

DTI-YOLO:改进YOLOv10s的交通标志检测模型

刘美辰, 李 杰, 陈廷伟⁺

辽宁大学 信息学部, 沈阳 110036

+ 通信作者 E-mail:twchen@lnu.edu.cn

摘 要:针对交通标志检测中,远景小目标特征易被弱化,难以与复杂背景区分的问题,提出了一种基于改进YOLOv10s的交通标志检测模型(DTI-YOLO)。提出膨胀卷积融合膨胀注意力模块(DDFM)替换PSA模块,设计局部和全局特征提取分支,通过聚焦局部细节与全局语义,抑制噪声干扰,增强模型在复杂背景中分离小目标特征的能力。构建基于二检测层的跨尺度特征融合网络(TDL-CCFN),利用跨尺度特征融合结构和针对小目标设计的二检测层结构,增强深浅层特征间的融合和小目标特征的保留,同时减少了模型的参数量。引入InnerMPDIoU损失函数替换CIoU损失函数,通过融合可调节尺度因子和顶点几何距离度量,增强模型对小目标位置和视角变化的敏感性,提升边界框回归效率与模型泛化能力。实验结果表明,DTI-YOLO模型具有良好的检测性能,相较于YOLOv10s,DTI-YOLO在TT100K和CCTSDB数据集上的mAP50分别提升5.5和4.8个百分点,分别达到90.9%和86.6%;同时,参数量减少约33.3%,降至 5.4×10^6 ,实现了模型轻量化。

关键词:交通标志检测;DTI-YOLO;复杂背景;小目标特征;多尺度特征

文献标志码:A **中图分类号:**TP391.41 **doi:**10.3778/j.issn.1002-8331.2501-0304

DTI-YOLO: Improved YOLOv10s Traffic Sign Detection Model

LIU Meichen, LI Jie, CHEN Tingwei⁺

Faculty of Information, Liaoning University, Shenyang 110036, China

Abstract: To address the issues in traffic sign detection where distant small object features are easily weakened and difficult to distinguish from complex backgrounds, this paper proposes an improved YOLOv10s-based model named DTI-YOLO. Firstly, the PSA module is replaced with the dilated convolution and dilated attention fusion module (DDFM), designed with local and global feature extraction branches that focus on local details and global semantics, respectively, helping to suppress background noise and enhancing the model's ability to separate small object features from complex scenes. Secondly, a two detection layer-based cross-scale feature fusion network (TDL-CCFN) is constructed, which integrates the cross-scale fusion structure and the two detection layer mechanism tailored for small objects, improving feature fusion between deep and shallow layers, enhancing small object feature retention, and reducing model parameters. Finally, the InnerMPDIoU loss function replaces the CIoU loss, combining an adjustable scale factor with a vertex-based geometric distance metric to enhance the model's sensitivity to the position and perspective changes of small objects, thereby improving bounding box regression performance and generalization capability. Experimental results demonstrate that, compared to YOLOv10s, the proposed DTI-YOLO achieves excellent detection performance and improves mAP50 by 5.5 and 4.8 percentage points on the TT100K and CCTSDB datasets respectively, reaching 90.9% and 86.6%, while the number of parameters is reduced by approximately 33.3%, to 5.4×10^6 , achieving significant model lightweighting.

Key words: traffic sign detection; DTI-YOLO; complex background; small object features; multi-scale characteristics

在人工智能技术快速发展和道路车辆持续增加的背景下,智能交通系统快速发展,道路交通标志的准确检测对于实现智能驾驶、交通管理和驾驶员辅助等应用至关重要^[1]。交通标志检测能够帮助驾驶员或驾驶系统

基金项目:辽宁省自然科学基金计划项目(2023-BS-085);辽宁省教育厅高校基本科研课题项目(LJ212410140012);辽宁大学基本科研课题项目(LJKLJ202421)。

作者简介:刘美辰(1999—),女,硕士研究生,CCF学生会会员,研究方向为目标检测、机器学习;李杰(2001—),男,硕士研究生,CCF学生会会员,研究方向为目标检测、机器学习;陈廷伟(1974—),男,博士,教授,CCF高级会员,研究方向为智能交通、机器学习。

收稿日期:2025-01-20 **修回日期:**2025-04-24 **文章编号:**1002-8331(2025)17-0112-11

更好地感知路况,尤其在复杂背景、夜间或极端天气等场景下,有效减少交通事故,提升道路通行效率和行车安全。然而,交通标志检测需要在现实场景中进行实时检测,环境多变是其特点之一,进而检测速度和精度往往会受到较大影响^[2]。因此,研究高效、可靠的交通标志检测模型,对于提升道路交通管理的效率具有重要意义。

近年来,深度学习技术在计算机视觉领域取得了显著进展,基于深度学习目标检测算法具备强大的特征提取能力,能够有效地获取细节信息和关键特征^[3],从而逐渐成为小目标检测领域的研究热点。当前主流的目标检测算法根据是否直接预测目标的类别和位置,可分为双阶段检测模型和单阶段检测模型。双阶段检测模型如 Cascade R-CNN^[4]等,首先生成候选区域,再进行分类与回归,虽然具备较高的检测精度,但训练和推理开销较大,难以满足实时性需求。相比之下,SSD^[5]、RT-DETR^[6]和 YOLO^[7]系列等单阶段检测模型将候选区域生成、分类和回归整合至同一阶段,具备更高的计算效率和更简洁的结构,更适合对实时性要求高、算力有限的场景。在小目标检测领域,模型需要在复杂背景中精准定位和识别尺寸较小的目标实例。为此,研究者对如何提高小目标检测性能展开了诸多研究。刘思元等人^[8]提出改进 RT-DETR 的航拍小目标检测模型,通过引入双向特征金字塔网络和轻量级 CARAFE 上采样算子优化多尺度特征融合,缓解语义信息丢失,同时利用 S2 特征增强小目标表达,并设计了基于部分卷积的 BasicBlock-PConv-Block 模块以减少参数量。Song 等人^[9]针对智能驾驶中的小目标检测,提出 HPRT-DETR 模型,设计 BasicIRMB-CGA(BIC)模块以有效提取特征并降低参数量,结合可变形注意力与尺度内特征交互构建 DAIFI 模块捕获丰富的语义特征,提高对遮挡小目标的检测能力。尽管这些改进在模型轻量化方面有所进展,但基于 Transformer 架构的 RT-DETR 计算开销较大,对于无人机等资源受限的设备而言仍存在优化空间。Mu 等人^[10]针对小目标检测难题提出 YOLOv8n 改进模型,通过融合检测头提取有效特征,并引入多头混合注意力机制(MMSA),结合通道注意力与空间信息,提升对小目标的检测性能。朱硕等人^[11]针对交通标志检测对 YOLOv5s 进行改进,设计 C3_DB 模块来增强多尺度特征处理能力,并通过上下文聚合模块减少复杂背景的干扰,提高对重点区域的关注。罗向龙等人^[12]利用全维动态卷积和 EMA 注意力机制改进 YOLOv8n 主干网络,增强网络对交通标志的特征提取的精准性和对上下文信息的表征能力,并引入 GhostBottleneckv2 进行轻量化处理,有效降低模型复杂度。Han 等人^[13]提出 YOLO-SG 交通标志检测模型,采用 GhostNet 作为特征提取架构并通过 SPD-Conv 进行下采样,减少下采样过程中的特征损失。上述方法主要侧重全局特征建模以加强上下文语义信息,而对局部特征的考虑并不充分,导致模型在

