

基于改进 YOLOv8n 的电动车头盔佩戴检测

唐皓阳¹ 肖小玲¹

TANG Haoyang XIAO Xiaoling

摘要

准确高效地检测电动车骑行者是否佩戴头盔，对于减少交通事故伤亡具有重要意义。文章提出一种改进的 YOLOv8n 模型，用于电动车骑行者头盔佩戴检测，有效提升了检测性能。首先，将原模型中的 SPPF 模块替换为 SPPELAN，增强特征提取能力；其次，采用 DWRSeg 优化 C2f 模块，增强了模型对头盔边界的定位能力；最后，采用融合 Focal 思想的 DIoU 损失函数，优化了边界框回归精度，并提升了整体检测性能。实验结果表明，该模型在公共数据集上取得了显著的性能提升，平均精度均值 mAP@0.5 提升了 3.4%，mAP@0.5:0.95 提升了 4.4%，证明了其有效性。

关键词

YOLOv8; 头盔检测; SPPELAN; C2f DWRSeg; Focal DIoU

doi: 10.3969/j.issn.1672-9528.2024.12.028

0 引言

电动车以其便捷性成为居民出行的重要交通工具之一，然而，部分骑行人员安全意识不足，未佩戴头盔现象普遍存在，严重威胁自身及他人生命安全。研究表明，佩戴头盔可有效降低交通事故中头部损伤风险，因此，加强头盔佩戴监管是提升交通安全的关键举措。

在实际交通环境中，天气、光照及雾霾等复杂因素均会影响骑行头盔佩戴检测的准确性。科研工作者致力于优化检测技术以增强其适应性和鲁棒性。Rubaiyat 等人^[1]融合 HOG、频域信息及颜色与 CHT 特征，实现了安全帽有效检测，但算法复杂度与计算成本高，影响实时性。Macalisang 等人^[2]通过微调 YOLOv3 参数提升头盔检测精度，然而数据集规模小限制了模型泛化能力。Bouhayane 等人^[3]使用 Swin Transformer 作为主干网络，使用 FPN 架构处理多尺度特征，最后使用 cascade-rcnn 进行检测并输出，但仍存在数据集局限性和模型复杂度增加等问题。JIA 等人^[4]融合注意力与 Soft-NMS 改进 YOLOv5，提升头盔检测性能与模型泛化，但数据集多样性和极端条件表现受限。陈扬等人^[5]在 YOLOv5 基础上添加跳层连接与注意力机制，用 GIoU 损失函数提升精度效率，但由于是自制数据集，多样性和均衡性不足。吴昊祺等人^[6]通过引入 SE 注意力机制和小样本处理技术优化 YOLOv7 主干网络，并引入边界框距离信息改进 IoU 损失函数，提升了小目标

别精度, 但召回率和模型体积有所下降。杨琰钱等人^[7]基于YOLOv8, 采用SPDConv优化卷积与下采样, 融合CG与C2f模块, 并引入组卷积改进head层, 提升了低分辨率图像及小目标检测精度, 但在强光等复杂场景下检测有误差。为进一步提高头盔佩戴检测的准确性和效率, 本文提出一种改进YOLOv8n的头盔检测算法, 实验表明, 在准确率、召回率和mAP上均优于现有方法。

1 YOLOv8 网络模型

YOLO (you only look once) 系列是单阶段目标检测算法, YOLOv8 由 Ultralytics 公司推出, 其网络结构如图 1 所示, 性能领先, 既保持了实时性又提升了精度, 其架构由骨干网络、颈部网络和检测头三大模块组成。

