

DOI:10.16136/j.joel.2022.12.0149

改进 YOLOv5 的车牌检测算法在林区中的应用

朱文超, 杨洁*, 卢成煜, 何超

(西南林业大学 机械与交通学院, 云南 昆明 650224)

摘要:针对林区环境中现有的交通监控系统目标检测算法在雾、雨、雪等恶劣天气条件下车牌定位困难、精度低和检测速度慢等问题,提出了一种新的车牌检测方法。该方法以 YOLOv5 (you only look once v5) 为基础模型,采用 K-means++ 的方法对实例标签信息进行聚类分析获取新的初始化锚框尺寸,在特征提取网络中融入 CBAM(convolutional block attention module)注意力机制提取到检测目标更多的特征信息,选取了 $CIoU$ 作为损失函数提高检测框定位精度。在预处理方面,模拟摄像头在采集图像时可能产生的干扰,使用 OpenCV-Python 编写脚本对图像进行处理,增加算法在林区复杂环境下检测的鲁棒性。实验分析表明,该方法的均值平均精度@0.5(mean average precision@0.5, $mAP@0.5$) 达 99.5%、均值平均精度@0.5:0.95($mAP@0.5:0.95$) 达 86.7%、检测速度达 128 帧/s、模型大小仅 14 M,与 YOLOv5 以及其他主流目标检测算法相比有更好的准确性、实时性和广泛可部署性。

关键词:林区防护;神经网络;YOLOv5;车牌检测;复杂天气环境

中图分类号: TP391.4;U495 **文献标识码:** A **文章编号:** 1005-0086(2022)12-1271-09

Application of improved YOLOv5 license plate detection algorithm in forest region

ZHU Wenchao, YANG Jie*, LU Chengyu, HE Chao

(College of Mechanics and Transportation, Southwest Forestry University, Kunming, Yunnan 650224, China)

Abstract: For the issues such as license plate positioning difficulties, low accuracy, and slow detection speed of the target detection algorithm of current traffic monitoring system in the forest environment in fog, rain, snow and other adverse weather conditions, this paper proposed a new license plate detection method, which used you only look once v5 (YOLOv5) as the base model. Firstly, this paper experimented with the K-means++ method and made cluster analysis of the label information of the instances to obtain the new anchor frame sizes. Next, the convolutional block attention module (CBAM) attention mechanism was incorporated into the feature extraction network to extract more feature information of the detection target. Finally, $CIoU$ was chosen as a loss function to improve the detection frame localisation accuracy. In terms of the pre-processing, the possible interference generated by the camera during image acquisition was simulated and the images were scripted using OpenCV-Python to increase the robustness of the algorithm for detection in complex environments in forest areas. As shown by the experimental analysis results, the mean average precision@0.5($mAP@0.5$) and the mean average precision@0.5:0.95($mAP@0.5:0.95$) of the improved model reached 99.5% and 86.7%, respectively. In addition, the detection speed reached 128 fps, but the model size was only 14 M. Compared with YOLOv5 and other leading target detection algorithms, this method had improved accuracy, real-time performance and broad deployability.

Key words: forest protection; neural network; you only look once v5 (YOLOv5); license plate detection; complex weather environment

* E-mail: 351725623@qq.com

收稿日期: 2022-03-11 修订日期: 2022-04-12

基金项目: 国家自然科学基金(51968065)和云南省教育厅科学研究基金项目(111722038)资助项目

1 引 言

森林是我国重要的自然资源,也是生态文明建设的重要组成部分。由于林区通常位于城市的边缘地带,基础建设相对薄弱,交通监控的智能化不够完善,且林区路况多变、环境复杂,常常伴有雾、风沙和雨雪等极端天气,使现有的交通监控设备难以有效地发挥监控任务,为不法分子潜入林区纵火、偷窃野生动植物提供了有机可乘的场所^[1]。车牌检测是智能交通监控系统中的重要组成部分,对进入林区的车辆进行识别、记录起到关键作用。

目前已有的车牌检测方法主要分为两类:一类是基于人工提取特征的传统车牌检测方法;另一类是基于深度学习的车牌检测算法。传统的车牌检测方法主要包括边缘检测^[2]、颜色特征提取^[3]、纹理特征提取^[4]和几何特征检测等^[5],这些方法单独或者互相结合,通过人工提取车牌特征的手段实现车牌检测,但提取的特征单一,受环境影响大,鲁棒性较差,在林区复杂环境下检测准确率低、检测速度慢。基于深度学习的目标检测算法则利用卷积神经网络(convolutional neural network, CNN)^[6]来替代传统的人工提取特征的方法,使模型适应性增强。与传统的车牌检测方法相比,后者不仅在高速公路、停车场等单一、简单的环境下有良好的检测效果,而且在光线不足、车牌倾斜、图像模糊等不利的条件下也有更稳定的鲁棒性,弥补了传统检测方法的不足,提高了在林区等复杂环境下的车牌检测的效果。目前基于深度学习的目标检测算法主要分为两阶段的目标检测算法(two-stage)和单阶段的目标检测算法(one-stage)^[7]。两阶段的目标检测算法首先是提取包含目标物体在内的候选框,然后再进行校准和分类,单阶段目标检测算法则直接对目标对象进行检测,有更快的检测速度,更适合对目标进行实时检测。

