

面向电能表检定流水线的轻量化目标检测算法

董贤光¹, 孙艳玲¹, 代燕杰¹, 邢宇¹, 翟晓卉¹, 孙凯¹, 吕玉超², 吴强², 刘璐²

(1. 国网山东省电力公司营销服务中心(计量中心), 济南 250001; 2. 山东大学信息科学与工程学院, 青岛 266237)

摘要: 在工业流水线场景中, 利用带有视觉信息的目标检测技术为故障发现及消缺提供决策信息已成为智能生产的新热点。针对电能表流水线检定场景中目标遮挡、小目标密集排列等问题, 在YOLOv8n的基础上, 提出了一种轻量化目标检测算法。通过引入O-GELAN模块, 在保持低计算量的同时获取更丰富的特征层次。利用特征收集-融合-分发的颈部架构和通道位置注意力机制实现特征跨层融合; 此外, 采用重参数化卷积检测头以进一步提高检测效率。在现场采集流水线数据上的实验表明, 改进后模型的mAP(0.5)和mAP(0.5:0.95)分别达到了0.994和0.828, 检测速度可达111.5帧/s, 能够满足工业场景下的高精度和高实时性需要, 可为故障消缺提供辅助决策。

关键词: 工业流水线; 密集目标检测; YOLOv8; 特征融合; 现场数据采集

中图分类号: TP391.4

文献标志码: A

Lightweight Object Detection Algorithm for Electric Meter Calibration Line

DONG Xianguang¹, SUN Yanling¹, DAI Yanjie¹, XING Yu¹, ZHAI Xiaohui¹, SUN Kai¹,
LYU Yuchao², WU Qiang², LIU Ju²

(1. Marketing Service Center (Metering Center), State Grid Shandong Electric Power Co. Ltd., Jinan 250001, China; 2. School of Information Science and Engineering, Shandong University, Qingdao 266237, China)

Abstract: In the industrial production line scenario, target detection technology with visual information has become a new hotspot for intelligent production, providing decision-making information for fault detection and elimination. In response to issues such as target occlusion and dense arrangement of small targets in the electric energy meter production line inspection scenario, this study proposes a lightweight target detection algorithm based on YOLOv8n. By introducing the O-GELAN module, the algorithm achieves richer feature levels while maintaining low computational complexity. The neck architecture of feature collection, fusion, and distribution, along with the channel position attention mechanism, enables cross-layer feature fusion. Furthermore, a re-parameterized convolutional detection head is employed to enhance detection efficiency. Experiments conducted on field-collected production line data demonstrate that the improved model's mAP(0.5) and mAP(0.5:0.95) have reached 0.994 and 0.828, respectively, with a detection speed of up to 111.5 frames per second. This meets the high precision and real-time requirements of industrial scenarios and can provide auxiliary decision-making for fault elimination.

Key words: industrial production line; dense object detection; YOLOv8; feature fusion; on site data collection

引言

目前,已有许多针对电能表检定流水线的故障消缺方案。然而,这些消缺方案大多依赖于传感器的信息传递。一方面,不同传感器之间的灵敏度有所差异,需要定时检修和校准。另一方面,对于不同的故障,传感器的报警信息大多相似,具体的故障原因仍需专业人员进一步检查。因此,亟待一种能够有效利用边缘设备的辅助检测方法,降低现有方案对传感器的依赖性,减少不必要的资源损耗。而随着深度学习技术的发展以及边缘设备计算能力的提高^[1-3],近年来,各类目标检测算法也被不断提出,在自动驾驶^[4]、无人机控制^[5]等场景中已有广泛的应用。在工业流水线的故障消缺方案中,利用目标检测方法补充视觉信息以进一步辅助决策也逐渐成为有效的思路。

聚焦到电能表的检定流水线环境。首先,工业流水线场景相对单一,视角基本固定,视野死角及目标遮挡问题难以避免。其次,故障消缺所关注的常见目标(如机械臂、流水线托盘等)之间尺度差别较大。第三,检定过程中,流水线上的目标通常密集排列为长序列,彼此之间难以分辨。基于上述分析,传统的目标检测算法难以直接应用于电能表的检定流水线环境。结合高实时性、高精度的工业场景要求,本文以YOLOv8网络作为基准模型,提出了一种改进的轻量化目标检测算法。具体来说,本文的主要贡献如下:

(1)针对原始YOLOv8中骨干网络特征提取能力有限、难以应对小目标和遮挡目标的问题,基于广义高效层融合网络(Generalized efficient layer aggregation network, GELAN)^[6]和在线重参数化(Online re-parameterization, OREPA)技术^[7]提出了一种O-GELAN模块,用于替换骨干网络中的C2F模块。该模块通过更宽的结构获取更丰富的特征层次,并利用重参数化技术以降低计算量。

(2)YOLOv8模型中的颈部结构在融合不相邻特征层时,通常需要借助中间层融合,不利于各层次特征的交互与传播,影响模型在不同尺度目标和密集排列目标上的综合表现。针对该问题,设计了一种基于GOLD-YOLO^[8]的特征收集-融合-分发(Collection-fusion-allocation, CFA)结构作为颈部结构,并在其中引入通道位置注意力机制(Channel position attention mechanism, CPAM)^[9]进一步强化目标的位置信息。

(3)对YOLOv8中解耦的检测头重新改用耦合结构,并采用了一种重参数化卷积(Re-parameterization convolution, RepConv)^[10],用于代替其中的卷积操作,防止计算量过大的头结构无法满足模型在工业环境下的实时性需求。