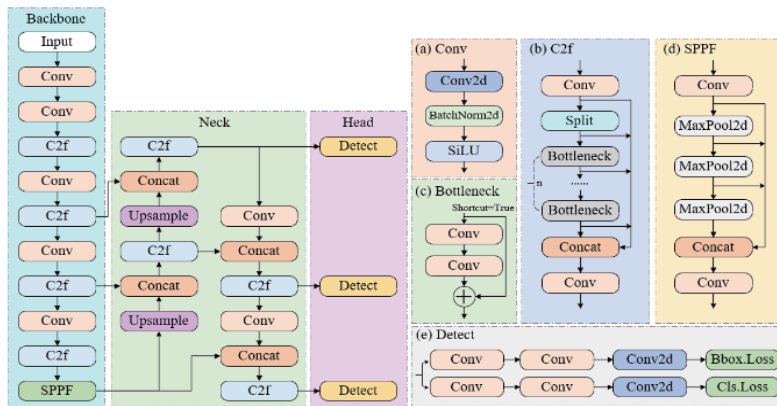


图 1 YOLOv8 网络结构

在骨干网络中，YOLOv8 基于 DarkNet-53 进行了优化，C2f 模块在保持模型轻量化的同时提供更丰富的梯度流信息，

1. 长江大学 湖北荆州 434100

并引入了 ELAN 的思想以增强模块性能。颈部网络的设计借鉴了 PAN-FPN 的设计理念，删除 1×1 卷积层并引入 C2f 模块，实现了多尺度特征的深度融合与高效利用。检测头部分，YOLOv8 引入了解耦头架构，将分类与检测任务分离，有效减少了任务间的相互干扰。此外，采用 Anchor-Free 机制替代传统的 Anchor-Based 方法，进一步提升了检测的灵活性与准确性，使模型能够更精准地捕捉各种尺寸与形状的目标。

2 改进的 YOLOv8 模型

由于 YOLOv8n 模型轻量及高效的特征契合本文的研究需求, 因此本文使用 YOLOv8n 作为基准架构, 在此基础上进行优化, 其网络架构如图 2。本文的改进主要在以下几个方面。

(1) 用 SPPELAN 模块替换 SPPF 模块, 使模型能够更精准地聚焦于图像的关键区域, 减少无关信息的干扰。

(2) 使用 DWRSeg 对 C2f 模块进行优化, 替换颈部网络中的 C2f 模块, 进一步优化特征的层次结构。

(3) 采用融合 Focal 思想的 DIoU 损失函数替换原损失函数, 更有效地解决样本不平衡问题。

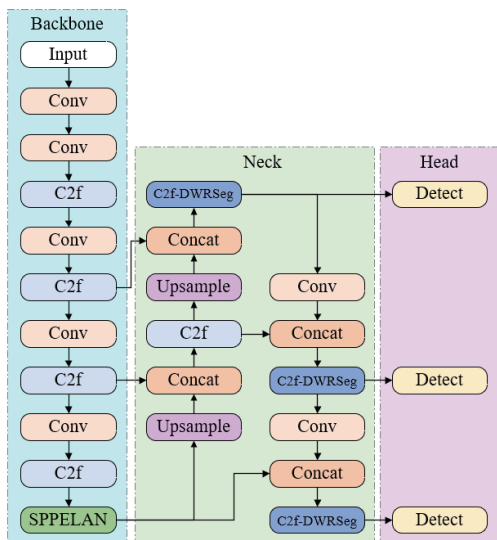


图 2 改进后的 YOLOv8 网络结构

2.1 SPPELAN 模块

SPPELAN 网络结构如图 3，是 YOLOv9 提出的模块，融合了空间金字塔池化（SPP）的层次化信息处理能力与轻量级聚合网络的效率优势。通过多尺度池化层，全面捕获特征图信息，避免信息局限与丢失，提升图像理解能力。SPPELAN 模块内的多层池化递进，聚焦于不同的感受野。每一层都独立地提取最为显著的特征点，并随着层数的增加，这些特征点逐渐覆盖了从局部细节到全局结构的广泛尺度范围。不仅丰富了特征表示的维度，还增强了模型对不同尺度

目标的检测能力。

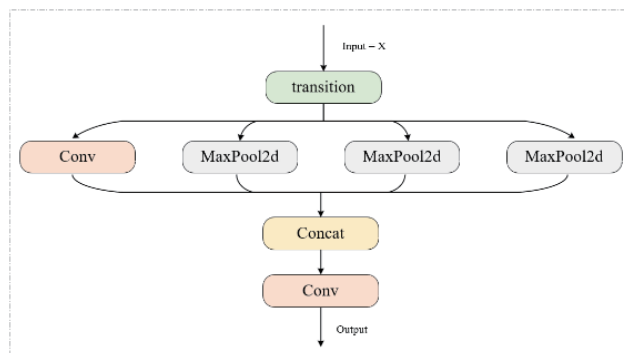


图 3 SPPELAN 结构图

此外，SPELAN 模块还通过后续的卷积层将不同尺度下的特征信息进行高效融合，不仅保留了各尺度特征的独特信息，还促进了特征之间的互补与协作，从而生成了一个更为全面的特征表示。进一步提升了模型在复杂场景下的检测精度和泛化能力。