文章提出的车牌检测方法就是以 YOLO(you only look once)系列为基础的单阶段目标检测算法。JAMTSHO 等^[8]提出利用 YOLOv3 对车牌进行定位,提高了检测精度。马巧梅等^[9]在 YOLOv3 的基础上改进了多尺度特征融合并且融入 Inception-SE(squeeze-and-excitation, SE)模块,提高了检测速度和精度。SELMIZ 等^[10]基于 YO-

LO 设计了车牌检测器并获得良好的识别准确率。为了在进一步提升检测器性能的同时降低检测算法部署所依赖的硬件环境,文章使用 YOLOv5 作为车牌检测的基础算法,结合车牌颜色、几何形状等特点,并针对易发生恶劣天气等对检测有更多不利因素的林区进行了算法改进,提高了其适应性和准确性。

2 YOLOv5 目标检测算法

YOLO 算法是由 REDMON 等^[11]于 2016 年提出的一种新的单阶段目标检测算法,与其他算法相比,在同等的尺寸下性能更强,并且也有很好的稳定性。从 YOLOv1 至今,YOLO 系列已经发展到 v5^[11-14],凭借着不断的创新和完善,在吸取了之前版本和其他算法的优点之后,v5 改变了之前 YOLO 目标检测算法的检测速度快但精度不高的缺点,在检测准确度以及实时性上都有所提升^[15],能够满足车牌检测所需的实时性和准确性要求。

此外,YOLOv5 通过调整 depth multiple 和 width multiple 控制网络的深度和宽度,设置了 5 种模型满足不同使用场景的需求,分别是 YOLOv5n、YOLOv5s、YOLOv5m、YOLOv5l、YOLOv5x,其中 YOLOv5s 在满足良好的检测精度和速度要求的情况下,对硬件环境的依赖较小,使检测算法能够广泛部署。实验选择 YOLOv5s 作为车牌检测的基础模型,改进其网络结构,增强模型在林区多雾、雨雪等恶劣环境下对车牌的检测能力。

YOLOv5s 的网络模型主要由 Input、Backbone、Neck 和 Prediction 4 个部分组成,整体的网络结构如图 1 所示。首先,在 Input 阶段,通过加入 Mosaic 数据增强方式扩展数据集;Backbone 阶段,通过 Focus 结构把尺度为 $608 \times 608 \times 3$ 的样本切片拼接成 $304 \times 304 \times 12$,融入 CSP 结构,再经过 32 个卷积核的卷积运算转换成 $304 \times 304 \times 32$ 的特征图;在 Neck 阶段,使用了 FPN(feature pyramid networks) + PAN(path aggregation network)结构,先完成自深层向浅层的特征传递,再自浅层向深层的传递特征,完成了不同层的特征融合;Prediction 阶段,输出 3 个尺度的特征图,分别为 19×19 、 38×38 、 76×76 的网格,对应检测大、中、小 3 种不同尺寸的物体,以每个网格为中心生成 3 个预测框,通过 NMS(non-maximum suppression)对预测框进行筛选,保留置信度最高的预测框信息,作为最后的检测结果。

