

控制与决策

Control and Decision



多目标小尺度车辆目标检测方法

柳长源, 王琪, 毕晓君

引用本文:

柳长源, 王琪, 毕晓君. 多目标小尺度车辆目标检测方法[J]. 控制与决策, 2021, 36(11): 2707–2712.

在线阅读 View online: <https://doi.org/10.13195/j.kzyjc.2020.0635>

您可能感兴趣的其他文章

Articles you may be interested in

基于MobileNet的多目标跟踪深度学习算法

Deep learning algorithm based on MobileNet for multi-target tracking

控制与决策. 2021, 36(8): 1991–1996 <https://doi.org/10.13195/j.kzyjc.2019.1424>

基于生成对抗网络学习被遮挡特征的目标检测方法

Object detection via learning occluded features based on generative adversarial networks

控制与决策. 2021, 36(5): 1199–1205 <https://doi.org/10.13195/j.kzyjc.2019.1319>

基于改进DenseNet网络的人体姿态估计

Improved DenseNet network for human pose estimation

控制与决策. 2021, 36(5): 1206–1212 <https://doi.org/10.13195/j.kzyjc.2019.1218>

改进YOLOv2的端到端自然场景中文字符检测

End-to-end Chinese character detection in natural scene based on improved YOLOv2

控制与决策. 2021, 36(10): 2483–2489 <https://doi.org/10.13195/j.kzyjc.2020.0270>

复杂背景下全景视频运动小目标检测算法

Panoramic video motion small target detection algorithm in complex background

控制与决策. 2021, 36(1): 249–256 <https://doi.org/10.13195/j.kzyjc.2019.0686>

多目标小尺度车辆目标检测方法

柳长源^{1†}, 王 琪¹, 毕晓君²

(1. 哈尔滨理工大学 电气与电子工程学院, 哈尔滨 150080;

2. 哈尔滨工程大学 信息与通信工程学院, 哈尔滨 150001)

摘要: 车辆目标检测是智能交通系统中的重要环节, 针对传统车辆目标检测方法效率低、小目标检测效果不好、漏检率高等问题, 提出一种基于改进的YOLOv3网络车辆目标检测算法. 为了提高车辆检测的效率, 利用轻量化模型MobileNet v2代替原YOLOv3中的特征提取网络, 使得网络计算量相比原算法有所降低. 为了有效提高网络对小尺度车辆目标的检测能力, 网络将由高到低不同尺度的特征层融合之后进行目标检测. 为了得到更丰富的语义特征信息和提高网络预测能力, 增加了特征增强模块. 同时针对车辆目标检测的特定应用, 利用K-means方法对锚框重新聚类以满足车辆目标检测的特定需求. 结合以上改进获得车辆目标检测网络YOLOv3-M2, 实验结果表明, 与YOLOv3相比, 改进方法平均检测准确率增加约9%, 时间减少约一半, 能够同时提高检测效率和小目标检测能力.

关键词: 智能交通; 深度学习; 目标检测; YOLOv3; 多尺度检测; 轻量化网络

中图分类号: TP391.4

文献标志码: A

DOI: 10.13195/j.kzyjc.2020.0635

开放科学(资源服务)标识码(OSID):



引用格式: 柳长源, 王琪, 毕晓君. 多目标小尺度车辆目标检测方法[J]. 控制与决策, 2021, 36(11): 2707-2712.

Multi-target and small-scale vehicle target detection method

LIU Chang-yuan^{1†}, WANG Qi¹, BI Xiao-jun²

(1. College of Electrical and Electronic Engineering, Harbin University of Science and Technology, Harbin 150080, China; 2. College of Information and Communication Engineering, Harbin Engineering University, Harbin 150001, China)