处理局部细节特征时存在不足,尤其是在远景场景下,小目标的像素占比小,易被背景噪声干扰,导致模型检测不稳定。而在交通标志检测中,箭头方向、边框形状、限速数字等局部细节特征,往往是目标分类和识别的关键。因此,交通标志领域的小目标检测更依赖于局部细节特征来确保检测的精准。此外,尽管现有改进方法在提升模型性能方面取得了一定成效,但往往伴随模型推理速度降低和参数量增大,限制了其在车载端上的实际应用。针对上述挑战,本文提出一种基于 YOLOv10s 改进的 DTI-YOLO 交通标志检测模型,旨在提升远景小目标在复杂背景下的检测精度,同时实现模型轻量化,主要工作如下:

(1)提出了联合局部膨胀卷积和全局膨胀注意力机制的特征提取模块 DDFM 替换 PSA 模块,通过构建多尺度感受野协同膨胀注意力机制,增强了模型对小目标局部细节特征的关注,并在全局感受野的基础上强化小目标的局部特征建模,从而有效提升模型在复杂背景下对小目标的识别能力。

(2)设计了 TDL-CCFN 特征融合网络,利用 CCFM 结构构建轻量级跨尺度特征传递通道,实现高效的多尺度特征融合。新增的 160×160 小目标检测层使模型更加关注小目标特征,减少了小目标特征的损失,并与 80×80 检测层共同构成二检测层,有效减少了模型的参数量。

(3)引入 InnerMPDIoU 损失函数替换 CIoU 损失函数,通过约束预测框与真实框的顶点几何距离和建立动态尺度感知,提升了模型对小目标的定位精度与视角变化的鲁棒性,同时提高了边界框回归效率和模型的泛化能力。

1 YOLOv10 模型

YOLOv10^[14]由清华大学团队开发,旨在解决前代版本在后处理和模型架构方面的不足。其核心架构包括主干网络、颈部和检测头三部分。主干网络利用大核卷积和部分自注意力机制(partial self-attention, PSA)增强特征提取。颈部网络通过 PAN 层汇聚不同尺度的特征,实现多尺度特征的高效融合。检测头采用一对多头和一对一头设计,训练时一对多头生成多个预测以提供丰富监督信号,推理时一对一头生成最终预测,无须 NMS 操作,降低延迟,提升检测效率。此外,为满足不同应用场景的需求,YOLOv10 提供了六种模型版本:n、s、m、b、l、x。考虑到交通标志检测任务对检测精度和模型轻量化的要求较高,本文选择 YOLOv10s 作为基础模型,针对交通标志检测任务进行改进设计。

2 DTI-YOLO 交通标志检测模型

2.1 改进 YOLOv10s 的 DTI-YOLO 模型

远景小尺寸交通标志易受背景干扰,影响检测精

度。本文基于YOLOv10s模型进行改进:首先在主干网络中,设计DDFM模块替换原有PSA模块,结合膨胀卷积与膨胀注意力机制,增强复杂背景中对小目标的特征提取能力。其次构建TDL-CCFN特征融合网络,引入轻量级跨尺度特征融合结构并优化检测层设计,新增 160×160 检测层,去除 20×20 和 40×40 检测层,增强多尺度特征融合能力,减少小目标特征损失。最后采用Inner-MPDIoU损失函数优化边界框回归过程。改进后的DTI-YOLO模型结构如图1所示。

2.2 DDFM特征提取模块

YOLOv10s主干网络中的PSA模块通过直接跳跃连接和注意力建模来保留原始特征信息,并强化全局特征提取。然而,由于PSA模块在局部特征捕捉上存在局限性,这导致模型在背景噪声较多的交通标志检测中检测精度下降。为了解决这一问题,本文借鉴卷积融合注意力模块(convolution and attention fusion module, CAFM)^[15]中卷积与注意力机制协同作用的思想,设计DDFM模块来增强主干网络的局部和全局特征提取能力,从而提高模型在复杂背景中的性能表现。

DDFM模块由局部分支和全局分支组成。局部分支通过膨胀卷积(dilated convolution)^[16]扩展感受野,有效捕捉局部细节和上下文关联;全局分支采用滑动窗口膨胀注意力机制(sliding window dilated attention, SWDA)^[17],结合滑动窗口策略与注意力机制,建模全局长距离依赖特征。通过局部分支与全局分支的协同作

