions of floating-point operations the processor can perform per second, evaluating the complexity and computational demand of the model or algorithm. FPS, frames per second, is used to gauge the speed of the model in processing images.

3.3. Ablation Study and Analysis of Algorithm Effectiveness

To demonstrate the effectiveness of the three approaches utilized, ablation studies were conducted under identical conditions to investigate the impact of each method on algorithm accuracy and computational complexity. The results of the ablation experiments are presented in Table 1. In the table, 'A' represents the use of C3_CBAM, 'B' denotes the addition of a small object detection layer, and 'C' indicates the replacement of the SPPF module with the RFB module.

As shown in Table 1, the use of each improvement method independently resulted in enhanced model accuracy, precision, and recall compared to the baseline YOLOv5s model. Notably, the addition of a small object detection layer showed the most significant improvement, with a 5.2% increase in mAP50, underscoring the importance of accurate positional information for detecting small objects. This method effectively utilizes shallow positional information to enhance the localization accuracy of small objects. When combining multiple methods, the highest detection accuracy was achieved using C3_CBAM, the small object detection layer, and the RFB module simultaneously. Compared to the baseline YOLOv5s model, this combination led to a 5.4% increase in mAP50, a 3.6% increase in mAP50-95, a 5.1% increase in precision, and a 4.2% increase in recall,

with the number of parameters increasing from 7.05M to 7.87M. These results indicate that a slight increase in the number of parameters effectively enhances detection accuracy.

Table 1. Results of ablation experiments.

YOLOv5s	A	B	C	mAP50/%	mAP50-95/%	P/%	R/%	Params/M	GFLOPs
√				33.8	18.7	45.4	34.5	7.05	16.0
√	√			34.4	19.0	48.8	34.6	6.73	15.0
√		√		39.0	22.1	49.8	39.1	7.19	18.9
√			√	34.4	19.1	47.6	34.9	7.71	16.6
√	√	√		39.0	22.1	50.3	38.5	7.21	19.0
√	√	√	√	39.2	22.3	50.5	38.7	7.87	19.5

Figure 10 is the comparative test results of baseline YOLOv5 and enhanced CRL-YOLOv5 in various scenes. The test images were selected from the VisDrone2019 test set and included scenes such as parking lots, squares, and roads, all of which contain a rich and dense population of small objects, meeting the requirements of the detection task. Figure 10(a) displays the detection results of the baseline YOLOv5, while Figure 10(b) shows the results from the enhanced CRL-YOLOv5. It is evident that YOLOv5 missed many targets in these scenarios, whereas the improved algorithm successfully detected the majority of small objects, significantly reducing the number of misses. Additionally, in the first and third images, YOLOv5 incorrectly identified buildings as trucks, and in the third image, a van in the lower-left corner was mistakenly recognized as a truck. The enhanced algorithm did not exhibit these issues, effectively lowering the false positive rate. In densely populated areas within these images, YOLOv5 failed to detect most targets, and the detection performance was even poorer in the darker environment of the second image. The improved CRL-YOLOv5 algorithm demonstrated substantial enhancements in detection accuracy in both densely populated areas and under low-light conditions. In conclusion, the upgraded CRL-YOLOv5 algorithm exhibited superior performance in detecting small objects, significantly reducing both misses and false detections, making it suitable for practical detection scenarios where targets are densely packed or lighting is insufficient