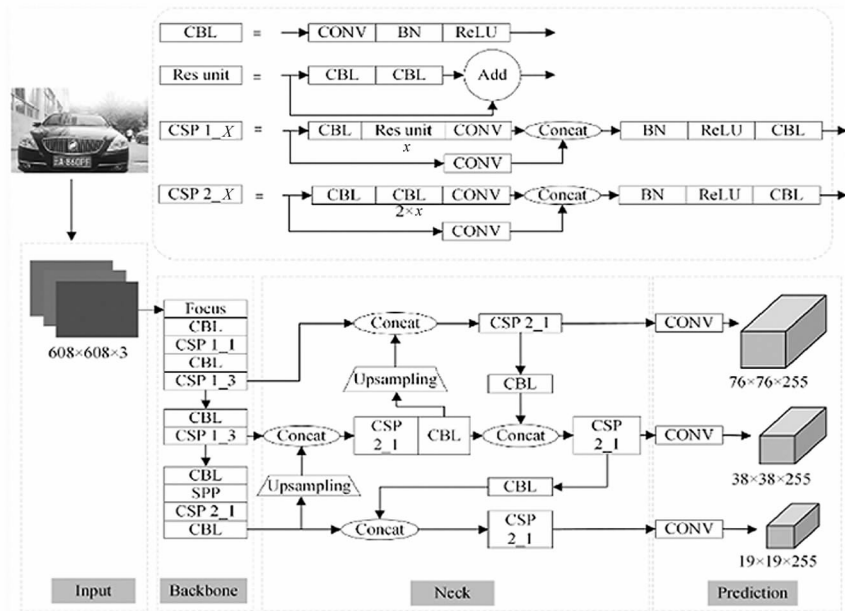


图 1 YOLOv5s 的网络结构
Fig. 1 YOLOv5s network structure

3 改进 YOLOv5

3.1 改进初始锚框参数

YOLOv5 在每种尺度的特征图上分别初始化了 3 个不同固定宽高的锚框,初始的锚框设定对网络在训练过程中的收敛情况和最终目标检测效果产生较大影响。原始的锚框参数由 COCO 数据集聚类产生,由于该数据集的目标种类有 80 个分类,物体种类丰富,尺寸大小跨度较大,导致其聚类产生的初始化锚框有较强的普遍性。本次研究的检测任务只有车牌一个类别,通过对数据集的标注信息进行分析,从图2中可以看出数据集中物体的尺寸分布较为集

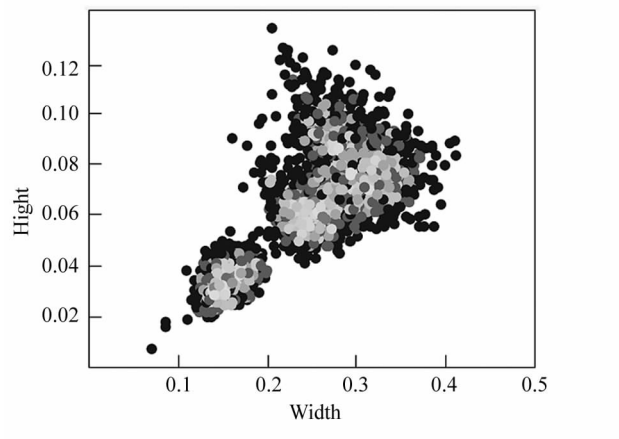


图 2 归一化目标大小图
Fig. 2 Map of normalized target size

中,且中小尺寸的物体占比较大。原生 YOLOv5 的初始化锚框无法很好地契合车牌检测任务的需要,本文采用了 K-means++ 聚类算法针对车牌数据重新设计了初始化锚框的大小,多次聚类后得到 9 组新的大小不同的锚框,改进后 3 种尺度的初始锚框为[108,11,171,18,175,32]、[190,27,209,27,217,21]、[221,29,246,25,248,21]。实验表明,将先验锚框重新聚类后,使模型的检测效果有了一定的提升。

3.2 融入 CBAM 注意力机制

目前,注意力机制已经广泛应用于计算机视觉领域,它借鉴了人的视觉原理,即人在看到一个物体时,首先会扫描整个物体,然后会得到物体的重点部位,也就是我们常说的焦点,通过识别焦点特征来判定该物体属于哪个分类。在卷积神经网络中,注意力机制用于提取特征图上可利用的注意力信息。CBAM(convolutional block attention module)^[16]表示卷积模块中的注意力机制模块,结合了空间和通道的注意力信息,通过 CAM(channel attention module)和 SAM(spatial attention moudule)两个子模块对特征图进行重组,提升空间位置信息和通道信息特征的重要性,抑制无用信息,从而达到网络模型对目标检测效果提升的目的,其整体结构如图 3 所示。

YOLOv5 网络中提取图像特征关键的地方在于

Backbone,实验把 CBAM 融合在 Backbone 之后, Neck 之前,目的是当图像在 Backbone 中完成特征提取之后,及时进行注意力重组,再经过 Neck 特征融合完成在不同尺度上的特征图输出和预测。融入 CBAM 注意力机制的网络可视化效果对比如图 4 所示,经过对比可以发现,检测网络融入 CBAM 后,特征覆盖到了待识别物体的更多部位,从而获得了更多的目标特征,使最终检测到目标物体的概率提高,这表明注意力机制的确让网络学会了关注重点信息。

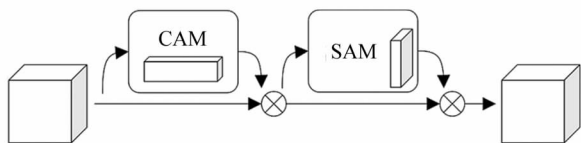


图 3 CBAM 结构示意图

Fig. 3 Diagram of CBAM structure



图 4 注意力可视化对比:(a) 原生;(b) 改进后

Fig. 4 Attentional visualization comparison:

(a) Original; (b) Improved

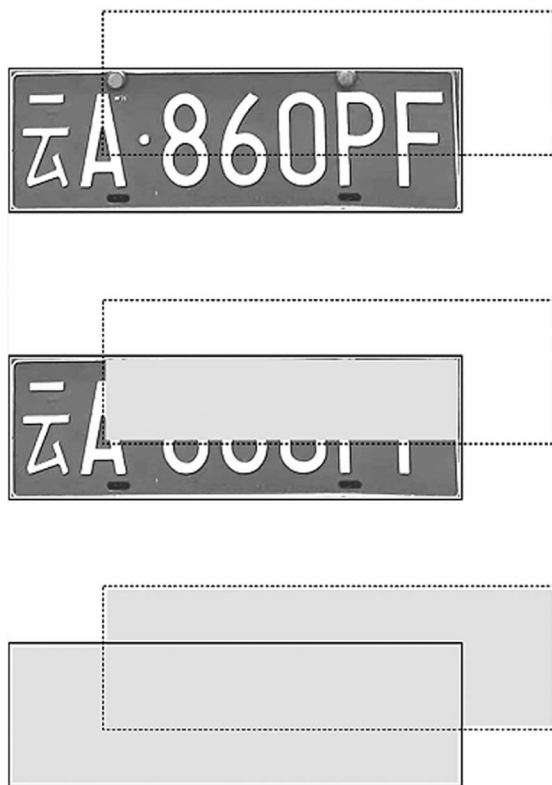
3.3 边界框回归损失函数优化

IoU (intersection over union)是目标检测任务中常见的指标,用来衡量目标框和预测框之间的距离,其数值大小等于目标框和预测框的交集与并集的比值。从图 5 中可以看出, IoU 越大,预测框的预测效果越好。但若以 $Loss_{IoU}=1-IoU$ 作为模型的损失函数,则会有以下缺点:1) 当目标框与预测框不重叠时, $IoU=0$,无法反映两个框的距离远近,此时损失函数不可导,导致无法优化;2) 当两个预测框大小相同,且 IoU 也相同时,无法区分目标框与预测框的相交情况。

在 YOLOv5 的设计中,作者把 $Loss_{GIoU}$ 作为模型的损失函数,其计算方法如下所示:

$$Loss_{GIoU} = 1 - (IoU - \frac{|C - (B \cup B^g)|}{|C|}), \quad (1)$$

式中, C 是指最小外接矩形面积(可以把目标框和预测框包含在内,如图 6 所示), B 和 B^g 分别表示目标框和预测框的面积。

图 5 IoU 示意图Fig. 5 IoU diagram

C : minimum external rectangle

图 6 $GIoU$ 示意图Fig. 6 $GIoU$ diagram

$GIoU$ (generalized intersection over union)的融入解决了当目标框与预测框不重叠时导致的损失函数无法优化问题,而一些问题仍然存在。通过图 7 分析可以发现,若预测框在目标框的内部且预测框的大小相同,此时 3 种状态下的 $GIoU$ 值大小相同且与 IoU 值相同,依然无法区分相交情况。

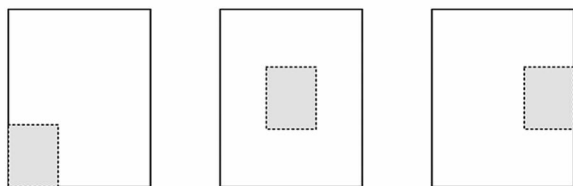


图 7 IoU 对比

Fig. 7 IoU comparison

基于上述对 $GIoU$ 损失局限性的分析,实验采用 $CIoU$ (complete intersection over union)损失作为边界框回归的损失函数。 $CIoU$ 不仅考虑了目标的重叠面积,而且把中心点距离和长宽比加入指标的影响因素,使得预测框的回归效果更好。计算式如下:

$$Loss_CIoU = 1 - (IoU - \frac{\rho^2(b, b^{gt})}{c^2} - \alpha v), (2)$$

$$v = \frac{4}{\pi^2} (\arctan \frac{\omega^{gt}}{h^{gt}} - \arctan \frac{\omega}{h})^2, (3)$$

$$\alpha = \frac{v}{(1 - IoU) + v}, (4)$$

式中, b 与 b^{gt} 分别表示目标框与预测框的中心点, ρ 表示欧氏距离, c 表示包含目标框和预测框的最小矩形框对角线的长度, v 用来衡量检测框长宽比的相似性, α 代表一个权重参数。