用,DDFM模块实现了局部与全局特征联合建模,从而增强模型在复杂背景中对交通标志的感知能力。DDFM结构如图2所示。

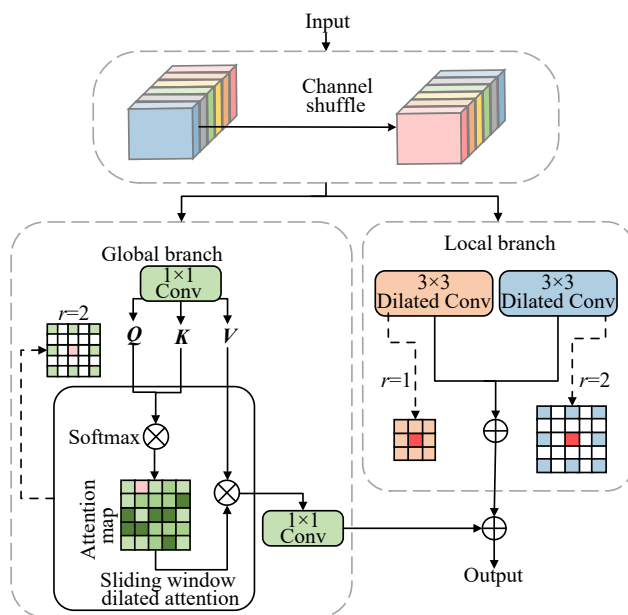


图2 DDFM结构

Fig.2 DDFM structure

输入特征首先经过通道洗牌操作进行重新分组和排列,然后分别传递到局部分支和全局分支进行独立处理,以增强跨通道特征的交互和融合能力。

在局部分支中,首先并行使用两个 3×3 膨胀卷积,

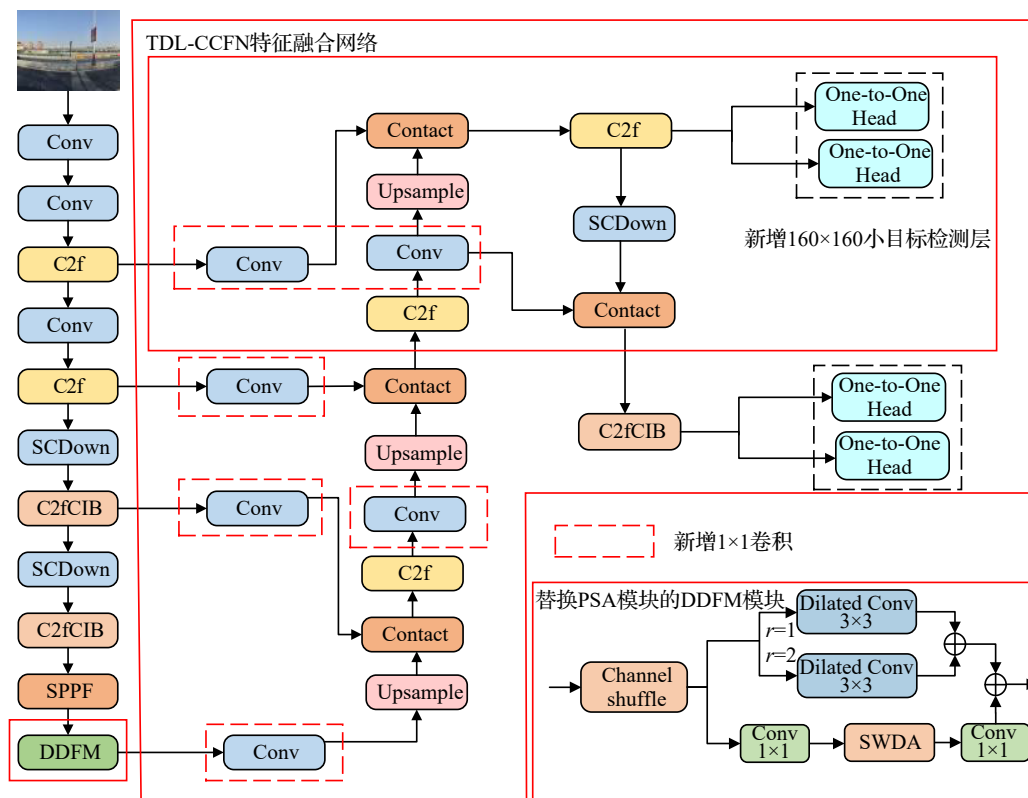


图1 DTI-YOLO 模型结构

Fig.1 DTI-YOLO model structure

膨胀率分别为1和2,分别对分支输入特征 Y_{CS} 进行独立处理,并将结果相加后整合为新的特征表示。膨胀率的设置引入多尺度感受野,在不增加参数数量的情况下扩大了感受野,同时保留特征图的分辨率,使得卷积核能够在更大空间范围内捕获更多细节特征。这有助于分离交通标志与复杂背景的细微特征,特别是在交通标志尺寸较小或检测距离变化显著时,能够有效应对多尺度问题以提升特征提取的鲁棒性与准确性。局部分支表述如式(1):

$$F_{\text{conv}} = \sum Y_{\text{DConv}} \quad (1)$$

$$Y_{\text{DConv}} = \sum_k Y_{CS}[i+r \times k] \times w[k] \quad (2)$$

其中, F_{conv} 表示局部分支的输出, Y_{CS} 表示分支输入特征, i 和 k 表示位置信息, r 表示膨胀率, $w[k]$ 表示权重。

在全局分支中,首先分支特征 Y_{CS} 通过 1×1 卷积生成查询矩阵(Q)、键矩阵(K)和值矩阵(V),然后进行SWDA。在原始特征图中,SWDA以 (i,j) 查询为中心,在大小为 $w \times w$ 的滑动窗口中稀疏选择键和值,并对所有查询块进行自注意力操作,用膨胀率 $r \in \mathbb{N}^+$ 来控制稀疏性的程度。SWDA通过满足局部性和稀疏性的特性,在减少计算冗余的同时,保持模型对长距离依赖关系的建模能力,从而能够捕捉到更广泛的全局上下文信息。接着对SWDA的输出进行 1×1 卷积。全局分支的表述如式(3):

$$F_{\text{att}} = W_{1 \times 1} Y_{DA} \quad (3)$$

$$(Q, K, V) = W_{1 \times 1}(Y_{CS}) \quad (4)$$

$$Y_{DA} = \text{SWDA}(Q, K, V, r) \quad (5)$$

$$y_{ij} = \text{Attention}(q_{ij}, K_r, V_r) =$$

$$\text{Softmax}\left(\frac{q_{ij} \times K_r^T}{\sqrt{d_k}}\right) \times V_r \quad (6)$$

$$\{(i', j') | i' = i + p \times r, j' = j + q \times r\} \quad (7)$$

其中, Y_{DA} 为SWDA的输出特征, y_{ij} 为输出特征 Y_{DA} 对应的分量; (i, j) 为 y_{ij} 对应的位置, K_r 和 V_r 表示从特征图 K 和 V 中选择的键和值, d 是可学习的缩放参数; (i', j') 为在滑动窗口中选择键和值的坐标集合; p 和 q 是相对于查询位置 (i, j) 的偏移量。

最后DDFM模块综合两个分支输出,实现全局和局部特征的联合建模,输出计算表述如式(8):

$$F_{\text{out}} = F_{\text{att}} + F_{\text{conv}} \quad (8)$$

2.3 TDL-CCFM特征融合网络

2.3.1 CCFM模块

YOLOv10s颈部网络通过自底向上和自顶向下的双向融合机制,将浅层和深层的相邻层级特征结合起来。然而,随着网络逐渐加深,逐层传递的特征更多地转换为深层语义信息,而包含大量交通标志细粒度信息的浅层特征可能被弱化甚至丢失。如果缺乏有效的深浅层特征融合,模型在抑制噪声和区分目标方面的能力

会减弱,容易在检测时出现交通标志的误检或漏检。为充分利用各层级的尺度特征,本文借鉴轻量级跨尺度特征融合模块(cross-scale feature-fusion module, CCFM)的结构对YOLOv10s颈部网络进行了改进,以加强多尺度特征之间的联系并增强特征融合能力。CCFM由基于Transformer架构的实时端到端检测器RT-DETR提出,将不同尺度的特征通过独立卷积处理与跨层级特征交互融合,从而在降低模型参数数量的同时,增强模型对尺度变化的适应性以及小尺度目标的检测能力。CCFM结构如图3所示。

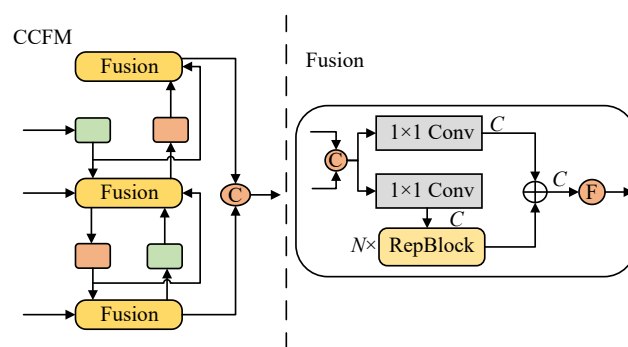


图3 CCFM结构

Fig.3 CCFM structure

2.3.2 重组检测层

YOLOv10s原始网络通过特征融合生成3种尺度分别为 20×20 、 40×40 、 80×80 的特征图用于检测不同大小的目标。但在远距离视角下,交通标志因尺寸较小而成为小目标,因此增强浅层特征中轮廓和位置信息的保留能力对于提高交通标志的检测精度至关重要。为此,本文对YOLOv10s的检测层结构进行了改进,新增了 160×160 检测层以保留浅层特征的细节信息,同时移除了原有的 20×20 和 40×40 检测层,以减少参数数量和计算开销,使模型聚焦于 80×80 和 160×160 检测层。重组检测层结构如图4所示。

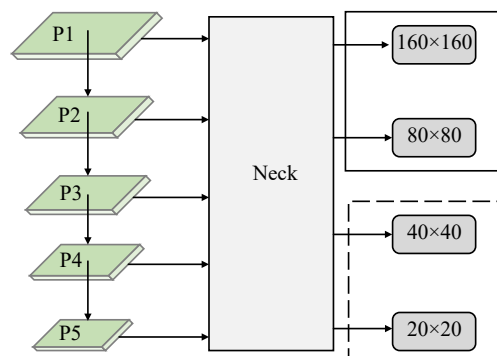


图4 重组检测层

Fig.4 Recombination detection layer

2.3.3 基于二检测层的跨尺度特征融合网络

结合改进的颈部网络与重组的检测层设计出基于二检测层的跨尺度特征融合网络(two detection layer-

based cross-scale feature-fusion network, TDL-CCFN) 特征融合网络,在增强小目标特征保留能力的同时,实现跨尺度特征融合并进一步轻量化检测模型。TDL-CCFN结构如图5所示。