Abstract: Vehicle target detection is an important link in intelligent transportation systems. Aiming at the problems of low efficiency, poor detection effect of small targets and high miss rate of traditional vehicle target detection methods, a vehicle target detection algorithm based on the improved YOLOv3 network is proposed. In order to improve the efficiency of vehicle detection, the lightweight model MobileNet v2 is used to replace the feature extraction network in the original YOLOv3, and the calculation amount is reduced compared with the original algorithm. In order to effectively improve the network's ability to detect small-scale vehicle targets, feature layers of different scales are fused and target detection is carried out on feature maps of different scales. At the same time, in order to obtain more abundant semantic feature information and improve the prediction ability of network, a feature enhancement module is proposed. For the specific application of vehicle target detection, K-means is used to re-cluster anchor frames to meet the requirements of vehicle target detection. Combined with the above improvements, the vehicle target detection network YOLOv3-M2 is obtained. Experimental results show that compared with YOLOv3, the proposed method not only improves the detection efficiency, but also improves the small target detection capability, increasing the average detection accuracy of the network by about 9%.

Keywords: intelligent transportation; deep learning; object detection; YOLOv3; multi-scale detection; lightweight network

0 引言

随着道路上的车辆迅速增多, 交通监管逐渐成为一个具有挑战性的问题. 在智能交通中, 车辆目标检测是关键的一步, 对后续的车辆追踪和车型识别等均

有重要意义, 一直是国内外学者的研究热点. 小型车辆在道路中较为常见, 在固定监控摄像头正面拍摄的情况下, 有些汽车离摄像头距离较远使得目标尺寸较小, 进而使得车辆目标在图像上占很小的像素, 对应

收稿日期: 2020-05-26; 修回日期: 2020-08-03.

基金项目: 国家自然科学基金项目(51779050); 黑龙江省自然科学基金项目(F2016022).

[†]通讯作者. E-mail: liuchangyuan@hrbust.edu.cn.

区域所含信息量较少,检测时容易发生漏检,影响算法的检测精度.因此,对小尺度车辆目标进行识别与定位是目标检测领域中的难点^[1].

近年来,随着对深度学习的深入研究,学者们已经将其应用到不同领域^[2],例如在计算机视觉领域的目标检测^[3]、图像语义分割、人脸识别^[4]等.基于深度学习的目标检测算法通常分为两类:一类为通过分类区域建议检测目标的算法,如RCNN^[5]、Fast-RCNN^[6]、Faster-RCNN^[7],这类算法将目标检测分为两步,首先利用滑动窗法在图片上获得候选区域,然后提取候选区域的特征向量,利用分类器的评分结果判别候选区域的目标类别;另一类为YOLO^[8]、SSD^[9]等算法,采用直接回归目标类别的方式很大程度地提升了检测速度,但是检测准确率有所下降.这类算法主要将输入图像先划分为网格后进行检测,并输出物体的类别和存在的概率. MobileNet网络用于提取图像的特征信息,是一种轻量化卷积神经网络,其独特的卷积结构可优化网络模型的大小并提升运算速度,在不降低算法精度的情况下提升效率.

利用深度神经网络提取特征比传统的特征提取方法效果更好,本文也使用深度学习的方法检测大小不同的车辆.为了提升YOLOv3算法小尺度车辆的检测能力,减少漏检率并提高检测效率,在YOLOv3网络的基础上进行改进,将YOLOv3与轻量化模型MobileNet v2^[10]相结合,将检测尺度扩展到4种尺度.同时增加了特征增强模块,用来增强4种尺度的车辆特征信息,为网络的预测层提供了丰富的语义信息,从而有效地提升了算法的小目标检测能力.另外,算法针对车辆目标检测这一特定应用,对anchor box进行重新聚类.在对比实验中发现,改进后的目标检测网络YOLOv3-M2相比YOLOv3在小尺度车辆目标检测上取得了不错的效果,明显提高了对车辆目标检测的准确率,减少了漏检情况的发生;由于减少了大量网络参数,检测效率也得到进一步提升.

1 相关网络算法与模型

1.1 YOLOv3算法

YOLO算法直接回归出目标的位置,不再选择候选区域,直接通过回归生成每类目标的边界框坐标和置信度.因此,YOLO算法的计算速度远远超过Faster R-CNN算法.相比YOLO和YOLOv2,YOLOv3算法首先使用了特征图金字塔网络(feature pyramid networks, FPN),实现了3种尺度预测尺度,分别是 13×13 、 26×26 、 52×52 ,其检测精度相比于YOLO和YOLOv2均有所提升.