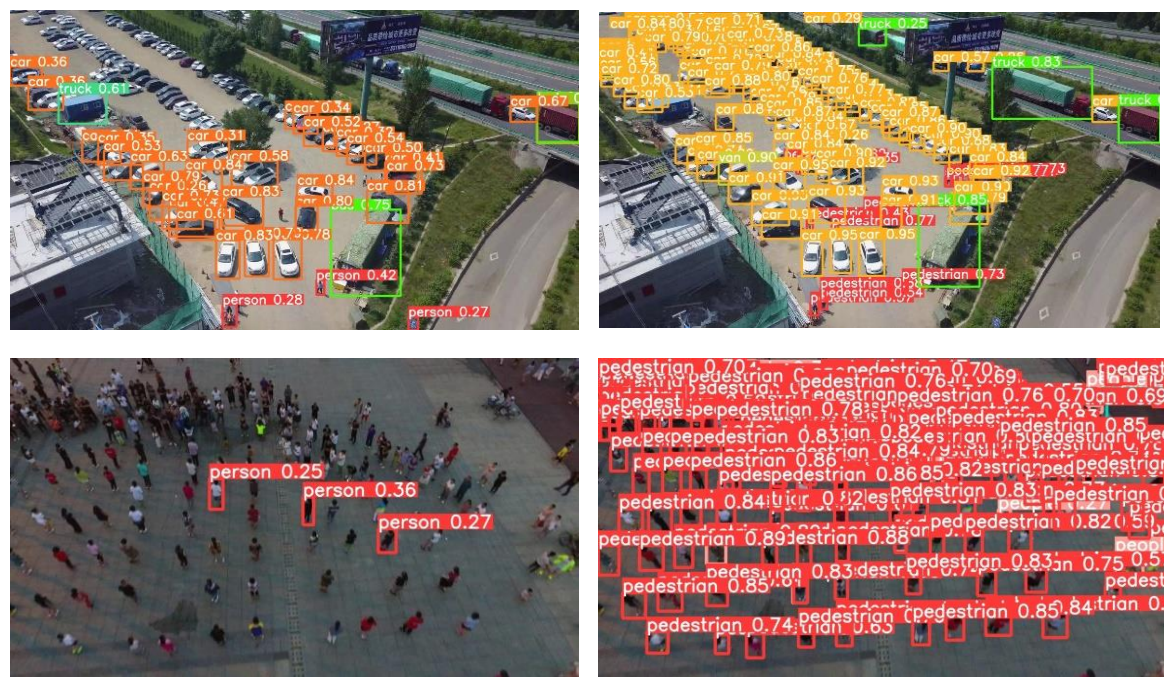




Figure 10. Comparison of test results.

3.4. Comparative Experiments

Finally, the enhanced CRL-YOLOv5 algorithm was compared with major algorithms in the field of object detection, using detection accuracy and speed as evaluation metrics. The results of these comparative experiments are shown in Table 2. The table indicates that, compared to previous iterations of YOLOv5 and other typical models, our algorithm achieves higher accuracy, with an mAP50 that is 8.5% higher than Faster-RCNN and 6.1% higher than YOLOv4. Although the FPS (frames per second) is 27, which is a decrease compared to YOLOv4 and YOLOv5, it still maintains real-time detection capability. Compared to TPH-YOLOv5, a model specifically designed for small object detection, our model shows a 1.9% higher mAP50 and also exhibits improved FPS. Furthermore, it also outperforms the latest models in the YOLO series, with an mAP50 that is 0.3% higher than YOLOv8. Overall, the results demonstrate that the CRL-YOLOv5 algorithm possesses significant advantages in handling small object detection tasks.

Table 2. Comparison experiment results.

Algorithm	mAP50/%	mAP50-95/%	FPS
Faster-RCNN	30.7	16.1	17
YOLOv4 [22]	33.1	18.1	36
YOLOv5s	33.8	18.7	47
TPH-YOLOv5 [23]	37.3	20.2	15
YOLOX [24]	35.2	19.4	77
YOLOv7 [25]	36.7	20.8	89
YOLOv8	38.9	22.1	86
CRL-YOLOv5	39.2	22.3	27

4. Conclusions

Addressing the challenges of low accuracy, missed detections, and false positives in small object detection in images, this paper proposes an improved algorithm, CRL-YOLOv5. Firstly, the algorithm integrates the CBAM (Convolutional Block Attention Module) within the residual structures of the C3 module in the backbone network to redistribute feature weights, enhancing the localization and detection capabilities for small objects. Secondly, a new small object detection layer is added to augment multi-scale feature fusion, aiding in the precise localization of small objects. Lastly, the substitution of the SPPF (Spatial Pyramid Pooling-Fast) module with an RFB (Receptive Field Block) module expands the model’s receptive field, improving the network’s ability to capture features of small objects. Experiments conducted using the VisDrone2019 dataset demonstrate that the enhanced model achieves improved detection accuracy, exhibits effective performance on small objects, and is suitable for everyday small object detection scenarios.

Author Contributions: Conceptualization, Z.W.; methodology, Z.W.; software, Z.W.; validation, Z.W.; formal analysis, L.Z.; investigation, Z.W.; resources, L.Z.; writing—original draft preparation, Z.W.; writing—review

and editing, S.M., Y.B., Y.Y., J.W., K.W. and L.Z.; visualization, Z.W.; supervision, L.Z.; project administration, Z.W. All authors have read and agreed to the published version of the manuscript.

Funding: This work was partially supported by Open Fund of State Key Laboratory of Information Photonics and Optical Communications (Beijing University of Posts and Telecommunications), P. R. China under Grant IPOC2021B06.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: The original data presented in this study are openly available at <https://github.com/VisDrone/VisDrone-Dataset>, (accessed on 24 September 2023).