在具体实现中,基于CCFM的跨尺度融合结构将从主干网络提取的特征图层P2、P3、P4、P5通过独立的融合路径输入到颈部网络中。每条融合路径中的融合块依次对每个尺度的特征图进行轻量级 1×1 卷积以压缩通道数。随后,通过融合块中的上采样操作,将卷积处理后的P5特征与P4特征拼接融合,映射为同一维度的特征。经过C2f模块进一步强化特征表达能力,融合后的P4特征自底向上传递到P3融合块,通过轻量级 1×1 卷积和上采样操作后,与P3特征拼接融合,使浅层特征包含更丰富的上下文信息。为了进一步加强小目标特征的保留能力,新增的 160×160 高分辨率检测层从浅层特征中提取更多细节信息。融合后的P3特征继续向上传递到P2融合块,进行卷积与上采样操作后,与P2特征拼接,实现跨尺度特征融合。最终,通过 80×80 和 160×160 的检测头进行两路输出,在增强模型对行驶中交通标志的尺度变化适应性的同时,提高对小尺寸交通标志的检测能力。

2.4 改进损失函数

在目标检测任务中,损失函数的组成通常包括边界框回归损失、分类损失和置信度损失,分别优化预测框的位置、类别和置信度。通过设计合理的边界框回归损失函数,可以减少模型的定位误差,从而提高推理效率与检测精度。YOLOv10s采用CIoU^[18]损失函数计算预测框与真实框之间的位置差异。CIoU的定义如下:

$$L_{CIoU} = 1 - IoU + \frac{\rho^2(b, b^{gt})}{c^2} + \alpha v \quad (9)$$

$$\alpha = \frac{v}{(1 - IoU) + v} \quad (10)$$

$$v = \frac{4}{\pi^2} \left(\arctan \frac{w^{gt}}{h^{gt}} - \arctan \frac{w}{h} \right)^2 \quad (11)$$

其中, b^{gt} 和 b 分别表示真实框和预测框的中心点, ρ 为

两框中心点间的欧氏距离, c 为覆盖两框的最小封闭框的对角线长度, w^{gt} 和 h^{gt} 表示真实框的长度与宽度, w 和 h 表示预测框的长度与宽度, v 为两框长宽比的惩罚项, α 为权重参数。

在交通标志检测中,目标尺度随距离变化较大。当预测框和真实框的长宽比接近或相等时,长宽比惩罚项的作用不显著,CIoU难以准确反映位置误差对总损失的影响,从而限制边界框回归的精度和收敛速度,并可能增加漏检现象。此外,CIoU对不同IoU范围内的样本缺乏相应调整,限制了模型的泛化能力。为了提升边界框回归性能,本文提出用改进的InnerMPDIoU(inner minimum point distance based IoU)损失函数替代CIoU, InnerMPDIoU计算方法如图6所示。

Inner-IoU^[19]通过引入可调节的尺度因子ratio,生成不同尺度的辅助边界框,其取值范围为 $[0.5, 1.5]$,对高IoU和低IoU样本分别施加不同的梯度影响以提升模型性能和泛化能力。针对高IoU样本,生成较小尺度的辅助边界框($ratio \in [0.5, 1)$)强化局部梯度信号,加速回归;而对于低IoU样本,生成较大尺度的辅助边界框($ratio \in (1, 1.5]$)增强负梯度信号,缩短收敛时间。Inner-IoU的定义如下:

$$L_{Inner-IoU} = 1 - IoU^{inner} \quad (12)$$

$$IoU^{inner} = \frac{inter}{union} \quad (13)$$

$$b_l^{gt} = x_c^{gt} - \frac{w^{gt} \times ratio}{2}, b_r^{gt} = x_c^{gt} + \frac{w^{gt} \times ratio}{2} \quad (14)$$

$$b_t^{gt} = y_c^{gt} - \frac{h^{gt} \times ratio}{2}, b_b^{gt} = y_c^{gt} + \frac{h^{gt} \times ratio}{2} \quad (15)$$

$$b_l = x_c - \frac{w \times ratio}{2}, b_r = x_c + \frac{w \times ratio}{2} \quad (16)$$

$$b_t = y_c - \frac{h \times ratio}{2}, b_b = y_c + \frac{h \times ratio}{2} \quad (17)$$

$$inter = (\min(b_r^{gt}, b_r) - \max(b_l^{gt}, b_l)) \times (\min(b_b^{gt}, b_b) - \max(b_t^{gt}, b_t)) \quad (18)$$

$$union = (w^{gt} \times h^{gt}) \times (ratio)^2 + (w \times h) \times (ratio)^2 - inter \quad (19)$$

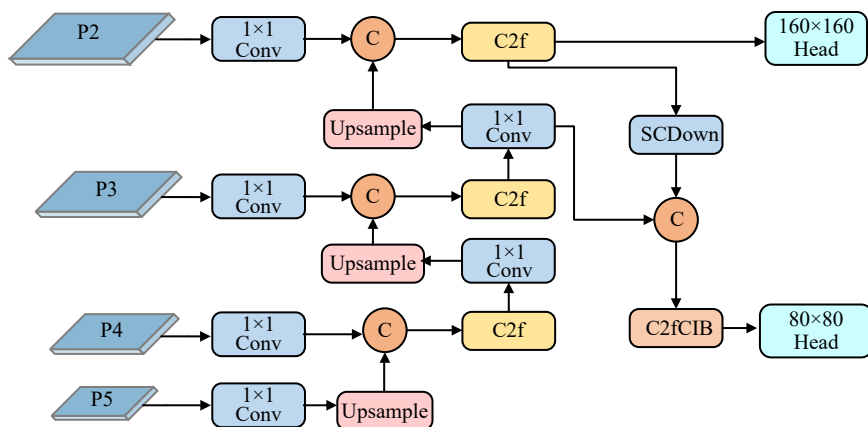


图5 TDL-CCFN 结构

Fig.5 TDL-CCFN structure

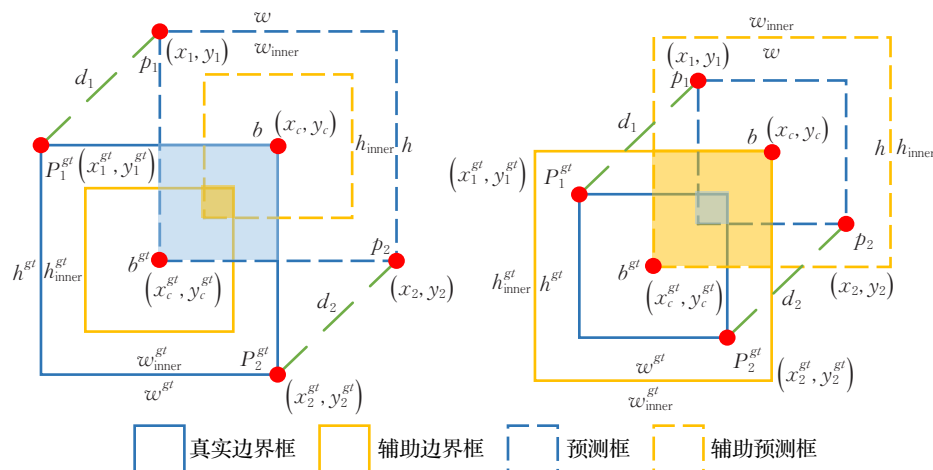


图6 InnerMPDIoU 计算方法

Fig.6 InnerMPDIoU calculation method

其中, (x_c^{gt}, y_c^{gt}) 和 (x_c, y_c) 分别表示辅助边界框和辅助预测框的中心点坐标。 w^{gt} 和 h^{gt} 、 w 和 h 、 w_{inner}^{gt} 和 h_{inner}^{gt} 、 w_{inner} 和 h_{inner} 分别表示真实框、预测框、辅助边界框和辅助预测框的长度与宽度。 (b_l^{gt}, b_t^{gt}) 、 (b_r^{gt}, b_b^{gt}) 、 (b_l, b_t) 、 (b_r, b_b) 和 (b_l, b_t) 、 (b_l, b_b) 、 (b_r, b_t) 、 (b_r, b_b) 分别表示辅助边界框和辅助预测框的四个顶点坐标。