YOLOv3主要分为特征提取和目标预测两部分,特征提取由Darknet53网络完成. YOLOv3网络首先将输入图像缩放到 416×416 ,送入卷积神经网络进行特征提取,在目标预测过程中利用非极大值抑制方法消除重叠的预测框,保留一个框作为最终的检测框. YOLOv3网络将输入图像划分为 $S \times S$ 个网格,每个网格单独预测目标并可以预测出3个尺寸不同的边界框以及边界框的4个偏移坐标和1个置信度值,所以每个网格得到的张量为

$$S \times S \times [3 \times (4 + 1 + N)]. \quad (1)$$

其中:4代表预测的边界框坐标 (t_x, t_y, t_w, t_h) ,1代表目标置信度, N 为数据集中目标类数.

1.2 MobileNet v2

轻量化网络的卷积与标准卷积的区别在于卷积的计算方式不同,前者的优点是能够在保持精度的同时减少网络参数和计算量. MobileNet v2是基于MobileNet v1^[11]改进的网络, MobileNet主要运用深度可分离卷积,其是许多高效神经网络架构的关键结构块,通过其来减少运算量以及参数量.深度可分离卷积由两部分组成:第1部分是深度卷积层(depthwise),通过对特征图的各个通道应用单个卷积滤波器进行卷积操作;第2部分是 1×1 的点卷积,将多个卷积层线性结合. MobileNet v2网络借鉴残差网络中的残差模块提出了倒立残差结构模块.残差模块对特征图先“压缩”再“扩张”,倒立残差结构则相反,先采用卷积作扩张后连着深度卷积层,最后卷积作压缩,因此称为倒残差结构.如图1所示, MobileNet v2有stride = 1和stride = 2两种结构不同的模块.

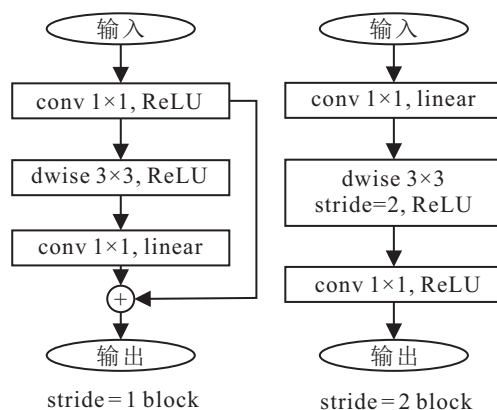


图1 MobileNet v2的基本结构块

由图1可见,在stride = 1模块中使用shortcut连接两个卷积层,提高了梯度跨层的传播能力,为了使维度匹配, stride = 2块中不采用shortcut,以上两种结构构成MobileNet v2的基本结构块.表1为

MobileNet v2 基本结构块输入与输出的关系. 其中: 输入图像尺寸为 $h \times w$, 通道数为 d , 输出通道数为 d' , 深度卷积层卷积核大小为 k , 卷积步长为 s , 扩展影响因子为 $t(0 < t < 1)$.

表 1 扩展影响因子为 t 的模块结构

输入	操作	输出
$h \times w \times d$	1×1	$h \times w \times td$
$h \times w \times td$	$k \times k, \text{dwise} = s$	$\frac{h}{s} \times \frac{w}{s} \times td$
$\frac{h}{s} \times \frac{w}{s} \times td$	linear 1×1	$\frac{h}{s} \times \frac{w}{s} \times d'$

基本结构块的计算量为三层卷积相加的结果, 即

$$S_1 = (1 \times 1 \times h \times w \times d \times td) + (k \times k \times h \times w \times td \times 1) + (1 \times 1 \times h \times w \times td \times d') = h \cdot w \cdot d \cdot t(d + k^2 + d'),$$

标准卷积计算量为

$$S_2 = h \cdot w \cdot w \cdot d \cdot d' \cdot k^2,$$

其比值为

$$\frac{S_1}{S_2} = \frac{td}{d'k^2} + \frac{1}{d'} + \frac{1}{k^2}.$$

当 $k = 3$ 且 $d = d'$ 时计算量可减少 8~9 倍.

