

基于 YOLOv5 的雾霾天气下交通标志识别^{*}朱开^{1,2} 陈慈发^{1,2}

(1. 三峡大学计算机与信息学院 宜昌 443002; 2. 湖北省建筑质量检测装备工程技术研究中心 宜昌 443002)

摘要: 针对雾霾天气下道路交通标志识别难度大、精确度较低的问题,提出一种基于 YOLOv5 的雾霾天气交通标志识别模型。首先在 YOLOv5 原始模型上融入卷积注意力机制,在空间维度和通道维度上进行特征增强,抑制雾霾天气对模型的干扰;然后将 BiFPN 作为 neck 层中的特征融合结构,更加充分地融合多尺度特征,减少目标信息丢失;并选用 CIoU 作为 YOLOv5 的损失函数提高定位能力;使用 K-means 聚类算法在 TT100K 和 CODA 数据集重新获取锚框值,加快模型收敛速度。实验结果表明,改进后模型识别精度达到 92.5%,比 YOLOv5 提升 5.6%,在雾霾天气下仍能准确识别交通标志,速度达 27 FPS,能够进行实时检测。

关键词: 目标检测;交通标志识别;YOLOv5;注意力机制

中图分类号: TP391.4 **文献标识码:** A **国家标准学科分类代码:** 520.2040

Traffic sign recognition under fog weather based on YOLOv5

Zhu Kai^{1,2} Chen Cifa^{1,2}

(1. College of Computer and Information, China Three Gorges University, Yichang 443002, China;

2. Hubei Province Engineering Technology Research Center for Construction Quality Testing Equipment, Yichang 443002, China)

Abstract: Aiming at the problem of high difficulty and low accuracy in road traffic sign recognition under haze weather, a traffic sign recognition model based on YOLOv5 was proposed. Firstly, the convolutional attention mechanism was integrated into the original YOLOv5 model to enhance features in the spatial dimension and channel dimension to suppress the interference of haze weather on the model. Then, BiFPN is used as the feature fusion structure in neck layer to more fully fuse multi-scale features and reduce the loss of target information. CIoU is used as the loss function of YOLOv5 to improve the positioning ability. K-means clustering algorithm was used to re-obtain anchor frame values in TT100K and CODA datasets to accelerate the convergence speed of the model. The experimental results show that the recognition accuracy of the improved model reaches 92.5%, which is 5.6% higher than that of YOLOv5, and it can still accurately identify traffic signs in haze weather, and the speed can reach 27 FPS, which can be used for real-time detection.

Keywords: target detection; traffic sign recognition; YOLOv5; attention mechanism

0 引言

现有的交通标志识别任务大多是在简单场景下进行,而雾霾天气下标志牌的图像信息模糊,使得难以对其定位和识别。目前的交通标志检测模型有着不错的识别效果,但是应用于雾霾场景中时往往会出现大量漏检且准确率不高。因此,研究雾霾天气下兼顾实时性与准确性交通标志识别方法就显得十分重要。传统的交通标志检测方法主要是根据目标的颜色和形状对其进行分类。戴学瑞等^[1]提出的基于颜色和最大稳定极值算法能够消除光照对目标特征

的影响。乔敏等^[2]提出的基于道路维纳复原的方法能对六类交通标志进行检测。胡聪等^[3]依靠区域推荐算法形成推荐区域等方法消除了光照对检测的影响。近年来基于深度学习的交通标志检测效果表现突出,可以有效解决传统方法中存在检测速度慢、泛化能力弱、易受周围环境和天气干扰等缺点。基于深度学习的目标检测算法可以分为两类,一种是以 Mask R-CNN^[4]、Faster R-CNN^[5]为代表的二阶段检测算法,这类算法根据目标特征提取预选框,在此基础上用卷积神经网络进行分类和定位得到检测结果,二阶段算法具有较高的准确率,但结构复杂,检测速度难以满足实