MPDIoU^[20]通过最小化预测框与真实框对应的左上角与右下角点间距离,简化两个边界框之间的相似性比较,能够很好地处理重叠或非重叠的边界框回归,提高收敛速度和回归结果的准确性,实现了更高效的边界框回归。MPDIoU的定义如下:

$$L_{MPDIoU} = 1 - IoU + \frac{d_1^2 + d_2^2}{h^2 + w^2} \quad (20)$$

$$d_1^2 = (x_1^{gt} - x_1)^2 + (y_1^{gt} - y_1)^2 \quad (21)$$

$$d_2^2 = (x_2^{gt} - x_2)^2 + (y_2^{gt} - y_2)^2 \quad (22)$$

其中, (x_1^{gt}, y_1^{gt}) 、 (x_2^{gt}, y_2^{gt}) 和 (x_1, y_1) 、 (x_2, y_2) 分别表示真实框和预测框的左上角点坐标和右下角点坐标; d_1^2 和 d_2^2 分别表示两框的左上角点和右下角点间欧氏距离的平方。

InnerMPDIoU的定义如下:

$$L_{InnerMPDIoU} = L_{MPDIoU} + IoU - IoU_{inner} \quad (23)$$

InnerMPDIoU 通过可调节尺度因子和顶点几何距离度量,增强模型对不同 IoU 样本的适应性,降低检测框的定位误差,从而有助于增强模型对复杂背景中小目标位置的敏感性和对视角变化所带来的目标定位误差的鲁棒性,同时有效提升边界框回归的速度和精度,增强模型泛化性。

3 实验结果及分析

3.1 实验数据集

清华-腾讯 100K 数据集(Tsinghua-Tencent 100K,

TT100K)是一个真实场景下的交通标志检测数据集。TT100K^[21]由来自腾讯地图收集的街景数据图像构成,共包含 10 万张图像,其中包含 3 万个交通标志实例,涵盖指示、警告、禁止 3 大类共 128 小类的中国交通标志。在真实交通场景中,该数据集的类别存在不均衡的问题,例如,“山体滑坡”标志的数量相对于“禁止停车”等类别的数量明显较少。因此本文对 TT100K 进行筛选,选择 45 类实例数大于 100 的常见类别进行实验,有效避免图像类别不均衡的问题。

为了进一步验证本文模型在交通标志检测中是否具有泛用性,选取中国交通标志检测数据集(CSUST Chinese traffic sign detection benchmark, CCTSDB)进行实验。CCTSDB^[22]由长沙理工大学的学者团队制作而成,共 17 856 张真实交通场景下的交通标志图像。该数据集标注指示(mandatory)、警告(warning)和禁止(prohibitory)3 类标志,并融入丰富的道路背景,涵盖晴天、阴天、雨雪、大雾和黑夜等复杂交通环境。

3.2 实验环境及相关参数

本文实验软件环境为 Ubuntu20.04, Python3.9.0, Pytorch2.0.1。硬件环境包括 NVIDIA A10 24 GB GPU 和 Intel Xeon (Ice Lake) Platinum 8369B CPU, 使用 CUDA11.8 对 GPU 进行加速,具体实验参数配置如表 1 所示。

表1 实验参数配置

Table 1 Experimental parameter configuration

参数	配置
输入尺寸	640×640
训练轮数	300
批量大小	16
线程数量	8
优化器	SGD
初始学习率	0.001
动量	0.937
权重衰减系数	0.000 5

3.3 实验评价指标

本文实验选取精确率(precision, P)、召回率(Recall, R)、平均精度均值(mean average precision, mAP)、模型参数量(params)以及每秒传输帧数(frames per second, FPS)作为评价指标来评估模型的整体性能。

精确率越高,模型的误检率越低。召回率越高,模型的漏检率越低。AP是精确率-召回率曲线下的面积,用于衡量模型在不同阈值下的表现。mAP50是在IoU阈值为0.5时的所有类别AP的平均精度均值;mAP50-95是在IoU阈值为0.5到0.95时的平均精度均值,mAP越高表明模型在多个类别上的检测效果越好。各个评价指标的计算公式如下:

$$P = \frac{TP}{TP + FP} \quad (24)$$

$$R = \frac{TP}{TP + FN} \quad (25)$$

$$AP = \int_0^1 P(R) dR \quad (26)$$

$$mAP = \frac{\sum_{i=0}^n AP(i)}{N} \quad (27)$$

其中, TP(true positives):模型正确检测出的正样本数量。FP(false positives):模型错误检测出的正样本数量(误检)。FN(false negatives):模型未能检测出的正样本数量(漏检)。N为物体类别的总数。

3.4 特征融合网络对比实验分析

DTI-YOLO模型通过引入CCFM和二检测层结构,设计TDL-CCFN特征融合网络。为了验证TDL-CCFN特征融合网络的检测效果,在TT100K数据集上与当前主流的网络BiFPN^[23]、Slim-neck^[24]进行对比实验。为了保证实验的一致与公平,对基于BiFPN和Slim-neck改进后的颈部网络结合与TDL-CCFN相同的二检测头结构,分别记为TDL-BiFPN和TDL-Slim-neck,实验结果如表2所示。与基础模型相比,TDL-BiFPN通过动态加权和双向融合充分融合不同层级的特征,有效缓解特征逐层弱化的情况,召回率提升3.6个百分点,但精确率提升微弱仅为0.1个百分点;TDL-Slim-neck的召回率提升1.5个百分点,但对小目标特征检测不敏感,精确率明显下降。而TDL-CCFN对精确率和召回率均有改善,其中召回率提升4.9个百分点,可见引入CCFM结构后模型对交通标志的检测更为精准和全面。与对改善多尺度融合具有良好效果的TDL-BiFPN相比,TDL-CCFN的mAP50和mAP50-95分别提升3.1、3.0个百分点,由此证明TDL-CCFN对提升模型检测性能具有优势。

3.5 损失函数对比实验分析

在TT100K数据集上,逐一使用CIoU、Inner-IoU、MPIoU和InnerMPDIoU损失函数进行训练,检测结果如表3所示。与CIoU相比,单独使用Inner-IoU和MPIoU在精确率和平均精度均值上均有提升。虽然在召回率

表2 特征融合网络对比实验

Table 2 Feature fusion network comparison experiment 单位:%

实验	P	R	mAP50	mAP50-95
YOLOvLOs	86.0	76.4	85.4	66.9
TDL-BiFPN	86.1	80.0	87.4	68.6
TDL-Slim-neck	84.7	77.9	85.5	66.1
TDL-CCFN	87.3	84.9	90.5	71.6

上Inner-IoU下降0.3个百分点,但MPIoU很好地弥补了这一略微劣势,召回率提升1.1个百分点。因此,证明Inner-IoU和MPIoU对CIoU的不足均有改善之处。

表3 损失函数对比实验

Table 3 Loss function comparison experiment 单位:%

损失函数	P	R	mAP50	mAP50-95
CIoU	86.0	76.4	85.4	66.9
InnerIoU	86.8	76.1	85.7	67.0
MPIoU	86.2	77.5	85.8	67.7
InnerMPDIoU	87.7	76.5	86.2	68.1

由图7曲线可以看出,在相同训练轮数下InnerMPDIoU收敛速度更快,最终损失值更低。在训练初期(0~50轮),InnerMPDIoU损失下降幅度更大,表明其能更高效地优化边界框回归,减少误差,提高梯度更新效率,使模型更快收敛并提升检测性能。在300轮时,InnerMPDIoU的最终损失值低于CIoU和其他损失函数,进一步验证了其优化效果,更有助于精准拟合目标框,提高检测精度。