2.2 改进 C2f 模块

DWRSeg^[8] (dilated-wise residual segmentation) 主要针对小目标检测, 其核心策略是通过整合扩张式残差模块与简易反转残差模块, 增强模型在复杂场景下的检测性能。DWRSeg 的核心思想在于其双阶段特征提取框架: 区域残差化与语义残差化。第一阶段使用 3×3 卷积核, 结合批量归一化 (BN) 和 ReLU 激活函数生成基础特征图。第二阶段用深度可分离卷积优化各区域特征, 减少了无关信息的干扰, 保留了关键语义, 使特征表达更加精准高效。

DWRSeg 的核心组件之一 DWR 模块（如图 4）是性能提升的关键。

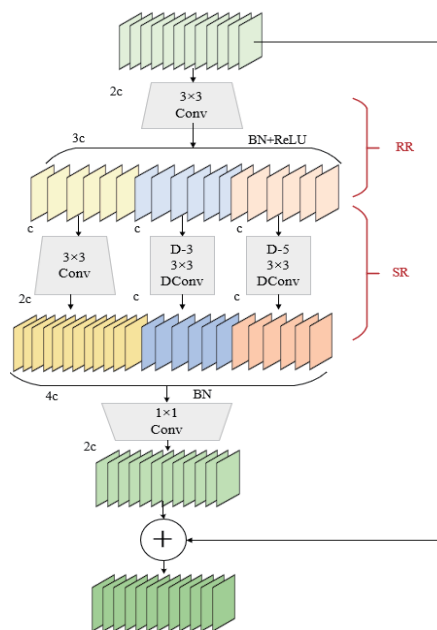


图 4 DRW 模块结构图

该模块采用多尺度扩张卷积并行捕捉不同尺度的上下文信息。随后,通过 1×1 卷积融合跨尺度信息。DWR模块还通过残差连接将融合后的特征与原始输入直接相连,不仅促进了信息的流通,同时还增强了模型的鲁棒性和泛化能力,使得模型在面对复杂多变的检测任务时能够展现出更加稳定的性能。

本文将DWRSeg融入Bottleneck结构中以优化C2f模块(如图5),其双阶段特征提取与高效的DWR模块显著提升了目标检测性能。

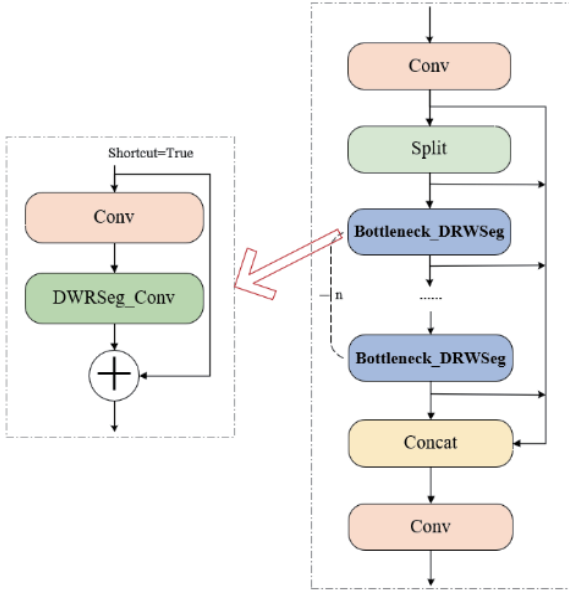


图5 C2f_DRWSeg 结构图

2.3 Focal_DIoU 损失函数

YOLOv8所使用的损失函数CIoU^[9],如式(1),虽然能够有效评估预测边界框与真实边界框的重叠程度,但在处理复杂场景时存在局限性,例如低重叠度情况下的梯度消失问题。

$$\text{CIoU} = 1 - (\text{IoU} + \frac{\rho^2(b, b^{\text{gt}})}{c^2} + \alpha\theta) \quad (1)$$