4 实验设计及实现

4.1 实验环境

实验中的模型搭建、训练以及结构的测试都是在深度学习框架 Pytorch 中完成,并使用 CUDA 和 cudnn 对 GPU 进行加速,提高计算机的运算能力,实验所需的环境具体见表 1。

表 1 实验运行环境

Tab. 1 Experimental operating environment

Parameters	Value
Operating systems	Ubuntu20.04
CPU	AMD Ryzen 7 4800H
GPU	GeForce GTX 1080 Ti
Python	3.8

4.2 评价指标

在目标检测任务中,常见的评价指标有精度(Precision)和召回率(Recall),精度表示预测目标中正确的目标占有所有预测目标的比例,召回率表示预测目标中正确的目标占有所有正确目标的比例。其计算方法如下:

$$Precision = \frac{TP}{TP + FP}, (5)$$

$$Recall = \frac{TP}{TP + FN}, (6)$$

式中, TP 指预测目标中正确的目标, FP 指预测目标中错误的目标, FN 指未被预测出的正确目标。其中,当 $IoU \geq$ 阈值时,认为目标预测正确, $IoU <$ 阈值时,认为目标预测错误。实际实验中采用综合了精度和召回率的均值平均精度(mean average precision, mAP)作为评价指标, mAP 是目标检测任务中一个最为常见的评价指标^[17],取值范围在 0—1 之间, mAP 越大表示网络模型的检测效果越好,其计算方法如下:

$$mAP = \frac{1}{c} \sum_{k=1}^N P(k) \Delta R(k), (7)$$

式中, N 表示测试集中的样本数, P 表示精度, k 表示第 k 个样本, $\Delta R(k)$ 表示召回率 R 由第 $k-1$ 个样本到第 k 个样本发生的变化, c 则表示多分类检测任务中的所要识别的类别数。

文章采用 $mAP@0.5$ 、 $mAP@0.5 : 0.95$ 作为模型检测效果的衡量指标。 $mAP@0.5$ 定义为当 $IoU = 0.5$ 时的 mAP 值,数值越高表示模型在高召回率下越容易保持高准确率。 $mAP@0.5 : 0.95$ 定义为 IoU 阈值从 0.5 以 0.05 的步长到 0.95 并取所有 mAP 的均值,可衡量模型在不同 IoU 阈值下的综合性能,数值越高表示预测框与真实框拟合越精准。

4.3 数据预处理

4.3.1 数据集及数据增强

在深度学习目标检测领域,数据集的质量对最终任务的检测效果有重要影响。实验使用的车牌数据主要来源于中国城市停车数据集(CCPD)和林区自采集的车牌数据。

为了解决车牌目标较小导致难以检测的问题, YOLOv5 使用了 Mosaic 的数据增强方法。Mosaic 数据增强的主要思想是将 4 张图片通过随机裁剪、缩放和拼接在同一张图片上,从而增大了原有的数据集和提高了训练效率。这有效地解决了模型训练中小目标不如大目标那样准确被检测到的问题,适用于车牌等小目标的检测任务。

此外,为了在增强特殊条件下检测效果的同时,降低样本不均衡对训练的不良影响,对原始数据中部分图片进行前期处理,经处理后的数据集见表 2。通过使用 OpenCV-Python 编写脚本分别对数据进行增强亮度、减弱亮度、添加盐噪声和蒙版的处理,

模拟图像采集时可能出现的强光、暗光、雨雪天、大雾等干扰,增强模型在林区易出现恶劣天气下检测的鲁棒性,数据增强效果如图 8 所示。

表 2 数据集情况

Tab. 2 Datasets information

Type	Normal	Strong light	Less light	Rain& snow	Fog	Total
Original	4 200	1 125	1 325	893	843	8 386
Ours	4 200	2 000	2 000	1 500	1 500	11 200

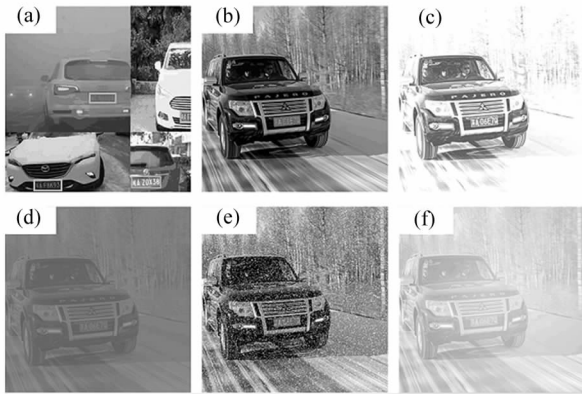


图 8 图像增强样本:(a) Mosaic; (b) 正常; (c) 强光;
(d) 暗光; (e) 雨雪; (f) 雾

Fig. 8 Sample of image enhancement:

(a) Mosaic; (b) Normal; (c) Strong light;
(d) Less light; (f) Rain and snow; (e) Fog

4.3.2 数据集标记

在工作环境中使用 Labelme 工具用以对数据集中的车牌位置进行标注,标注后生成的.json 格式的标签文件不能直接用于 YOLOv5 的训练。实验通过 Python 编写脚本将.json 格式转为.txt 格式,并对标注的信息进行归一化处理,归一化后的标注信息格式为(class_id, C_x , C_y , w , h),具体的运算方法如下:

$$C_x = \frac{1}{2}(x_{\min} + x_{\max}), \quad (8)$$

$$C_y = \frac{1}{2}(y_{\min} + y_{\max}), \quad (9)$$

$$w = x_{\max} - x_{\min}, \quad (10)$$

$$h = y_{\max} - y_{\min}, \quad (11)$$

式中, C_x 、 C_y 表示标注的车牌检测框的中心坐标, $x_{\min}$ 、 $x_{\max}$ 、 $y_{\min}$ 、 $y_{\max}$ 表示原始标注数据检测框的边界信息。class_id 表示多分类任务中类别的编号,由于实验研究的检测任务只有车牌这一类,故将 class_id 默认为 0。

4.4 整体实现流程

首先准备实验所需要的车牌数据集,并对数据集进行车牌标注并获得标签文件,使用数据集中的训练集和验证集进行模型训练,得到的最佳权重模型在测试集上验证识别效果,最后画出车牌在图片上的位置,具体处理流程见图 9。

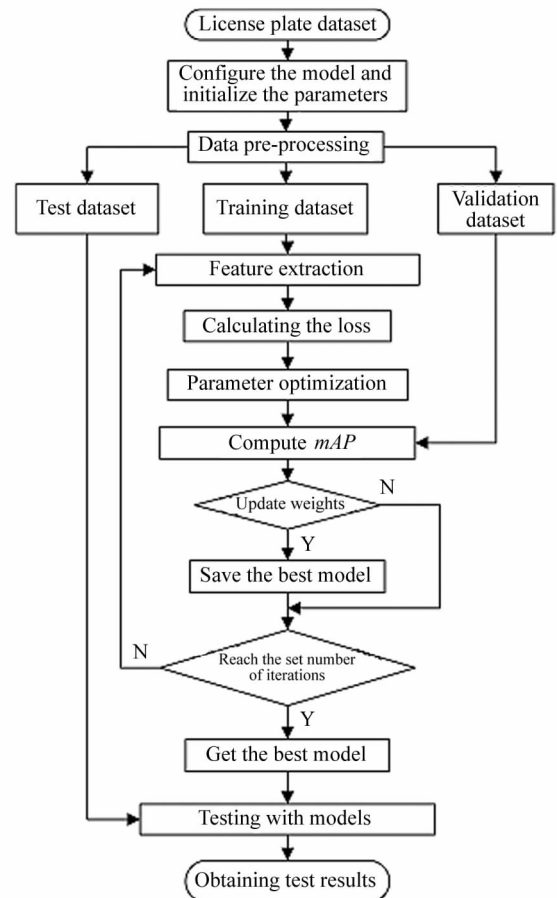


图 9 整体实现流程

Fig. 9 Flow chart of overall implementation

4.5 模型的训练

将实验使用各类型数据集按照 8 : 1 : 1 的比例随机划分成训练集、验证集、测试集。为了使模型的训练达到理想性能,实验使用的部分超参数设置如下:将 Epoch 设置为 200, batch-size 为 16, 学习率为 0.01, momentum 为 0.937。

在网络模型训练中,模型结构的损失函数值越小表示训练后的权重与真实值拟合效果越好。改进前后损失的变化情况如图 10 所示,从图中可以看出,在训练前期损失下降较快,但随着 Epoch 的增加,损失值逐渐趋于稳定,直到在第 150 个 Epoch 时模型基本完成收敛,且在训练过程中没有出现拟合现象。相较于 YOLOv5 算法,实验算法收敛更快、