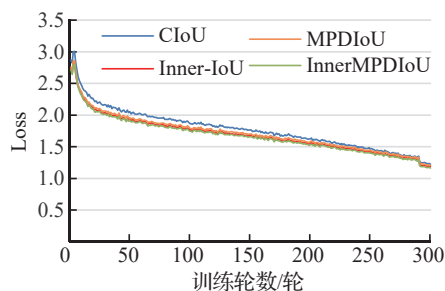


图7 损失值对比曲线

Fig.7 Loss value comparison curve

为了进一步验证引入InnerMPDIoU的模型是否具有泛化能力,本文在实际场景中对交通标志进行了拍摄。为评估模型对小目标的位置和拍摄视角变化的敏感性,选取了不同视角(侧面视角)下,呈现标志尺寸较小的图片进行检测,图片从左到右依次为原图、基础模型和改进损失函数后模型的检测结果。由于图片中交通标志较小,为了清晰呈现,本文对原图的检测区域进行了放大处理。

如图8所示,引入InnerMPDIoU的模型在未见过的检测场景下依然能够保持较优的检测效果,展现了其良好的泛化能力。在图8(a)中,即使在标志尺寸较小且为

非正面视角的情况下,由于InnerMPDIoU对小目标位置的高度敏感性,改进后的模型成功检测到了pl30和pn标志,而基础模型却出现了漏检现象。在图8(b)中,基础模型对pl40和pn标志的检测置信度明显低于改进损失函数后的模型,后者提供了更高的置信度评分,进一步验证了InnerMPDIoU的优越性。

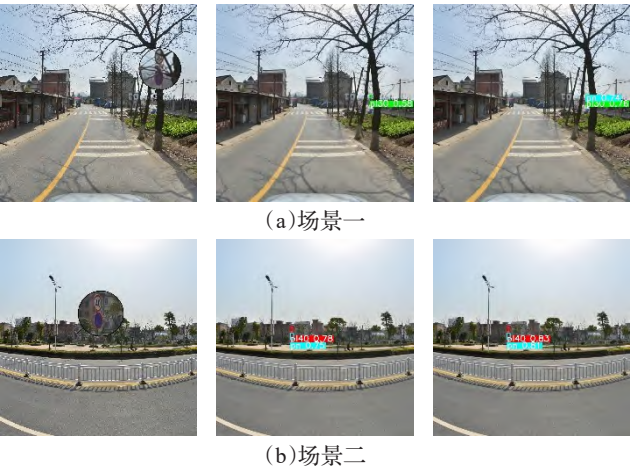


图8 损失函数检测效果对比

Fig.8 Comparison of loss function detection effects

3.6 消融实验

为了验证DDFM特征提取模块、TDL-CCFN特征融合网络和InnerMPDIoU损失函数对YOLOv10s模型检测性能的提升效果,本文在TT100K数据集上对DTI-YOLO模型进行消融实验。依次将上述改进方法记为A、B、C,实验结果如表4所示。由表4中的数据可知,在主干网络特征提取阶段,A方法设计卷积和注意力融合模块DDFM后,在保持参数量不变的情况下,由于局部分支的膨胀卷积扩大了感受野,加强了模型对局部标志和整体背景的把握,精确率和召回率都有明显提升,分别提升2.8和2个百分点,有效降低由于检测的交通标志尺寸较小难以从较为复杂的背景中分离出来,而造成的误检与漏检情况。B方法设计特征融合网络TDL-CCFN后,跨层级特征交互融合和160×160小目标检测层增强了小尺度目标特征的保留能力和模型对多尺度变化的适应性,明显降低交通标志的漏检情况,召回率提升8.8个百分点,mAP50和mAP50-95分别提升5.1,4.7个百分点,模型能够更全面准确地检测出交通标志。此外,由于轻量化CCFM结构的引入以及原有20×20和40×40检测层的移除,模型参数量从8.1×10⁶降到5.4×10⁶,降低33.3%,达到轻量化模型的效果。C方法采用InnerMPDIoU损失函数替换CIoU,在精确率、召回率、mAP50和mAP50-95方面较基础模型均有提升。上述表明,单独使用A、B、C方法对模型的检测性能均有不同程度的改善。将A与B方法组合起来后,由于局部特征提取与多尺度特征融合的加强,模型的检测效果进一步提升,相较于单独使用B,mAP50和mAP50-95提升

0.3和0.2个百分点。最后将A、B、C全部组合起来得到DTI-YOLO模型,由实验数据可知,与基础模型相比,精确率、召回率、mAP50和mAP50-95分别提升4.7、8.2、5.5和5.1个百分点。综上表明,在3个方法的协同作用下,显著提升了模型对检测目标特征的捕捉与分类能力,有效降低漏检与误检的情况。因此,DTI-YOLO模型适用于对检测精度要求较高的交通标志检测任务,同时实现了模型的轻量化设计,便于在道路车辆的移动端设备中部署。

表4 DTI-YOLO模型在TT100K上的消融实验

Table 4 Ablation experiment of DTI-YOLO model on TT100K

实验	A	B	C	Params/10 ⁶	P/%	R/%	mAP50/%	mAP50-95/%
1	—	—	—	8.1	86.0	76.0	85.4	66.9
2	✓	—	—	8.1	88.8	78.0	86.2	67.3
3	—	✓	—	5.4	87.3	84.8	90.5	71.6
4	—	—	✓	8.1	87.7	76.5	86.2	68.1
5	✓	✓	—	5.4	87.7	84.8	90.8	71.8
6	✓	✓	✓	5.4	89.4	84.9	90.9	72.0

3.7 对比实验

为了验证本文提出的改进模型DTI-YOLO在交通标志检测任务中的高效性,在TT100K数据集上选取当前主流的单阶段和双阶段目标检测模型进行对比实验,采用上述5个评价指标作为衡量标准,各模型检测结果如表5所示。

表5 各模型在TT100K上的性能对比实验

Table 5 Performance comparison experiments of each model on TT100K

模型	P/%	R/%	mAP50/%	Params/10 ⁶	FPS
Cascade-RCNN	—	—	65.6	69.30	30.8
SDD	—	—	63.7	26.30	48.9
YOLOv8s ^[25]	83.0	76.4	84.2	11.10	126.1
YOLOv9s ^[26]	83.4	74.4	83.4	7.20	97.6
YOLOv10n	78.6	69.0	75.1	2.70	135.7
YOLOv10s	86.0	76.4	85.4	8.10	101.7
YOLO11s	85.3	78.9	85.5	9.40	111.3
YOLOv12s ^[27]	87.6	76.1	85.6	9.20	98.0
RT-DETR-R34	82.7	75.9	85.9	30.30	69.0
文献[8]	89.1	83.5	87.6	17.30	—
BH-RT-DETR ^[28]	—	—	87.8	18.20	—
文献[10]	—	—	88.1	10.90	—
文献[11]	70.0	67.9	71.8	1.90	294.0
YOLO-SG	—	—	75.8	4.00	131.6
文献[29]	—	—	90.5	7.62	39.4
DTI-YOLO	89.4	84.9	90.9	5.40	115.3

从表中数据可知,双阶段模型Cascade-RCNN和传统单阶段模型SDD在检测精度方面有明显短板,二者FPS较低,且Cascade-RCNN参数量过高,难以满足轻量化部署需求。YOLOv8s^[25]、YOLOv9s^[26]、YOLOv10s、YOLO11s和YOLOv12s^[27]通过结构优化,提供了较优的