2 网络设计

2.1 YOLOv3-M2 网络

通过使用轻量化网络模型减少网络参数和网络运算量可以有效提升车辆目标检测的效率. MobileNet v2 为轻量化网络, 主要用于图像的特征信息提取任务, 其与标准卷积的计算方式不同, 优点是能够在保持精度的同时减少网络参数和计算量, 利用 MobileNet v2 提取图像的特征信息, 利用 YOLOv3 的多尺度检测部分进行检测, 在保持精度的同时提升检测效率^[12-13]. 图 2 为目标检测的网络结构. YOLOv3-M2 网络分为两部分, 第 1 部分通过 MobileNet v2 提取目标特征, 第 2 部分通过 YOLOv3 检测出车辆目标. 首先将图片分辨率调整至 416×416 大小后输入进 MobileNet v2 网络提取特征, MobileNet v2 含有 17 个基本结构块, 图片经过 MobileNet v2 网络后得到一个 $13 \times 13 \times 1024$ 维的张量. 因为该网络预测的目标只有汽车这一类, 由于式(1)中 $N = 1$, 通过一个 1×1 的卷积核进行卷积操作得到一个 $S \times S \times 18$ 维的张量, 通过此张量预测车辆目标的位置.

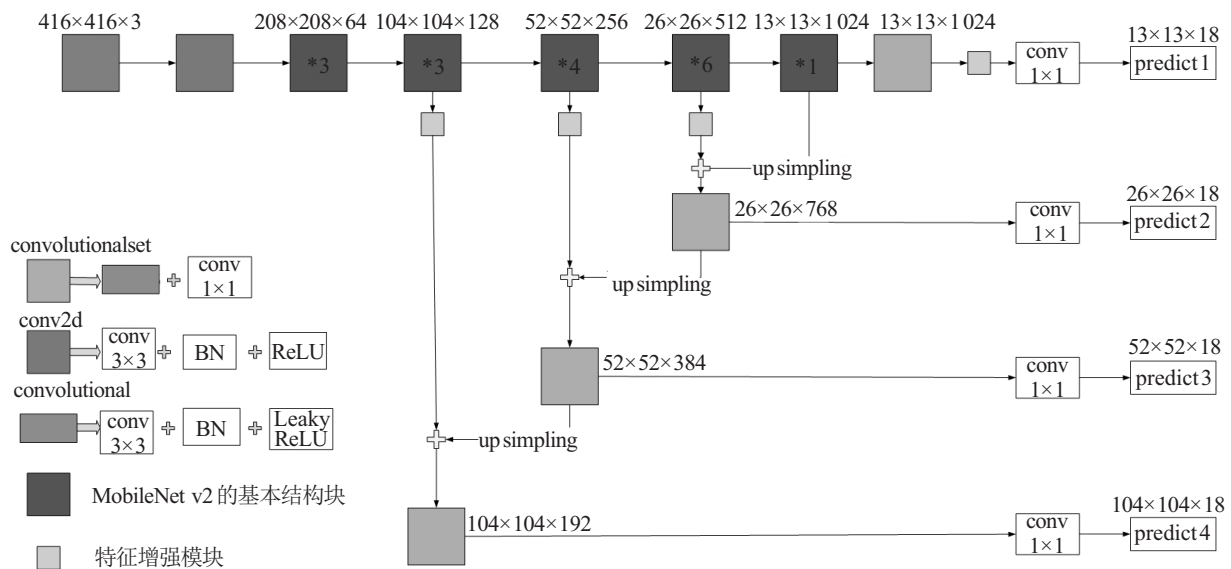


图 2 YOLOv3-M2 网络结构

在车辆目标检测时, 由于车辆和摄像机的距离有远有近, 车辆在图片上呈现的大小也不等, 最后一层特征层的尺寸仅为 13×13 , 是原输入图像的 $1/32$, 这使得特征层会丢失一些较小物体的特征信息. 在深度神经网络中, 级数越高特征图尺寸越小, 所包含的语义信息更为丰富; 而低层的特征层具有更大的分辨率, 并保留了更多原始图像中的细节信息, 有利于确定物体的位置. 本文采用高层特征与低层特征相

融合并在多个尺度特征图上预测的方法提高网络对小目标车辆目标检测的能力. 第 1 次预测采用尺寸为 13×13 的特征图, 将其上采样变为 26×26 的特征图后与卷积过程中尺寸为 26×26 的特征图结合作为第 2 次预测的基础. 利用这样的方式分别得到尺寸为 52×52 和 104×104 的特征图做第 3 次和第 4 次预测.