收稿日期:2022-08-04

^{*} 基金项目:国家自然科学基金新疆联合基金重点项目(U1703261)资助

际应用需求。另一种是以 YOLO^[6]、SSD^[7] 为代表的一阶段算法,这类算法取消了目标预选框的提取,直接对被测物体进行定位和预测分类。一阶段算法拥有较快检测速度的同时,保持着不错准确率。闫钧华等^[8]在 YOLO 模型上实现了小目标检测。吕禾丰等^[9]使用改进 YOLOv5 进行交通标志检测。Liu 等^[10]提出的 TsingNet 网络通过构建双边特征金字塔学习子网的特征,可以有效的检测交通标志。张上等^[11]对 YOLOv5 模型进行 FPGM 剪枝,引入注意力机制,有效减小了模型体积,剪枝后网络在 GTSRB 数据集上准确率达 92.64%。张达为等^[12]提出了改进的 YOLOv3 算法,能够在复杂环境下完成对交通标志的检测。王文胜等^[13]设计了基于 YOLOv5 模型进行交通标志识别的智能小车。解宇虹等^[14]通过合成有雾的数据集来验证先验知识对目标检测的影响。上述方法在雾霾天气下存在检测速度慢、识别效果差、识别种类少和特征提取能力差等问题,不适合应用于实际的目标检测任务中。针对这些问题,本文提出了一种基于 YOLOv5 的雾霾天气下交通标志识别模

型,消除了雾霾天气对模型的影响,主要工作如下:

1) 在模型中引入注意力机制加强特征提取能力,抑制雾霾天气对模型的干扰;改进特征金字塔充分融合多尺度特征;优化损失函数提升定位能力。

2) 对 TT100K 和 CODA 数据集进行预处理,使用 K-means 聚类算法在数据集上重新获取初始锚框值。

3) 对改进模型进行验证,与其他主流算法和原模型进行对比实验,证明本文方法先进性。设置消融实验验证不同改进方法对模型性能影响。选取部分样本展示改进算法的检测效果。

1 YOLOv5 算法原理

YOLOv5^[15] 是 YOLO 系列最新目标检测算法,具有结构灵活、检测速度快等优点。按照大小可分为 YOLOv5s、YOLOv5m、YOLOv5l、YOLO5x 四个版本。网络结构如图 1 所示,主要由如下部分组成:Input 输入端、Backbone 主干网络、Neck 颈部网络、Head 输出端。