式中:IoU是预测边界框 b 和真实边界框 b^{gt} 的交并比; $\rho^2(b, b^{\text{gt}})$ 是预测边界框和真实边界框中心点之间的欧氏距离的平方; c 是能够同时包含预测边界框和真实边界框的最小闭合框的对角线长度; α 是用于平衡中心点距离和宽高比一致性项权重的参数,定义为:

$$\alpha = \frac{\theta}{(1 - \text{IoU}) + \theta} \quad (2)$$

式中: θ 是衡量宽高比一致性的参数,具体形式为:

$$\theta = \frac{4}{\pi^2} (\arctan \frac{\omega^{\text{gt}}}{h^{\text{gt}}} - \arctan \frac{\omega}{h})^2 \quad (3)$$

式中: ω 和 h 是预测边界框的宽和高; ω^{gt} 和 h^{gt} 是真实边界框的宽和高。

为了克服上述问题,本文采用DIoU^[10]损失函数替代CIoU,并结合Focal Loss^[11-12]的思想,提出了Focal_DIoU损失函数。DIoU、Focal Loss的计算公式分别为:

$$\text{DIoU} = 1 - (\text{IoU} - \frac{\rho^2(b, b^{\text{gt}})}{c^2}) \quad (4)$$

$$\text{FL}(p, t) = -(1 - p_t)^\gamma \log p_t \quad (5)$$

式中: γ 是Focal Loss引入的可调超参数,用于调节易分类样本的权重,降低其权重,使难以分类的样本得到更多关注; p' 是模型对样本的预测概率。

DIoU有效缓解了IoU在低重叠度情况下的梯度消失问题,使模型能够更加精准地调整边界框的位置。同时,借鉴Focal Loss的思想,动态调整不同样本在损失计算中的权重,确保模型在训练过程中更加聚焦于难以分类的困难样本。因此,采用融合Focal思想的DIoU作为本文算法的损失函数,提升模型在复杂场景下的学习效率 and 检测性能。

3 实验结果与分析

3.1 实验环境与参数设置

本文实验环境为:操作系统Ubuntu 20.04, GPU NVIDIA GeForce RTX 3080 Ti, 深度学习框架PyTorch 2.1.2, CUDA版本12.1。模型训练使用SGD优化器,批处理大小设置为8,训练轮数为100,初始学习率为0.01,动量项为0.937。

3.2 数据集

本文实验使用的数据集来自Kaggle平台^[13],包含骑手信息、头盔佩戴情况以及车牌信息等,可用于头盔佩戴检测、车牌识别等任务。

为了提高模型的泛化能力,本文对数据集进行了数据增强,包括添加噪声、亮度调整、图像裁剪、平移变换、随机裁剪以及水平翻转等操作。数据增强后的数据集包含732张图片,按照8:1:1的比例划分为训练集、验证集和测试集。

3.3 评价指标

为了全面且细致地评估本文所提出的模型在电动车头盔佩戴检测任务上的性能,本文选取精确度 P 、召回率 R 、平均精度均值mAP以及mAP50-95等指标对实验结果进行考量。公式为:

$$P = \frac{\text{TP}}{\text{TP} + \text{FP}} \quad (6)$$

$$R = \frac{\text{TP}}{\text{TP} + \text{FN}} \quad (7)$$

$$\text{mAP} = \frac{1}{n} \sum_{i=1}^n \text{AP}_i \quad (8)$$

$$\text{AP} = \int_0^1 P(R) dR \quad (9)$$

式中:TP(true positive)表示正确检测的正样本个数;FP

(false positives) 表示被错误检测为正样本的负样本个数; FN (false negatives) 表示被错误检测为负样本的正样本个数; n 表示数据集中不同类别的数量; AP 是平均精度, 表示精确度 - 召回率 (P-R) 曲线下的面积。