检测精度与速度,但与DTI-YOLO相比,在模型轻量化部署方面还存在可提升的空间。文献[10]虽然提供较高的检测精度,但未考虑FPS指标。YOLOv10n、文献[11]和YOLO-SG^[13]进一步压缩模型,将参数量控制在 4×10^6 以下,且具有较高的FPS值。其中,文献[11]的FPS甚至达到了294.0,参数量也降低至 1.9×10^6 。但这3个模型的mAP50仅为75%左右,因此mAP50达到90.9%的DTI-YOLO对检测复杂场景下的小目标交通标志更具优势,能够保证检测的全面性和准确性。RT-DETR-R34、文献[8]和BH-RT-DETR^[28]的mAP50能够达到85%以上,但牺牲了一定的参数量。虽然文献[8]和BH-RT-DETR将参数量降至 18×10^6 左右,但仍为DTI-YOLO参数量的3倍左右。文献[29]在检测精度上具备一定竞争力,mAP50达到90.5%,接近DTI-YOLO,但其FPS只有39.4,实时性较差,难以满足实际需求。本文提出的DTI-YOLO在多个评价指标上表现最优。DTI-YOLO的mAP50达到90.9%,相较于YOLOv10s提升5.5个百分点;精确率和召回率分别提升3.4和8.5个百分点。此外,DTI-YOLO的模型参数量压缩至 5.4×10^6 ,与YOLOv10s相比降低33.3%。同时,DTI-YOLO的推理速度达到115.3FPS,能够在保证高精度的前提下实现实时检测。实验证明,DTI-YOLO通过轻量级跨尺度特征融合架构和膨胀卷积融合膨胀注意力设计,在保持模型高效推理的同时,有效解决了远景和复杂背景下小目标漏检问题。相较于现有模型,其在检测精度、模型参数量和实时性3个维度实现更优平衡,为车载端交通标志实时检测提供了兼具鲁棒性与实用性的检测模型。

YOLOv10s模型和DTI-YOLO模型在TT100K数据集上的平均精度均值曲线对比分别如图9所示,其中图(a)为mAP50曲线对比,图(b)为mAP50-95曲线对比。蓝色曲线代表原始模型YOLOv10s,橙色曲线代表改进模型DTI-YOLO。通过曲线的变化趋势可知,随着训练轮数递增,模型的检测精度趋于稳定。在训练50轮后橙色曲线始终呈现在蓝色曲线上方,表明DTI-YOLO模型具有更优的检测性能。

图10展示了模型改进前后在TT100K数据集上的部分检测效果对比。从与YOLOv10s模型的比较中可以看出,DTI-YOLO模型在检测效果上表现出明显优势。在第一列检测场景下,由于光线充足且视野开阔,模型改进前后均能准确检测出目标标志的类别。然而DTI-YOLO明显对远距离小目标的适应性更强,具备更高的检测精度。第二列的对比图展示了交通量较为密集的道路场景,此时背景噪声和待测目标较多,模型分离目标标志的难度加大。YOLOv10s误将无关的地铁指示牌检测为po类别标志,而DTI-YOLO在复杂背景下未出现误检或漏检的情况,并且在检测pl30类别标志时具备较高的置信度分数。在第三列的检测场景中,拍摄角度发生偏斜,DTI-YOLO能够更好地适应角度变

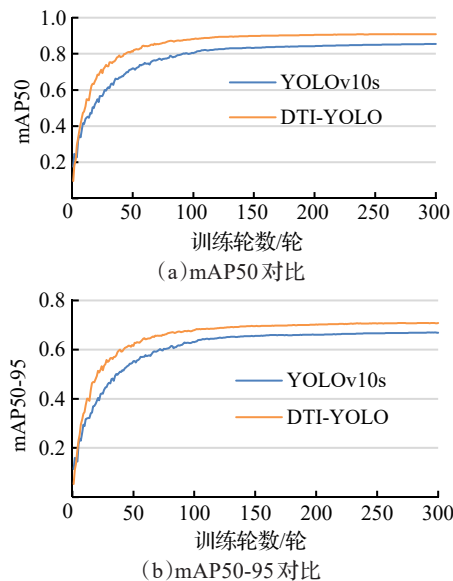


图9 模型改进前后的mAP对比

Fig.9 Comparison of mAP before and after model improvement. DTI-YOLO provides higher detection accuracy, and the detection confidence score for the p1140 category reached 0.93. The comparison results show that the DTI-YOLO model can effectively deal with complex backgrounds and extreme angles, accurately capture the feature information of traffic signs, and show more excellent detection performance.



图10 模型改进前后在TT100K上的检测效果对比

Fig.10 Comparison of detection effects on TT100K before and after model improvement

3.8 泛化性实验

黑夜或恶劣天气会加大背景噪声的干扰,导致模型对关键目标的定位与特征提取的难度增加,从而容易出现漏检和误检现象。一个良好的交通标志检测模型应该具备应对各种检测场景的能力,为了检验DTI-YOLO模型在复杂交通场景下是否具有高效性及泛用性,在CCTSDB数据集上与原始模型进行对比实验,实验结果如表6和表7所示。

由表6可知,改进模型较原始模型召回率达到80.2%,mAP50达到86.6%,召回率、mAP50和mAP50-95分别提升5.8、4.8和5.0个百分点。

表6 模型改进前后在CCTSDB上的性能对比

Table 6 Performance comparison on CCTSDB before and after model improvement

实验方法	Params/10 ⁶	P/%	R/%	mAP50/%	mAP50-95/%
YOLOv10s	8.0	89.6	74.4	81.8	53.3
DTI-YOLO	5.4	89.8	80.2	86.6	58.3

表7展示了模型改进前后对指示(mandatory)、警告(warning)和禁止(prohibitory)3类标志的检测效果对比。与YOLOv10s相比,DTI-YOLO能够更准确地检测出mandatory和prohibitory类别标志,mAP50分别提升5.1和6.4个百分点;mAP50-95分别提升5.4和7.3个百分点,精确率和召回率也均有改善。

为了更直观地对比模型改进前后在复杂场景下的检测效果,进行可视化分析。图11分别为在雾天、雪天和黑夜的交通场景下,原始YOLOv10s模型和改进DTI-YOLO模型的检测效果图。由各检测效果图可知,改进模型较原始模型具有更优的检测效果。

在雾天环境中,受光照不足的影响交通标志的清晰度较低,YOLOv10s的检测精度不高,并出现了误检prohibitory类别标志的现象。在雪天环境中,大面积白色背景的干扰和雪层的遮挡会减弱交通标志的对比度和可见度,加大特征提取难度。但DTI-YOLO能够关注标志本身,有效放大前景特征并减少背景噪声干扰。在检测warning类别标志时,DTI-YOLO的置信度分数相

较YOLOv10s提升0.35。在黑夜环境中,昏暗的背景加上灯光的反射造成标志模糊不清,导致YOLOv10s难以区分交通标志与周围环境,没有检测出prohibitory类别标志,出现漏检现象,而DTI-YOLO仍能够捕捉并保留充分的交通标志的细节特征,提供精准的定位信息。

4 结束语

为了解决远景下小目标特征不明显且难以与背景分离导致检测精度低的问题,本文对YOLOv10s模型的主干网络、颈部网络和损失函数进行了改进。在特征提取阶段,设计了DDFM模块替代PSA模块,有效分离背景噪声与关键目标,增强了模型的特征提取能力。在特征融合阶段,构建了TDL-CCFN特征融合网络,采用轻量化的CCFM跨尺度融合结构和小目标检测层协同设计,提升了多尺度特征的交互和小目标特征的保留,同时减少了冗余参数。采用InnerMPDIoU损失函数替代CIoU,增强了模型对小目标位置的感知能力,提高了应对不同场景的泛化能力。在TT100K和CCTSDB数据集上,DTI-YOLO模型表现出良好的检测性能,在保持高检测精度的同时实现了模型轻量化,参数量仅为 5.4×10^6 。尽管如此,本文提出的模型仍有优化空间,未来的研究将聚焦于提高检测精度的同时,进一步减少计算开销,加快检测速度,以更好地满足交通场景对实时性的要求,推动模型在复杂环境中的实际应用。