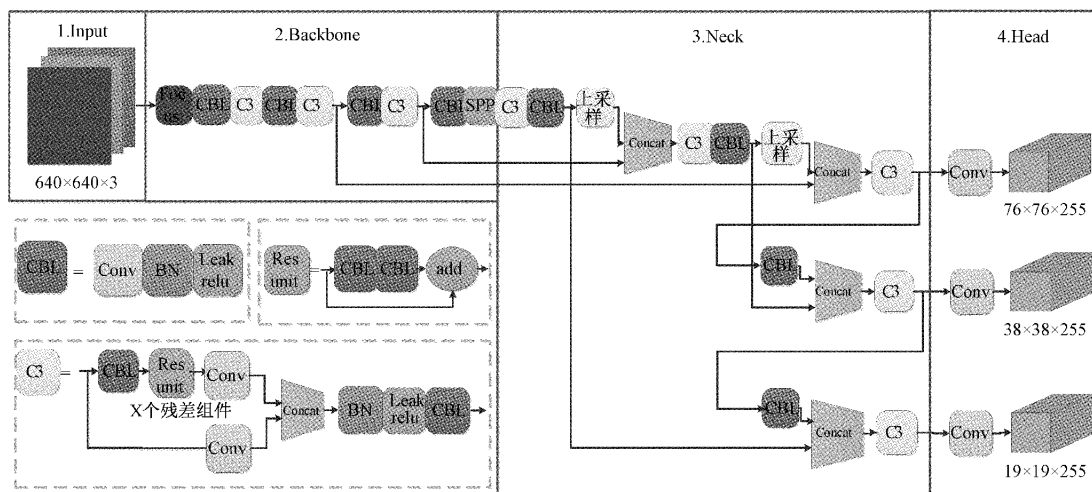


图 1 YOLOv5 网络结构

Input 输入端完成了对图像的预处理,包括 Mosaic 数据增强、自适应锚框计算、自适应图像填充。Backbone 主干网络用于特征提取,网络的开头为 Focus 模块,如图 2 所示:输入的 $4 \times 4 \times 3$ 图片长与宽缩减一半,通道则扩充成 4 倍,进行卷积、池化等操作后得到特征图。CBL 由卷积、归一化和 SiLU 激活函数 3 部分组成,C3 包含的 CBL 和残差模块能够记录梯度变化,增强网络学习能力。SSP 为空间金字塔池化层,对于不同尺度的特征能够提取出固定大小的特征向量。Neck 层用于特征融合,包含特征金字塔 (feature Pyramid networks, FPN)^[16]、路径聚合网络 (path aggregation network, PAN)^[17]。FPN 把高层的语义信息传向低层,在 FPN 层的基础上添加 PAN 层用于自底而上传递定位特征。Head 输出端为最后的预测部分,用于生成目标图像的预选框。

由于原始 YOLOv5 实验的 COCO 数据集多为大目

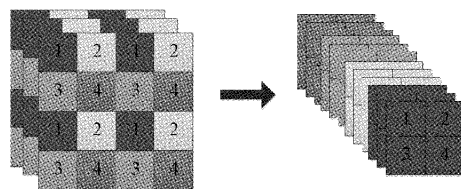


图 2 Focus 切片操作

标,检测远距离小目标和模糊目标时效果不佳,因此本文选择 YOLOv5s 作为基础模型并对其进行改进,得到更加适用于雾霾天气的交通标志检测模型。

2 改进 YOLOv5 网络

2.1 融合注意力机制

注意力机制^[18]能获取中特征图更关键的信息并忽略不相关信息,本文在主干网络中融入卷积注意力机制

(convolutional block attention module, CBAM), 其结构如图3所示, 包括通道注意力和空间注意力两个子模块。

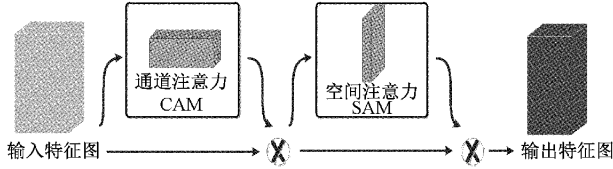


图3 卷积注意力机制

对于卷积网络中某一层特征图 F , 先在通道注意力中推导出特征图 M_c 。将 M_c 和输入图 F 相乘生成 F' 后送入空间注意力机制中生成特征图 M_s , M_s 与 F' 相乘得到 F'' 。其计算公式如下, 其中 $\otimes$ 表示逐元素相乘。

$$F' = M_c(F) \otimes F \quad (1)$$

$$F'' = M_s(F') \otimes F' \quad (2)$$

通道注意力模块如图4所示, 其作用是关注特征图中更加关键的信息。

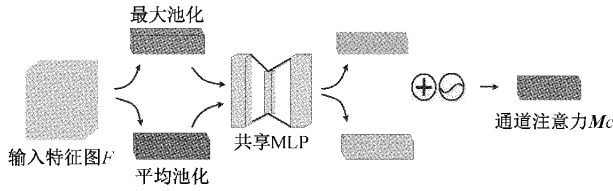


图4 通道注意力模块

将输入的特征图 F 进行通道池化, 送入多层感知机进行运算后, 对特征向量进行逐元素相加和 Sigmoid 函数激活, 最终得到通道注意力向量 M_c 。计算公式如式(3)所示, σ 为 Sigmoid 激活函数, W_0 和 W_1 为 MLP 中的隐藏层和输出层权重, F_{avg}^c 和 F_{max}^c 表示全局平均池化特征和最大池化特征。

$$M_c(F) = \sigma(MLP(AvgPool(F)) + MLP(MaxPool(F))) = \sigma(W_1(W_0(F_{avg}^c)) + W_1(W_0(F_{max}^c))) \quad (3)$$