3.4 消融实验

为了全面评估本文所提出算法的实际效能, 本文以 YOLOv8n 作为基础框架, 设置了多组消融实验, 其实验结果如表 1 所示。

表 1 消融实验结果

序号	YOLOv8	SPPELAN	改进 C2f	Focal_Loss	P/%	R/%	mAP/%	mAP50-95/%	GFLOPs
1	√				76.6	77.7	84.9	61.4	8.2
2	√	√			86.2	78.1	85.6	61.7	8.3
3	√		√		86.0	80.0	87.8	65.6	8.5
4	√			√	87.8	78.0	86.0	64.5	8.2
5	√	√	√		87.9	80.0	87.7	64.6	8.6
6	√	√	√	√	87.2	80.4	88.3	65.8	8.6

分析表 1 中的各项关键性能指标, 可以观察到, 通过替换 SPPELAN 模块, 模型的平均精度均值 mAP 得到了提升。针对 C2f 模块的优化策略, 不仅增强了模型的特征提取能力, 同时使 mAP 增长了 2.9%, 展现了显著的改进效果。此外, 将 CIoU 损失函数替换为 Focal_Loss, 促进了模型在多个维度上的性能提升, 使模型定位与分类更加精准。综上所述, 改进后的 YOLOv8 模型, 相较于原始的 YOLOv8n 模型, 在整体性能上实现了增强, 充分验证了本文算法设计的有效性与优越性。

3.5 对比实验

为了证明本文提出的方法在电动车头盔佩戴检测方面的有效性, 本文将改进后的 YOLOv8 模型与部分主流的一阶段目标检测模型: YOLOv3-tiny^[14]、YOLOv5^[15]、YOLOv6^[16] 以及 YOLOv8n 在相同的环境中运行, 进行对比实验, 实验结果如表 2 所示。

表 2 不同网络模型对比实验结果

网络模型	P/%	R/%	mAP/%	mAP50-95/%	GFLOPs
YOLOv3-tiny	82.8	75.5	84.2	60.3	19.0
YOLOv5	81.3	70.8	81.8	56.2	7.2
YOLOv6	78.2	73.4	81.5	59.7	11.9
YOLOv8n	76.6	77.7	84.9	61.4	8.2
Ours	87.2	80.4	88.3	65.8	8.6

从实验结果中发现, 本文提出的模型相较于 YOLOv3、v5、v6 和 v8n 模型在检测性能上均有提升, mAP 分别提升了 4.1%、6.5%、6.8% 和 3.4%。综合对比可知, 本文所提出的模型对于电动车头盔佩戴的检测具有较好的效果。

3.6 结果分析

为了更加直观地展示本文提出模型的检测性能, 选择两个测试场景, 并运用 YOLOv3-tiny、YOLOv5、YOLOv6、YOLOv8 及本文提出的模型进行对比实验, 检测结果如图 6

所示。在 A 组实验场景中, 本文通过调节图像亮度来模拟暗光环境, YOLOv5 与 YOLOv6 均出现了漏检的情况, 而 YOLOv8 不仅漏检还伴有误检情况。相比之下, 本文模型在此场景下依然能够准确识别目标, 并有效提升了检测的置信水平, 展现出对光照

变化的高度鲁棒性; 在 B 组实验场景中, 本文针对复杂场景进行了测试, YOLOv3-tiny、YOLOv5、YOLOv6 均出现了不同程度的误检问题, 而本文模型有效克服了这一难题, 实现了更为精准的检测, 彰显了其在复杂场景下的卓越性能。综上所述, 本文提出的模型不仅在准确性上有所突破, 而且在复杂环境下的鲁棒性和适应性方面也表现出色, 为目标检测领域的研究与应用提供了新的思路与有力支持。



图 6 检测结果

4 总结

本文提出了一种基于改进 YOLOv8n 的电动车骑行者头盔佩戴检测算法。该算法通过使用 SPPELAN 模块、DWRSeg 优化 C2f 模块以及 Focal_DIoU 损失函数, 有效提升了模型的检测性能。实验结果表明, 本文算法在公共数据集上取得了显著的性能提升, 证明了其在头盔佩戴检测领域的有效性和应用潜力。

本文所提出的方法在提升检测性能和泛化能力的同时,

也不可避免地导致了模型参数数量的增加。因此,如何在提升性能的同时,有效降低模型的参数量,即实现模型的轻量化,更好地满足实时性、低能耗等要求,将成为未来研究的一个重要方向。