由于数据集只有车辆一类目标, 损失函数由回归框损失和置信度损失组成. 损失函数为

loss =

$$\begin{aligned} & \lambda_{\text{coord}} \sum_{i=0}^{S^2} \sum_{j=0}^B I_{ij}^{\text{obj}} [(x_i - \hat{x}_i^j)^2 + (y_i - \hat{y}_i^j)^2] + \\ & \lambda_{\text{coord}} \sum_{i=0}^{S^2} \sum_{j=0}^B I_{ij}^{\text{obj}} [(\sqrt{w_i^j} - \sqrt{\hat{w}_i^j})^2 + (\sqrt{h_i^j} - \sqrt{\hat{h}_i^j})^2] - \\ & \sum_{i=0}^{s^2} \sum_{j=0}^B I_{ij}^{\text{obj}} [\hat{C}_i^j \log(C_i^j) + (1 - \hat{C}_i^j) \log(1 - C_i^j)] - \\ & \lambda_{\text{noobj}} \sum_{i=0}^{S^2} \sum_{j=0}^B I_{ij}^{\text{obj}} [\hat{C}_i^j \log(C_i^j) + (1 - \hat{C}_i^j) \log(1 - C_i^j)]. \end{aligned} \quad (2)$$

其中: x_i, y_i, w_i, h_i 为模型的预测值; $\hat{x}_i, \hat{y}_i, \hat{w}_i, \hat{h}_i$ 为人工事先标记的真实值; λ_{coord} 为加权系数; I_{ij}^{obj} 表示第 i 个网格第 j 个 anchor box 是否负责该 object, 若负责, 则 $I_{ij}^{\text{obj}} = 1$, 否则为 0; C_i^j 为置信度预测值, $\hat{C}_i^j$ 为实际值, 取值由网格 anchor box 是否负责预测某个对象决定, 若负责, 则 $\hat{C}_i^j = 1$, 否则为 0; λ_{noobj} 为权重系数, 一般取值为 0.5.

2.2 特征增强模块

YOLOv3-M2 网络进行多尺度的车辆检测, 在特征提取网络中高层的小尺度特征图进行上采样操作后与低层的大尺度特征图进行融合. 为了丰富 4 种尺度特征图的特征信息提出特征增强模块, 特征提取网络输出的各个尺度的特征图在特征增强模块中与感受野大小不同的卷积核进行卷积后再进行融合, 减少卷积过程中的信息损失, 并为目标预测部分提供丰富的特征信息, 提升 YOLOv3-M2 网络的目标检测能力. 特征增强模块结构如图 3 所示.

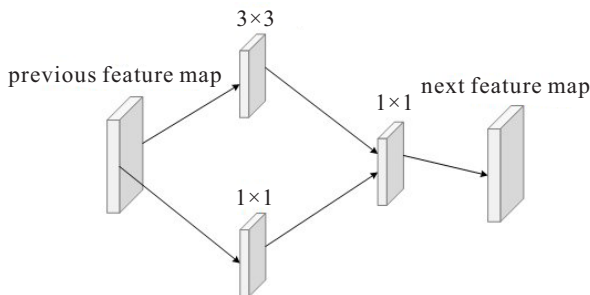


图 3 特征增强模块结构

特征增强模块由 1×1 和 3×3 大小的卷积核构成, 不同大小的卷积核可以获得图像中不同感知域的信息, MobileNet v2 提取的多尺度特征图通过特征增强模块后再进行特征融合, 可以获得丰富的语义信息, 提高了多尺度特征的提取能力. 特征增强模块能够发掘各个尺度特征图的感知域和有意义的语义信息, 有利于小尺度车辆目标的检测.