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Su J, Qin Y C, Jia Z, et al. Small Object Detection Algorithm Based on ATO-YOLO[J]. *Journal of Computer Engineering & Applications*, 2024, 60(6):68-77.
2. Girshick R, Donahue J, Darrell T, et al. Rich feature hierarchies for accurate object detection and semantic segmentation[C]//*Proceedings of the IEEE conference on computer vision and pattern recognition*. 2014: 580-587.
3. Girshick R. Fast r-cnn[C]//*Proceedings of the IEEE international conference on computer vision*. 2015: 1440-1448.
4. Ren S, He K, Girshick R, et al. Faster r-cnn: Towards real-time object detection with region proposal networks[J]. *Advances in neural information processing systems*, 2015, 28.
5. He K, Gkioxari G, Dollár P, et al. Mask r-cnn[C]//*Proceedings of the IEEE international conference on computer vision*. 2017: 2961-2969.
6. Redmon J, Divvala S, Girshick R, et al. You only look once: Unified, real-time object detection[C]//*Proceedings of the IEEE conference on computer vision and pattern recognition*. 2016: 779-788.
7. Redmon J, Farhadi A. YOLO9000: better, faster, stronger[C]//*Proceedings of the IEEE conference on computer vision and pattern recognition*. 2017: 7263-7271.
8. Redmon J, Farhadi A. Yolov3: An incremental improvement[J]. *arXiv preprint arXiv:1804.02767*, 2018.
9. Liu W, Anguelov D, Erhan D, et al. Ssd: Single shot multibox detector[C]//*Computer Vision—ECCV 2016: 14th European Conference, Amsterdam, The Netherlands, October 11–14, 2016, Proceedings, Part I 14*. Springer International Publishing, 2016: 21-37.
10. Yang R, Li W, Shang X, et al. KPE-YOLOv5: an improved small target detection algorithm based on YOLOv5[J]. *Electronics*, 2023, 12(4): 817.
11. Zhang H, Hao C, Song W, et al. Adaptive slicing-aided hyper inference for small object detection in high-resolution remote sensing images[J]. *Remote Sensing*, 2023, 15(5): 1249.
12. Kim M, Jeong J, Kim S. ECAP-YOLO: Efficient channel attention pyramid YOLO for small object detection in aerial image[J]. *Remote Sensing*, 2021, 13(23): 4851.
13. Mahaur B, Mishra K K. Small-object detection based on YOLOv5 in autonomous driving systems[J]. *Pattern Recognition Letters*, 2023, 168: 115-122.
14. Guo Z, Wang C, Yang G, et al. Msft-yolo: Improved yolov5 based on transformer for detecting defects of steel surface[J]. *Sensors*, 2022, 22(9): 3467.
15. Wang M, Yang W, Wang L, et al. FE-YOLOv5: Feature enhancement network based on YOLOv5 for small object detection[J]. *Journal of Visual Communication and Image Representation*, 2023, 90: 103752.
16. Woo S, Park J, Lee J Y, et al. Cbam: Convolutional block attention module[C]//*Proceedings of the European conference on computer vision (ECCV)*. 2018: 3-19.
17. Liu S, Huang D. Receptive field block net for accurate and fast object detection[C]//*Proceedings of the European conference on computer vision (ECCV)*. 2018: 385-400.
18. Lin T Y, Dollár P, Girshick R, et al. Feature pyramid networks for object detection[C]//*Proceedings of the IEEE conference on computer vision and pattern recognition*. 2017: 2117-2125.
19. Liu S, Qi L, Qin H, et al. Path aggregation network for instance segmentation[C]//*Proceedings of the IEEE conference on computer vision and pattern recognition*. 2018: 8759-8768.
20. Pang N Y, Du A Y Application research on surface defect detection based on YOLOv5s_Attention [J]. *Modern Electronic Technology*, 2023, 46(03): 39-46.
21. Feng Z Q, Xie Z J, Bao Z W, et al. Real-time dense small target detection algorithm for UAVs based on improved YOLOv5 [J]. *Acta Aeronautica et Astronautica Sinica*, 2023, 44(07): 251-265.
22. Bochkovskiy A, Wang C Y, Liao H Y M. Yolov4: Optimal speed and accuracy of object detection[J]. *arXiv preprint arXiv:2004.10934*, 2020.

23. Zhu X, Lyu S, Wang X, et al. TPH-YOLOv5: Improved YOLOv5 based on transformer prediction head for object detection on drone-captured scenarios[C]//Proceedings of the IEEE/CVF international conference on computer vision. 2021: 2778-2788.
24. Ge Z, Liu S, Wang F, et al. Yolox: Exceeding yolo series in 2021[J]. arXiv preprint arXiv:2107.08430, 2021.
25. Wang C Y, Bochkovskiy A, Liao H Y M. YOLOv7: Trainable bag-of-freebies sets new state-of-the-art for real-time object detectors[C]//Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. 2023: 7464-7475.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.