特征图经过通道注意力后会丢失部分信息, 而空间注意力模块关注特征图的位置信息, 图5展示了其计算过程。

对特征图 F' 进行最大值池化和平均值池化, 将得到的两个特征向量进行拼接和卷积操作, 并使用 Sigmoid 激活函数生成空间注意力向量 M_s 。计算公式如式(4)所示, σ

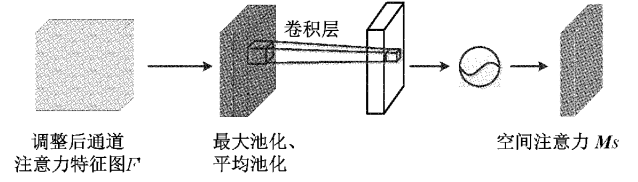
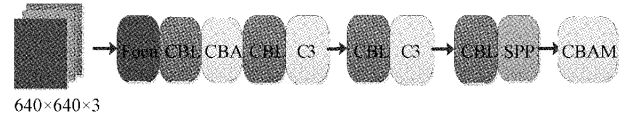


图5 空间注意力模块

为 Sigmoid 激活函数, $f^{7 \times 7}$ 为 7×7 卷积核的卷积层, F_{avg}^s 和 F_{max}^s 表示平均池化特征和最大池化特征。

$$M_s(F) = \sigma(f^{7 \times 7}([AvgPool(F); MaxPool(F)])) = \sigma(f^{7 \times 7}([F_{avg}^s; F_{max}^s])) \quad (4)$$

CBAM 加入位置的不同会产生不同的效果, 实验证明, 本文在主干网络尾部加入 CBAM 并替换掉第一个 C3 结构, 可以加强网络的特征提取能力, 提升模型检测准确率, 融入 CBAM 后的主干网络结构如图6所示。



征。最后,将每条双向路径融和进一个特征网络中,把删除节点的输入边同时连接到下一个节点和下一层的节点上,实现更高层次的特征融合,如图 7(c)所示。

2.3 优化损失函数

YOLOv5 使用 GIoU_Loss 作为边界框损失函数来评判预测边界框和真实边界框的距离。

$$GIoU = IoU - \frac{|A_c - U|}{A_c} \quad (5)$$

式中: IoU 为交并比, U 为并集面积, A_c 为真实框与预测框的最小外接矩形面积。GIoU_loss 可以有效解决 IoU_loss 中当预测框和目标框不相交时,无法衡量两者位置关系导致无法优化的问题。而目标框完全包含预测框时,GIoU_loss 无法两者的位置关系,值仍不会改变,此时 GIoU_loss 退化为 IoU_Loss。针对这一问题,本文采用 CIoU_loss 作为预测框的损失函数,公式如下:

$$CIoU = IoU - \frac{\rho^2(b, b^{gt})}{C^2} - \alpha v \quad (6)$$

式中: α 表示权重系数, v 表示真实框与预测框之间的相似性。计算公式如下:

$$v = \frac{4}{\pi^2} \left(\arctan \frac{\omega^{gt}}{h^{gt}} - \arctan \frac{\omega}{h} \right)^2 \quad (7)$$

$$\alpha = \frac{v}{(1 - IoU) + v} \quad (8)$$

式中: $\frac{\omega^{gt}}{h^{gt}}$ 和 $\frac{\omega}{h}$ 分别表示真实框和预测框的宽高比。和 GIoU_Loss 相比 CIoU_Loss 着重考虑边界框宽高比的尺度,解决了原 YOLO 模型损失函数中目标框包含预测框时损失值不变的问题,使得收敛速度变快,回归更加准确,增强模型的定位能力。

3 实验及分析

3.1 数据集的构建与处理

本文实验数据集包括两部。一部分是清华大学和腾讯联合制作的数据集 TT100K,包含了 221 个类别交通标志,图像的分辨率为 2048×2048 。由于部分类别数量较少不利于进行模型训练,因此选取其中的 45 个类别标志 9 300 张样本进行实验。TT100K 数据集中缺少雾霾天气下的图像,本文采用 Python 的第三方库 pillow 对数据集进行雾霾增强,效果如图 8 所示。