2.3 Anchor 框重新聚类

原始 YOLOv3 算法中的 anchor box 尺寸经过 COCO 数据集和 PASCAL VOC 数据集训练时聚类得到, 在 PASCAL VOC 数据集中有 20 类目标, 在 COCO 数据集中有 80 类目标, 这些目标物体尺寸不一, 因此聚类出的 anchor box 形状不一. YOLOv3-M2 目标检测网络主要的检测目标只是车辆, 针对车辆目标数据集, 多数 anchor box 形状应该是矮胖的, 即 anchor box 宽度大于高度. 为了使得 YOLOv3-M2 网络更准确地预测目标位置, 利用 K -means 算法对车辆目标数据集重新聚类, 得到更精准、更具代表性的 anchor box. K -means 算法随机选取 k 个初始的聚类中心, 计算其他目标与聚类中心的距离, 并分配给最近的聚类中心成为 k 个群, 通过迭代调整使群中各个目标之间的距离变小, 群间距离变大. K -means 算法通常以欧氏距离作为计算的度量距离, 但在目标检测算法中更适合采用预测框和 anchor box 的面积重叠度 $\text{IOU}(B, C)$ 作为度量距离, 新的度量标准计算公式为

$$d(B, C) = 1 - \text{IOU}(B, C). \quad (3)$$

其中: B 为物体真实包围框集合, C 为聚类中心框集合. 对车辆数据集进行聚类, 改变聚类中心 k 的个数得到不同的平均 IOU 结果如图 4 所示. 由图 4 可见, 当 $k = 9$ 时曲线逐渐平缓, 因此选择 9 个 anchor box 分别为 (18.4 12.5)、(25.4 18.6)、(40.8 26.0)、(30.0 28.6)、(63.1 38.1)、(38.5 42.0)、(50.8 58.0)、(83.9 67.1)、(118.6 106.6).

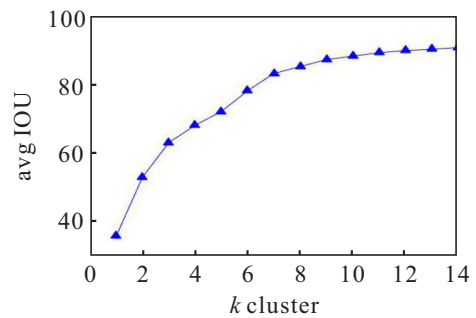


图 4 IOU 结果比较

3 实验分析

3.1 数据集与实验环境

本文采用 UA-DETRAC 数据集进行实验, 数据集取景于北京和天津的道路过街天桥, 包含图片 80 245 张, 随机抽取 70% 作为训练集, 剩余 30% 样本作为测试集. 实验操作系统为 Ubuntu 14.04, 程序设计平台为 Python, 使用 GPU 加速, 同时安装 CUDA 10.0 和 cudnn 7.4.2 以支持 GPU 的使用.

3.2 实验方案与结果分析

在对 YOLOv3-M2 网络模型进行训练时,采用批量机梯度下降法优化损失函数,将初始学习率设为 0.001,最大迭代次数为 50 000 次,权重衰减值设为 0.000 5,批量大小设为 64,网络迭代 40 000 次和 45 000 次后分别将学习率改为 0.000 1 和 0.000 01。

YOLOv3-M2 网络模型训练完成后利用测试集对其进行测试,图 5 所示为 YOLOv3 算法与 YOLOv3-M2 网络的车辆检测结果对比。由图 5 可见, YOLOv3 出现了严重的漏检现象, YOLOv3-M2 对于图片上的小尺度车辆目标检测的效果非常优秀。