参考文献:

- [1] RUBAIYAT A H M, TOMA T T, MASOUMEH K K, et al. Automatic detection of helmetuses for construction safety[C]//2016 IEEE/WIC/ACM International Conference on WebIntelligence Workshops (WIW). Piscataway:IEEE, 2016:135-142.
- [2] MACALISANG J R, ALON A S, JARDINIANO M F, et al. Drive-awake: a YOLOv3 machine vision inference approach of eyes closure for drowsy driving detection[C]//Proceedings of the IEEE International Conference on Artificial Intelligence in Engineering and Technology. Piscataway: IEEE, 2021: 1-5.
- [3] BOUHAYANE A, CHAROUH Z, GHOGHO M, et al. A swin transformer-based approach for motorcycle helmet detection[J]. IEEE access, 2023,11: 74410-74419.
- [4] JIA W, XU S Q, LIANG Z, et al.Real-time automatic helmet detection of motorcyclists in urban traffic using improved YOLOv5 detector[J].IET image processing, 2021,15(14):3623-3637.
- [5] 陈扬, 吕艳辉. 基于改进 YOLOv5s 的头盔佩戴检测算法[J]. 沈阳理工大学学报, 2023, 42(5): 11-17.
- [6] 吴昊祺, 刘小军, 周倩文, 等. 基于改进 YOLOv7 的骑行头盔佩戴识别方法[J]. 自动化应用, 2023, 64(21): 1-4.
- [7] 杨琚钱, 胡平, 戴家树. 基于 YOLOv8-scG 神经网络的电动车头盔佩戴检测算法[J/OL]. 重庆工商大学学报(自然科学版), 2024: 1-10[2024-07-23]. <http://kns.cnki.net/kcms/detail/50.1155.N.20240710.2030.002.html>.
- [8] WEI H R, LIU X, XU S C, et al. DWRSeg: Rethinking efficient acquisition of multi-scale contextual information for real-time semantic segmentation[DB/OL]. (2022-12-02)[2024-03-12]. <https://doi.org/10.48550/arXiv.2212.01173>.
- [9] ZHENG Z H, WANG P, REN D W, et al. Enhancing geometric factors in model learning and inference for object detection and instance segmentation[J]. IEEE transactions on cybernetics, 2021, 52(8): 8574-8586.
- [10] ZHENG Z H, WANG P, LIU W, et al. 2019. Distance-IoU loss: faster and better learning for bounding box regression[DB/OL]. (2019-11-19)[2024-06-19]. <https://doi.org/10.48550/arXiv.1911.08287>.
- [11] LLIN Z Y, PRIYA G, ROSS G, et al. Focal loss for dense object detection[C]//Proceedings of the IEEE International Conference on Computer Vision. Piscataway: IEEE, 2017: 2980-2988.
- [12] 鲍晓慧. 基于 YOLO V5 的输电线路鸟巢缺陷检测[J]. 电气技术与经济, 2024(6): 334-337.
- [13] KAGGLE. Rider, with helmet, without helmet, number plate [EB/OL]. [2024-07-26]. <https://www.kaggle.com/datasets/aneesarom/rider-with-helmet-without-helmet-number-plate>.
- [14] REDMON J, FARHADI A. YOLOv3: an incremental improvement[DB/OL]. (2018-04-08)[2024-05-12]. <https://doi.org/10.48550/arXiv.1804.02767>.
- [15] Ultralytics. YOLOv5[EB/OL]. (2020-06-26)[2024-06-01]. <https://github.com/ultralytics/YOLOv5>.
- [16] LI C Y, LI L L, JIANG H L, et al. YOLOv6: a single-stage object detection framework for industrial applications[DB/OL]. (2022-09-07)[2024-07-21]. <https://doi.org/10.48550/arXiv.2209.02976>.

【作者简介】

唐皓阳(2000—), 女, 湖北襄阳人, 硕士研究生, 研究方向: 目标检测与计算机视觉。

肖小玲(1973—), 女, 湖北荆州人, 博士, 教授, 研究生导师, 研究方向: 智能信息处理与网络安全。

(收稿日期: 2024-08-26)