另一个部分选用的是华为制作的 CODA 数据集。该数据集由 1 500 个精心挑选的真实世界驾驶场景组成,本文选择其中的 700 张有雾图像进行实验。为了满足模型训练要求,需保持两个数据集格式的一致性,用 Labeling 软件将 CODA 数据集与 TT100 同类型的标志设置成一样的标签。如图 9 所示,将靠右侧行驶标志标注成 i5,禁鸣标志标注成 p11。TT100K 中包含了 CODA 数据集中所有标志的类别,故实验仍为 45 个类别。



(a) 原图

(b) 雾霾增强后图像

图 8 雾霾增强图像



图 9 图像标注

改进后的 YOLOv5 网络对数据集进行饱和度调整,亮度调整以及 Mosaic 数据增强。Mosaic 数据增强从样本中任意选择 4 张图片,进行随机翻转,放缩变换等操作,再拼接成和原始图像尺寸相同的新图像。Mosaic 数据增强使目标的数量得到扩充,可以同时训练 4 张图像,提高网络训练速度。

3.2 重获锚框值

在目标识别任务中,初始锚框包含目标的先验知识,锚框的参数设置会对影响模型检测的速度和精度。本文数据集中目标尺寸更小,原模型对 COCO 数据集聚类生成的 9 种锚框参数不适合作为实验初始锚框。使用 K-means 聚类算法在数据集上重获锚框值,具体参数如表 1 所示。

表 1 锚框值

锚框值	特征层
(4,5),(6,7),(9,10)	小目标层
(11,21),(12,13),(15,16)	中目标层
(20,22),(29,31),(46,47)	大目标层

3.3 实验环境与评价指标

在本文实验中,将样本按照 8:2 划分为训练集和测试集。模型使用 640×640 和 1280×1280 两种尺度分别进行训练。运行环境是 Windows10 操作系统, NVIDIA GeForce RTX 2070 显卡, Pytorch 深度学习框架,

Python3.7 编程语言。实验初始参数设置如表 2 所示。

表 2 实验参数

参数名称	参数值
学习率	0.01
迭代次数	300
批次大小	8
衰减系数	0.000 5
动量因子	0.8

为评估本文改进后 YOLOv5 网络的有效性,使用平均精度均值(mean average precision, mAP)来衡量模型检测结果。对某个类别的平均精度(average precision, AP),求均值得到 mAP。公式如式(9)所示,其中 P 代表准确率, R 代表召回率, C 表示所有的类别数。

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

$$mAP = \frac{1}{C} \sum_{i=1}^C AP_i \quad (9)$$

另一个评价指标为每秒检测图片帧数(frames per second, FPS)。在实际应用场景中 FPS 值到达 25 可以进行实时检测。

3.4 模型训练

图 10 展示了本文模型训练过程中的损失值变化情况,随着训练轮数的增加损失值逐渐下降,在 250 轮时趋于平稳。mAP 的变化曲线如图 11 所示,前 100 轮 mAP 值上升较快,150 轮后缓慢上升并且稳定在 0.924 左右。从以上分析可以看出,模型在训练过程中效果越来越优。

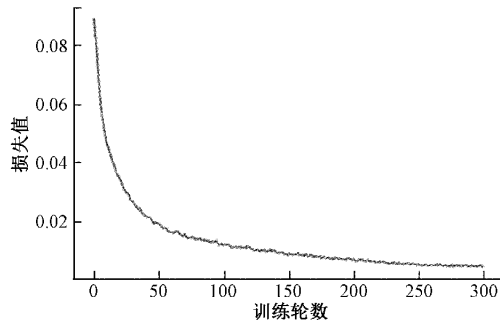


图 10 损失值曲线

3.5 实验结果分析

为了验证改进算法的先进性,与目前主流的算法在交通标志检测任务中进行对比,表 3 为不同算法在 mAP 和 FPS 的对比结果。在训练尺度为 640×640 时,改进 YOLOv5 的 mAP 值达到 80.7%,检测速度 27 FPS 满足实时检测任务的需求。与两阶段算法的 Faster RCNN 相比, mAP 提升了 8.2%,检测速度提升了 24 FPS。与一阶段的 YOLOv3、YOLOv4、RetinaNet 相比, mAP 值分别提高了 9.3%、6.1% 和 7.9%,检测速度也有一定优势,与