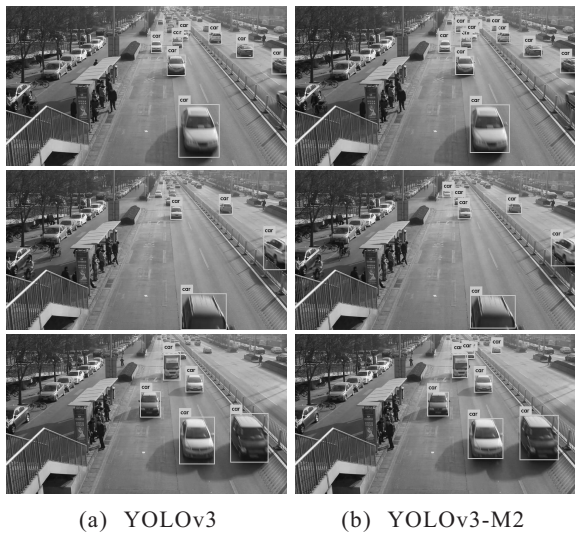


图 5 目标检测结果对比

实验采用精度 (precision)、召回率 (recall)、平均精度 (AP) 作为算法性能定量评价的标准,分别表示为

$$\text{precision} = \frac{N_{\text{true}}}{M}, \quad (4)$$

$$\text{recall} = \frac{N_{\text{true}}}{K}, \quad (5)$$

$$\text{AP} = \frac{\sum \text{precision}}{N}. \quad (6)$$

其中: N_{true} 为成功检测的个数, M 为检测的车辆总数, K 为测试集中的车辆目标总数, N 为图片总数。算法的检测精度、召回率与运算速度对比结果如表 2 所示。

表 2 检测结果对比

	precision	recall	AP / %	t / s
YOLOv3	87.3	83.5	85.3	0.332
YOLOv3-M2*	91.7	88.3	90.9	0.195
YOLOv3-M2**	93.8	91.2	93.4	0.188
YOLOv3-M2	95.5	92.6	94.8	0.193

表 2 中: YOLOv3-M2* 未对车辆数据集进行 anchor box 重新聚类,使用原始 YOLOv3 中的 anchor

box 尺寸; YOLOv3-M2** 未添加特征增强模块。由表 2 可见,由于 YOLOv3-M2 网络使用了多尺度特征检测的方法、特征增强模块和 anchor box 的重新聚类,整体的检测精度和召回率均高于 YOLOv3、YOLOv3-M2* 和 YOLOv3-M2**, YOLOv3-M2* 和 YOLOv3-M2** 的检测精度和召回率相比于 YOLOv3 算法同样有所提高。由于 YOLOv3-M2 网络使用了轻量化模型 MobileNet v2 提取图像特征,减少了网络运算量和参数数量,对于单张图片的检测时间, YOLOv3-M2 网络相比于 YOLOv3 的检测时间有所减少。为了使实验对比结果更直观地表现出来,图 6 绘制了 4 个网络的 P-R 曲线, P-R 曲线围起来的面积即为 AP 值,可以更直观地看出 YOLOv3-M2 结果更优。

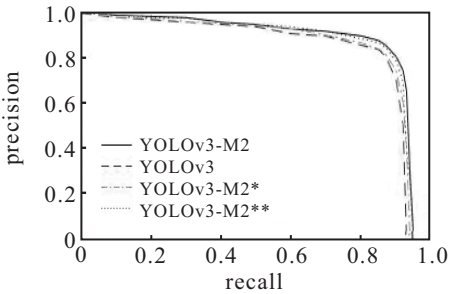


图 6 P-R 曲线

4 结 论

本文针对传统目标检测算法小尺度目标检测能力差、漏检率高、检测效率低的问题,提出了改进的 YOLOv3-M2 网络。利用 MobileNet v2 网络替换原 YOLOv3 中的 DarkNet53,减小了网络模型的规模和运算量,从而提高了 YOLOv3-M2 网络的检测效率。相比于原 YOLOv3 算法提取 3 种尺度特征图, YOLOv3-M2 网络提取了 4 种不同尺度的特征图,为算法的目标预测部分提供更多的细节信息,同时对这 4 种尺度的特征图均添加了特征增强模块,可以提供不同的感知域和有意义的语义信息,丰富了 4 种尺度特征图的特征信息,从而更有利于对小尺度目标进行检测。针对车辆目标检测的专用性,对 YOLOv3 算法中的 anchor box 进行重新聚类得到新的 anchor box 的尺寸,更符合车辆目标检测的应用,使得卷积神经网络更容易准确预测目标位置。通过实验对比,在检测效率和检测精度上, YOLOv3-M2 算法相比于 YOLOv3 均有明显提升。

参考文献 (References)

[1] 徐子豪, 黄伟泉, 王胤. 基于深度学习的监控视频中多类别车辆检测[J]. 计算机应用, 2019, 39(3): 700-705.