表7 模型改进前后三类标志的检测效果对比

Table 7 Comparison of detection effects of three types of signs before and after model improvement 单位:%

实验方法	mandatory				warning				prohibitory			
	P	R	mAP50	mAP50-95	P	R	mAP50	mAP50-95	P	R	mAP50	mAP50-95
YOLOv10s	87.7	66.9	73.0	49.7	89.2	84.6	88.1	54.6	91.8	71.8	84.4	55.5
DTI-YOLO	87.0	70.2	78.1	55.1	87.5	87.7	90.7	57.4	93.6	82.7	90.8	62.8



图11 模型改进前后在CCTSDB上的检测效果对比

Fig.11 Comparison of detection effects on CCTSDB before and after model improvement

参考文献:

- [1] BARODI A, BAJIT A, ZEMMOURI A, et al. Improved deep learning performance for real-time traffic sign detection and recognition applicable to intelligent transportation systems [J]. International Journal of Advanced Computer Science and Applications, 2022, 13(5): 712-723.
- [2] 陈飞, 刘云鹏, 李思远. 复杂环境下的交通标志检测与识别方法综述[J]. 计算机工程与应用, 2021, 57(16): 65-73.
CHEN F, LIU Y P, LI S Y. Survey of traffic sign detection and recognition methods in complex environment[J]. Computer Engineering and Applications, 2021, 57(16): 65-73.
- [3] 赵磊, 李栋. PMM-YOLO: 多尺度特征融合的交通标志检测算法[J]. 计算机工程与应用, 2025, 61(4): 262-271.
ZHAO L, LI D. PMM-YOLO: traffic sign detection algorithm with multi-scale feature fusion[J]. Computer Engineering and Applications, 2025, 61(4): 262-271.
- [4] CAI Z W, VASCONCELOS N. Cascade R-CNN: delving into high quality object detection[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE, 2018: 6154-6162.
- [5] LIU W, ANGUELOV D, ERHAN D, et al. SSD: single shot multiBox detector[C]//Proceedings of the European Conference on Computer Vision. Cham: Springer, 2016: 21-37.
- [6] ZHAO Y A, LYU W Y, XU S L, et al. DETRs beat YOLOs on real-time object detection[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE, 2024: 16965-16974.
- [7] 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. Piscataway: IEEE, 2016: 779-788.
- [8] 刘思元, 高凯, 雍龙泉. 改进 RT-DETR 的航拍小目标检测算法[J]. 计算机工程与应用, 2025, 61(4): 272-281.
LIU S Y, GAO K, YONG L Q. Improved RT-DETR algorithm for aerial small object detection[J]. Computer Engineering and Applications, 2025, 61(4): 272-281.
- [9] SONG X, FAN B, LIU H, et al. HPRT-DETR: a high-precision real-time object detection algorithm for intelligent driving vehicles[J]. Sensors (Basel), 2025, 25(6): 1778.
- [10] MU J H, SU Q H, WANG X Y, et al. A small object detection architecture with concatenated detection heads and multi-head mixed self-attention mechanism[J]. Journal of Real-Time Image Processing, 2024, 21(6): 184.
- [11] 朱硕, 梁吉丰, 孙佳豪, 等. 基于改进 YOLOv5s 的交通标志检测算法[J]. 无线电工程, 2024, 54(12): 2902-2912.
ZHU S, LIANG J F, SUN J H, et al. Traffic sign detection algorithm based on improved YOLOv5s [J]. Radio Engineering, 2024, 54(12): 2902-2912.
- [12] 罗向龙, 吕温馨, 石镇岳, 等. 改进 YOLOv8n 的轻量化交通标志检测算法[J]. 激光与光电子学进展, 2025, 62(12): 422-433.
LUO X L, LYU W X, SHI Z Y, et al. Improved lightweight traffic sign detection algorithm for YOLOv8n[J]. Laser & Optoelectronics Progress, 2025, 62(12): 422-433.
- [13] HAN Y J, WANG F P, WANG W, et al. YOLO-SG: small traffic signs detection method in complex scene[J]. The Journal of Supercomputing, 2024, 80(2): 2025-2046.
- [14] WANG A, CHEN H, LIU L H, et al. YOLOv10: realtime end-to-end object detection[J]. arXiv:2405.14458, 2024.
- [15] HU S, GAO F, ZHOU X W, et al. Hybrid convolutional and attention network for hyperspectral image denoising[J]. IEEE Geoscience and Remote Sensing Letters, 2024, 21: 1-5.
- [16] OPLE J J M, YEH P Y, SUN S W, et al. Multi-scale neural network with dilated convolutions for image deblurring[J]. IEEE Access, 2020, 8: 53942-53952.
- [17] JIAO J Y, TANG Y M, LIN K Y, et al. DilateFormer: multi-scale dilated transformer for visual recognition[J]. IEEE Transactions on Multimedia, 2023, 25: 8906-8919.
- [18] ZHENG Z H, WANG P, LIU W, et al. Distance-IoU loss: faster and better learning for bounding box regression[C]//Proceedings of the AAAI Conference on Artificial Intelligence, 2020: 12993-13000.
- [19] ZHANG H, XU C, ZHANG S J. Inner-IoU: more effective intersection over union loss with auxiliary bounding box[J]. arXiv:2311.02877, 2023.
- [20] MA S L, XU Y. MPDIoU: a loss for efficient and accurate bounding box regression[J]. arXiv:2307.07662, 2023.
- [21] ZHU Z, LIANG D, ZHANG S H, et al. Traffic-sign detection and classification in the wild[C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE, 2016: 2110-2118.
- [22] ZHANG J, ZOU X, KUANG L D, et al. CCTSDB 2021: a more comprehensive traffic sign detection benchmark[J]. Human-Centric Computing and Information Sciences, 2022, 12: 22967.
- [23] TAN M X, PANG R M, LE Q V. EfficientDet: scalable and efficient object detection[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE, 2020: 10778-10787.
- [24] LI H L, LI J, WEI H B, et al. Slim-neck by GSConv: a better design paradigm of detector architectures for autonomous vehicles[J]. arXiv:2206.02424, 2022.
- [25] REIS D, KUPEC J, HONG J, et al. Real-time flying object detection with YOLOv8[J]. arXiv:2305.09972, 2023.
- [26] WANG C Y, YEH I H, MARK LIAO H Y. YOLOv9: learning what you want to learn using programmable gradient information[C]//Proceedings of the European Conference on Computer Vision. Cham: Springer Nature Switzerland, 2025: 1-21.
- [27] TIAN Y, YE Q, DOERMANN D. YOLOv12: attention-centric real-time object detectors[J]. arXiv:2502.12524, 2025.
- [28] 秦伦明, 张云起, 崔昊杨, 等. 基于改进 RT-DETR 的极端天气下交通标志检测方法[J]. 电子测量技术, 2025, 48(9): 56-64.
QIN L M, ZHANG Y Q, CUI H Y, et al. Traffic sign detection method in extreme weather based on improved RT-DETR[J]. Electronic Measurement Technology, 2025, 48(9): 56-64.
- [29] 高翊轩, 李昕, 刘婧彤. 改进 YOLOv5 的小目标交通标志检测方法[J]. 计算机工程与设计, 2024, 45(12): 3639-3647.
GAO Y X, LI X, LIU J T. Improved YOLOv5 small target traffic sign detection method[J]. Computer Engineering and Design, 2024, 45(12): 3639-3647.