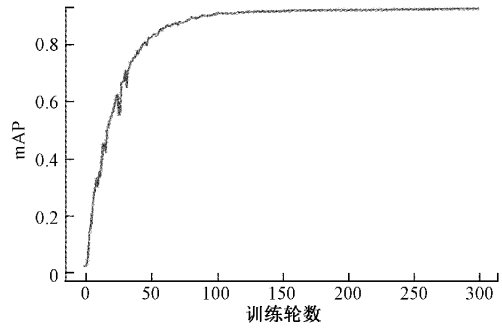


图 11 mAP 曲线

SSD 算法相比,虽然在检测速度低了 15 FPS,但是在检测精度上提高了 17%。综上所述,与其他算法相比,改进后 YOLOv5 对交通标志识别的综合性能最好。

表 3 改进模型与其他模型对比

模型	输入尺寸/px	mAP/%	FPS/s ⁻¹
FasterR-CNN	640×640	72.5	3
YOLOv3	640×640	71.4	20
YOLOv4	640×640	74.6	22
RetinaNet	640×640	72.8	16
SSD	640×640	63.7	42
改进 YOLOv5	640×640	80.7	27

为进一步验证所提出算法的优势,分别在 640×640 和 1280×1280 两种输入尺度上与原 YOLOv5 进行对比,结果如表 4 所示。改进 YOLOv5 与原算法相比,在输入尺寸为 640×640 时平均检测精度提升了 6.4%,速度下降 2 FPS,仍能满足实时检测需求。输入尺寸为 1280×1280 时,提升了 5.6%。在检测速度上改进 YOLOv5 与原 YOLOv5 相比下降 1 FPS。综上所述,改进 YOLOv5 在不同输入尺度对雾霾天气下交通标志的识别精度均高于原始模型。

表 4 改进模型与原模型对比

模型	输入尺寸/px	mAP/%	FPS/s ⁻¹
YOLOv5	640×640	74.3	29
改进 YOLOv5	640×640	80.7	27
YOLOv5	1280×1280	86.9	9
改进 YOLOv5	1280×1280	92.5	8

此外本文在输入尺度为 1280×1280 下进行消融实验,在 YOLOv5 模型的基础上分别融入卷积注意力、改进特征金字塔和优化损失函数,以此来验证不同改进方法对模型影响,实验结果如表 5 所示。融入卷积注意力后 mAP 值提升 1.7%,优化 YOLOv5 特征金字塔后 mAP 值提升 2.2%,优化 YOLOv5 损失函数 mAP 值提升 0.6%,三类改进方法同时应用于 YOLOv5 时, mAP 值提升 5.6%。

综上所述,本文改进方法都能提升检测精度,三者结合提升效果最佳。

表 5 消融实验

模型	CBAM	BiFPN	CIoU	mAP/%
YOLOv5	×	×	×	86.9
优化模型 1	✓	×	×	88.6
优化模型 2	×	✓	×	89.1
优化模型 3	×	×	✓	87.5
本文模型	✓	✓	✓	92.5

为了更加直观展示改进算法的检测效果,分别选取简

单目标、远距离小目标和密集目标对原 YOLOv5 和本文模型进行验证,结果如下所示。从图 12 中可以看出两种模型都能够识别出目标,YOLOv5 的置信度分别为 0.92 和 0.69,而本文模型略高于 YOLOv5 分别为 0.94 和 0.91。图 13 是对远距离小目标的检测,原模型没能识别出目标,而改进模型不仅能有效识别出小目标且保持着 0.81 的置信度。在图 14 对密集目标检测对比图中,目标特征模糊,检测难度大,原模型漏检严重且置信度较低,而本文模型能识别出更多的目标,置信度也高于原模型。综上所述,YOLOv5 模型在雾霾天气下识别精度低,对远距离小目标和密集目标检测效果差,而改进后 YOLOv5 检测更加精确,在雾霾天气中具备更优越的性能。