- (Xu Z H, Huang W Q, Wang Y. Multi-class vehicle detection in surveillance video based on deep learning[J]. Journal of Computer Applications, 2019, 39(3): 700-705.)
- [2] 张顺, 龚怡宏, 王进军. 深度卷积神经网络的发展及其在计算机视觉领域的应用[J]. 计算机学报, 2019, 42(3): 453-482.
(Zhang S, Gong Y H, Wang J J. Development of deep convolution neural network and its application in the field of computer vision[J]. Chinese Journal of Computers, 2019, 42(3): 453-482.)
- [3] 李会军, 王瀚洋, 李杨, 等. 一种基于视觉特征区域建议的目标检测方法[J]. 控制与决策, 2020, 35(6): 1323-1328.
(Li H J, Wang H Y, Li Y, et al. An object detector based on visual feature region proposal[J]. Control and Decision, 2020, 35(6): 1323-1328.)
- [4] Wang H, Gong D H, Li Z F, et al. Decorrelated adversarial learning for age-invariant face recognition[C]. IEEE Conference on Computer Vision & Pattern Recognition. Long Beach: IEEE, 2019: 3527-3536.
- [5] Girshick R, Donahue J, Darrell T, et al. Rich feature hierarchies for accurate object detection and semantic segmentation[C]. Proceedings of IEEE Conference on Computer Vision & Pattern Recognition. Columbus: IEEE, 2014: 580-587.
- [6] Girshick R. Fast R-CNN[C]. Proceedings of IEEE International Conference on Computer Vision. Santiago: IEEE, 2015: 1440-1448.
- [7] Ren S Q, He K M, Girshick R, et al. Faster R-CNN: Towards real-time object detection with region proposal networks[J]. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2017, 39(6): 1137-1149.
- [8] Redmon J, Divvala S, Girshick R, et al. You only look once: Unified, real-time object detection with region proposal networks[J]. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2017, 39(6): 1137-1149.
- [9] Liu W, Anguelov D, Erhan D, et al. SSD: Single shot multibox detector[C]. European Conference on Computer Vision. Amsterdam: Springer, 2016: 21-37.
- [10] Sandle M, Howard A, Zhu M L, et al. MobileNet v2: Inverted residuals and linear bottlenecks[C]. The 31st Meeting of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. Salt Lake City: IEEE, 2018: 4510-4520.
- [11] Howard A G, Zhu M, Chen B, et al. MobileNets: Efficient convolutional neural networks for mobile vision applications[C]. IEEE Conference on Computer Vision and Pattern Recognition. Honolulu: IEEE, 2017: 1-9.
- [12] 袁哲明, 袁鸿杰, 言雨璇, 等. 轻量化深度学习模型的田间昆虫自动识别与分类算法[J]. 吉林大学学报, DOI: 10.13229/j.cnki.jdxbgxb20200116.
(Yuan Z M, Yuan H J, Yan Y X, et al. Automatic recognition and classification algorithm of field insects based on lightweight deep learning mode[J]. Journal of Jilin University, DOI: 10.13229/j.cnki.jdxbgxb2020 0116.)
- [13] 张富凯, 杨峰, 李策. 基于改进YOLOv3的快速车辆检测方法[J]. 计算机工程与应用, 2019, 55(2): 12-20.
(Zhang F K, Yang F, Li C. Fast vehicle detection method based on improved YOLOv3[J]. Computer Engineering and Applications, 2019, 55(2): 12-20.)

作者简介

柳长源(1970—), 男, 副教授, 博士, 从事模式识别与图像处理等研究, E-mail: liuchangyuan@hrbust.edu.cn;

王琪(1996—), 女, 硕士生, 从事模式识别与计算机视觉的研究, E-mail: 1208401521@qq.com;

毕晓君(1964—), 女, 教授, 博士生导师, 从事机器学习与智能信息处理技术等研究, E-mail: bixiaojun@hrbeu.edu.cn.

(责任编辑: 郑晓蕾)