过程更加平滑,且收敛后的损失值更低,有利于获得检测效果更好的权重模型。

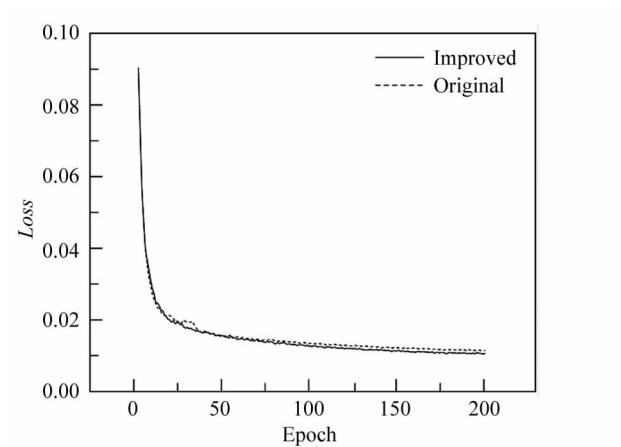


图 10 训练损失曲线对比

Fig. 10 Comparison of training loss curves

4.6 实验结果和分析

4.6.1 实验结果对比

为了更加直观地展示改进前后的车牌检测效果,文章设置了 5 组测试图像,其中包括正常、强光、暗光、雨雪、大雾等场景,分别使用改进前后的检测模型对场景中的车牌进行检测,检测对比结果见图 11(左图为改进前,右图为改进后)。从图 11(a)中可以看出,改进前后的车牌检测算法在正常环境下都取得不错的检测效果,实验算法除了准确地预测出车牌位置外,对小目标车牌也有更好的检测准确率;从图 11(b)、(c)、(d)和(e)中可以看出,改进前的车牌检测算法表现较差,频繁出现误检、漏检现象,而实验算法则较为准确地检测出在不同场景下的车牌位置。



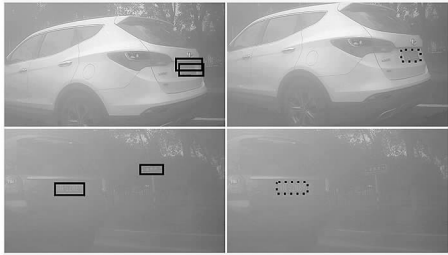
(a) Normal



(b) Strong light



(c) Less light



(d) Rain and snow



(e) Fog

图 11 车牌检测效果展示

Fig. 11 License plate detection effect display

4.6.2 消融实验

将训练后的权重模型在测试集上进行测试,为了分析不同改进部分对模型性能的影响,对整个过程进行消融实验,见表 3。消融实验以 YOLOv5s 为基础模型,评估指标为召回率、 $mAP@0.5:0.95$ 和推理时间。

表 3 消融实验结果

Tab. 3 Results of ablation experiments

Anchor	CBAM	CIoU	Recall	$mAP@0.5:0.95$	Inference time/ms
×	×	×	0.991	0.841	7.3
✓	×	×	0.994	0.846	7.4
✓	✓	×	0.997	0.861	7.7
✓	✓	✓	0.998	0.867	7.8

从表 3 中可以看出,实验算法较原始 YOLOv5 在 mAP 上提升了 2.6%。值得注意的是, CBAM 注意力机制使模型的 $mAP@0.5:0.95$ 值上升了 1.5%,对模型性能的提升最大,证明 CBAM 注意力机制确实可以提取到目标更多的特征信息,但同时

也使推理时间增加了 0.3 ms。在消融实验的逐步进行中可以发现,随着改进方法的叠加,模型的性能都会在上一个改进的基础上有一个小幅度的提升,证明了本文提出的改进 YOLOv5 方法在林区等易发生复杂天气环境下对提升车牌检测效果是有意义的。

4.6.3 主流目标检测模型性能对比

为了验证实验算法对车牌的检测性能,将改进的算法与主流的目标检测算法 SSD (single shot multiBox detector)、Faster R-CNN、YOLOv3 进行性能对比,采用召回率、 $mAP@0.5$ 和推理时间作为性能评估指标,并加入训练后模型大小的比较,用来评估检测模型对硬件的依赖情况。由表 4 分析可知,实验算法发挥了 YOLOv5 检测速度快的优势,使推理时间仅有 7.8 ms,远超过其他主流目标检测算法,并取得了良好的检测效果。此外,检测算法在训练后生成的权重模型对硬件环境依赖最小,降低了使用成本,便于检测算法后续在林区中广泛部署。

表 4 主流目标检测模型性能对比

Tab. 4 Mainstream target detection model performance

Model	Recall	$mAP@0.5$	Inference time/ms	Model size/M
SSD	0.812	0.842	35	207.1
Faster R-CNN	0.998	0.995	165	432.2
YOLOv3	0.911	0.939 5	22	62.1
Ours	0.998	0.995	7.8	14

5 结 论

针对林区易发生雾、雨、雪等复杂天气环境下的车牌检测任务,提出了一种改进的 YOLOv5 车牌检测算法。主要的改进有:优化先验锚框、融合 CBAM 注意力机制、使用 $CIoU$ 作为损失函数。实验结果表明,文章提出的算法可以获得较好的检测效果,满足林区复杂环境中对进入监控区域车辆车牌检测的准确性、实时性和广泛可部署性要求。下一步的工作计划是在提升检测算法性能的同时,将进一步丰富车牌在复杂天气环境下的特征,提高算法的鲁棒性。

参考文献:

[1] DING J, ZHU H Q. Study of license plate recognition in forest areas based on CNN multi-label classification[J]. Modern Information Technology, 2020, 4(12): 84-87.
丁键, 朱洪前. 基于 CNN 多标签分类的林区车牌识别研究[J]. 现代信息科技, 2020, 4(12): 84-87.