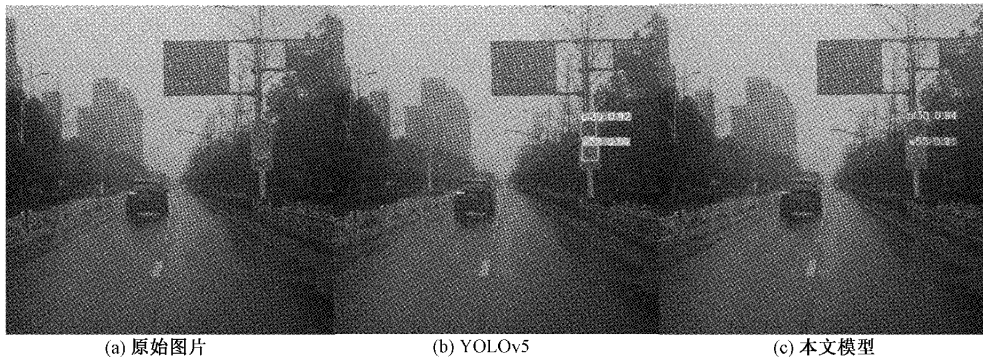


图 12 简单目标检测效果对比图

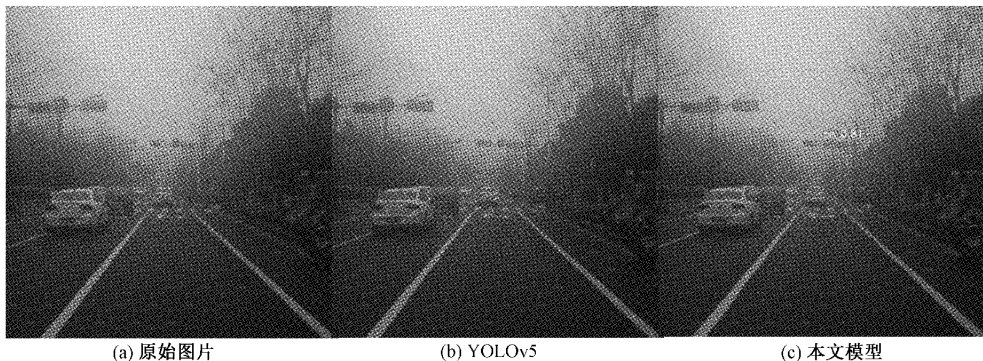


图 13 远距离小目标检测效果对比图

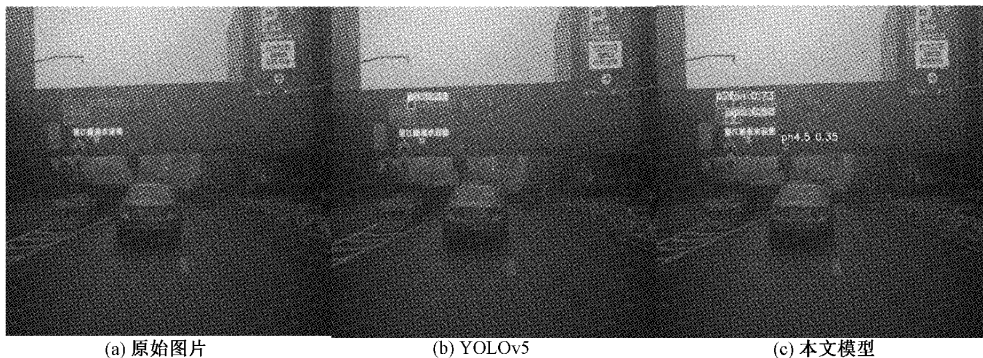


图 14 密集目标检测效果对比

4 结 论

本文针对YOLOv5模型在雾霾天气下交通标志识别中存在精度低、小目标漏检严重、目标特征模糊检测难度大等问题,通过引入注意力机制使其更加专注有效特征,采用BiFPN作为特征融合结构实现高效融合,使目标特征更有效的表达,CIoU作为损失函数来提升目标定位能力,并重新设置初始锚框值。实验结果表明,改进后YOLOv5在两种输入尺度下的准确率分别提升了6.4%和5.6%,速度满足交通标志实时检测要求。但是改进算法在部分类别的检测准确率上仍然不高,因此下一步本文将对网络做进一步优化,使其在雾霾天气下的检测达到更好效果。