[2] ZHENG D, ZHAO Y, WANG J. An efficient method of li-

cense plate location[J]. Pattern Recognition Letters, 2005, 26(15): 2431-2438.

[3] YANG D D, CHEN S Q, LIU J Y. License plate location algorithm based on license plate background and character color feature[J]. Computer Applications and Software, 2018, 35(12): 216-221.
杨鼎鼎, 陈世强, 刘静漪. 基于车牌背景和字符颜色特征的车牌定位算法[J]. 计算机应用与软件, 2018, 35(12): 216-221.

[4] NATHAN V S L, RAMKUMAR J, PRIYA S K. New approaches for license plate recognition system[C]//International Conference on Intelligent Sensing and Information Processing, 2004. Proceedings of IEEE, January 04-07, 2004, Chennai, India. New York: IEEE, 2004: 149-152.

[5] LI Q. A geometric framework for rectangular shape detection[J]. IEEE Transactions on Image Processing, 2014, 23(9): 4139-4149.

[6] ACHARYA U R, OH S L, HAGIWARA Y, et al. A deep convolutional neural network model to classify heartbeats[J]. Computers in Biology and Medicine, 2017, 89: 389-396.

[7] WU X, SONG X R, GAO S, et al. Review of target detection algorithms based on deep learning[J]. Transducer and Microsystem Technologies, 2021, 40(2): 4-7+18.
吴雪, 宋晓茹, 高嵩, 等. 基于深度学习的目标检测算法综述[J]. 传感器与微系统, 2021, 40(2): 4-7+18.

[8] JAMTSO Y, RIYAMONGKOL P, WARANUSAST R. Real-time Bhutanese license plate localization using YOLO[J]. ICT Express, 2020, 6(2): 121-124.

[9] MA Q M, WANG M J, LIANG H R. A license plate location detection algorithm based on improved YOLOv3 in complex scenes[J]. Computer Engineering and Applications, 2021, 57(7): 198-208.
马巧梅, 王明俊, 梁昊然. 复杂场景下基于改进 YOLOv3 的车牌定位检测算法[J]. 计算机工程与应用, 2021, 57(7): 198-208.

[10] SELMI Z, HALIMA M B, ALIMI A M. Deep learning system for automatic license plate detection and recognition [C]//2017 14th IAPR International Conference on Document Analysis and Recognition (ICDAR). IEEE, November 9-15, 2017, Kyoto, Japan. New York: IEEE, 2017, 1: 1132-1138.

[11] REDMON J, DIVVALA S, GIRSHICK R, et al. You only

- look once: Unified, real-time object detection[C]//IEEE Conference on Computer Vision and Pattern Recognition, June 27-30, 2016, Las Vegas, NV, USA. New York: IEEE, 2016: 779-788.
- [12] REDMON J, FARHADI A. YOLO9000: better, faster, stronger[C]//IEEE Conference on Computer Vision and Pattern Recognition, July 21-26, 2017, Honolulu, HI, USA. New York: IEEE, 2017: 7263-7271.
- [13] REDMON J, FARHADI A. YOLOv3: an incremental improvement[EB/OL]. (2018-04-08) [2022-03-11]. <https://arxiv.org/abs/1804.02767>.
- [14] BOCHKOVSKIY A, WANG C Y, LIAO H Y M. YOLOv4: Optimal speed and accuracy of object detection[EB/OL]. (2020-04-23) [2022-03-11]. <https://arxiv.org/abs/2004.10934>.
- [15] YU J, LUO S. Detection method of illegal building based on YOLOv5[J]. Computer Engineering and Applications, 2021, 57(20): 236-244.
- 于娟, 罗舜. 基于 YOLOv5 的违章建筑检测方法[J]. 计算机工程与应用, 2021, 57(20): 236-244.
- [16] WOO S, PARK J, LEE J Y, et al. CBAM: convolutional block attention module[C]//European Conference on Computer Vision (ECCV), September 8-14, 2018, Munich, Germany. Berlin: Springer, 2018: 3-19.
- [17] ZAIDI S S A, ANSARI M S, ASLAM A, et al. A survey of modern deep learning based object detection models [EB/OL]. (2021-04-24) [2022-03-11]. <https://arxiv.org/abs/2104.11892>.

作者简介:

杨洁 (1973—), 女, 工学博士, 副教授, 硕士生导师, 主要从事检测技术与自动化装置、视觉图像处理等方面的研究。