参考文献

- [1] 戴雪瑞,袁雪,乐国庆,等.复杂场景下基于颜色和MSER的交通标志检测方法[J].北京交通大学学报,2018,42(5):107-115.
- [2] 乔敏,孙国强.基于维纳复原的道路限速交通标志检测[J].软件导刊,2020,19(5):234-237.
- [3] 胡聪,何晓晖,邵发明,等.基于极大极稳定区域及SVM的交通标志检测[J].计算机科学,2022,49(S1):325-330.
- [4] COUTEAUX V, SIMOHAMED S, NEMPONT O, et al. Automatic knee meniscus tear detection and orientation classification with Mask-RCNN [J]. Diagnostic and Interventional Imaging, 2019, 100(4): 235-242.
- [5] QIAO L, ZHAO Y, LI Z, et al. Defrcn: Decoupled faster R-CNN for few-shot object detection [C]. Proceedings of the IEEE/CVF International Conference on Computer Vision, 2021: 8681-8690.
- [6] CHENG L, LI J, DUAN P, et al. A small attentional YOLO model for landslide detection from satellite remote sensing images[J]. Landslides, 2021, 18(8): 2751-2765.
- [7] SHA Z, CAI Z, TRAHAY F, et al. Unifying temporal and spatial locality for cache management inside SSDs [C]. DATE 2022: 25th Design, Automation and Test in Europe, IEEE, 2022: 1-6.
- [8] 闫钧华,张琨,施天俊,等.融合多层次特征的遥感图像地面弱小目标检测[J].仪器仪表学报,2022,43(3):221-229,DOI:10.19650/j.cnki.cjsi.J2108699.
- [9] 吕承丰,陆华才.基于YOLOv5算法的交通标志识别技术研究[J].电子测量与仪器学报,2021,35(10):137-144,DOI:10.13382/j.jemi.B2104449.
- [10] LIU Y, PENG J, XUE J H, et al. TSingNet: Scale-aware and context-rich feature learning for traffic sign detection and recognition in the wild [J]. Neurocomputing, 2021, 447: 10-22.
- [11] 张上,王恒涛,冉秀康.基于YOLOv5的轻量化交通标志检测方法[J].电子测量技术,2022,45(8):129-135,DOI:10.19651/j.cnki.emt.2108726.
- [12] 张达为,刘绪崇,周维,等.基于改进YOLOv3的实时交通标志检测算法[J].计算机应用,2022,42(7):2219-2226.
- [13] 王文胜,李继旺,吴波,等.基于YOLOv5交通标志识别的智能车设计[J].国外电子测量技术,2021,40(10):158-164,DOI:10.19652/j.cnki.femt.2102913.
- [14] 解宇虹,谢源,陈亮,等.真实有雾场景下的目标检测[J].计算机辅助设计与图形学学报,2021,33(5):733-745.
- [15] YANG G, FENG W, JIN J, et al. Face mask recognition system with YOLOV5 based on image recognition [C]. 2020 IEEE 6th International Conference on Computer and Communications (ICCC), IEEE, 2020: 1398-1404.
- [16] LIN T Y, DOLLAR 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.
- [17] 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.
- [18] FUKUI H, HIRAKAWA T, YAMASHITA T, et al. Attention branch network: Learning of attention mechanism for visual explanation[C]. Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, 2019: 10705-10714.

作者简介

朱开,硕士研究生,主要研究方向为人工智能、计算机视觉等。

E-mail:582928797@qq.com

陈慈发(通信作者),教授、研究员,主要研究方向为嵌入式系统、物联网、计算机测控系统。

E-mail:chcf0415@126.com