1 相关工作

目前,目标检测方法大多可分为两类。研究者通常根据图像的基本纹理特征来定制目标区域检测算法,而随着深度学习的发展,以卷积为基础的各类深度学习模型开始逐渐受到青睐。

1.1 传统的目标检测算法

与语义分割这类需要精确描述目标边界的任务不同,目标检测任务的核心是定位和分类,即在目标区域处标定边界框并标识类别。传统的目标检测算法遵从三段式的设计理念,将目标检测方法拆分为目标区域定位、特征提取、分类识别3个阶段,完成从定位到表示再到分类的完整逻辑链。

在目标区域定位阶段,由于不同对象具有不同的大小和纵横比,相关研究方法往往采用以多尺度滑动窗口^[11]为基础的方法。在特征提取阶段,则根据不同场景定制不同的特征表示算法,如梯度方向直方图(Histogram of oriented gradient, HOG)^[12]、尺度不变特征转换(Scale-invariant feature transform, SIFT)^[13]等。最终,通过Adaboost^[14]、支持向量机(Support vector machine, SVM)^[15]等分类方法进行识别。然而,上述方案的定位阶段通常具有较高的计算成本,在背景杂乱、视角差异大的环境中特征提

取的效果也不尽如人意。随着深度学习技术的发展,传统目标检测的劣势也越发明显。

1.2 基于深度学习的目标检测算法

深度学习技术的发展有效解决了传统目标检测算法所存在的弊病,在特征表示阶段的效率及精度得到了显著提升。目前,现有的深度学习目标检测算法主要有两类。一类是以 Faster R-CNN^[16]为代表的双级检测,另一类则是以 YOLO^[17]为代表的单级检测。

双级检测先在图像中产生多个候选区域,再通过对各个候选区域进行偏移和缩放以达到收敛。Girshick等^[18]提出了 R-CNN,通过选择性搜索产生候选框,并利用 AlexNet^[19]和 SVM 分别进行特征提取与分类识别。He等^[20]提出了空间金字塔池化网络(Spatial pyramid pooling network, SPP-net),通过固定特征维度提高了检测速度。Girshick等^[21]在前期工作基础上实现了一定程度的并行,提出了 Fast-RCNN;Ren等^[16]在此基础上加以改进,提出了第一个实时速度的 Faster-RCNN。双级检测的方案通常具有较高的准确率,然而,其检测速度却难以满足工业环境的高实时性需求。

与之相对,单级检测则直接对目标进行分类和定位,以实现低时延的实时检测。Joseph等^[17]提出了 YOLO,实现了第一个单级检测器。Liu等^[22]提出的 SSD 则首次在网络的不同层检测不同尺度对象。得益于 YOLO 的端到端检测理念,其模型的可扩展性较其他模型更强,针对小目标等通用场景的改进更多样更密集^[23-25]。胡一帆等^[26]通过对 YOLOv3 改进初始化、优化其骨干网络,提出了一种轻量化的人脸检测算法。贾子豪等^[27]则在 YOLOv5 的基础上针对自动驾驶的实时场景设计了一种更快捷的交通标志检测模型。李生辉等^[28]改进 YOLOv5 特征融合方式,设计了一种更精确的卫星图像中船舶检测方法。

1.3 工业检测中的 YOLO 方法

在工业相关环境中,具有更低时延的 YOLO 方法备受青睐。研究者们通常通过改进其中的不同结构来实现和任务的针对性适配,由此派生出许多不同任务环境下的目标检测模型。

Kou等^[29]在 YOLOv3 中堆叠密集残差块来提高对带钢缺陷检测的准确率,然而,这也带来了更大的时延。Xie等^[30]在 YOLO 中引入深度可分离卷积来降低模型参数量,满足工业缺陷检测的实时性需要。Wang等^[31]在 YOLOv7 的骨干网络上添加注意力机制来提高在带钢缺陷检测任务中对小目标缺陷的检测精度。姚景丽等^[32]就 YOLOv8 中添加 3D 模块来更加深入表示轴承表面缺陷的特征。冉庆东等^[33]则对 YOLOv5 融合可变形卷积,使之能够在锂电池极片缺陷检测中实现较高的检测精度。然而,这些方法的适用环境往往是检测实体目标中的某一区域,在工业应用过程中,目前仍缺少在流水线环境中对实体目标进行检测定位的针对性算法。

2 改进 YOLOv8 模型

针对电能表检定流水线中目标尺度差异大、密集目标过多、目标遮挡等问题,提出了一种面向电能表流水线目标检测的改进 YOLOv8 算法。考虑到工业环境中高实时性需求,该模型选用体量最小的 YOLOv8n 作为基础网络,模型的整体结构如图 1 所示。与 YOLOv8 相同,模型可划分为骨干网络(Backbone)、颈部结构(Neck)以及分类检测头(Head)三部分。其中,C2F 是 YOLOv8 中用于挖掘深层特征的核心模块,主要通过叠加不同的 BottleNeck 实现特征的多维提取和深层挖掘。C2F 模块、SPPF 模块、瓶颈层(BottleNeck)结构以及基础 CBS 模块的结构也在图 1 中标出。原始图像以三通道格式为输入,经过骨干网络进行特征提取,并于颈部结构进行特征融合后,得到 3 组尺寸分别为 80×80 、 40×40 以及 20×20 的特征图,用于在分类检测头中检测不同尺度目标。

为实现不同尺度不同层次的特征提取和融合,基于 YOLOv9 中的 GELAN 结构,提出了一种轻量化模块 O-GELAN,用于代替原始 YOLOv8 骨干网络部分的 C2F 模块,在获取多层次特征的同时实现

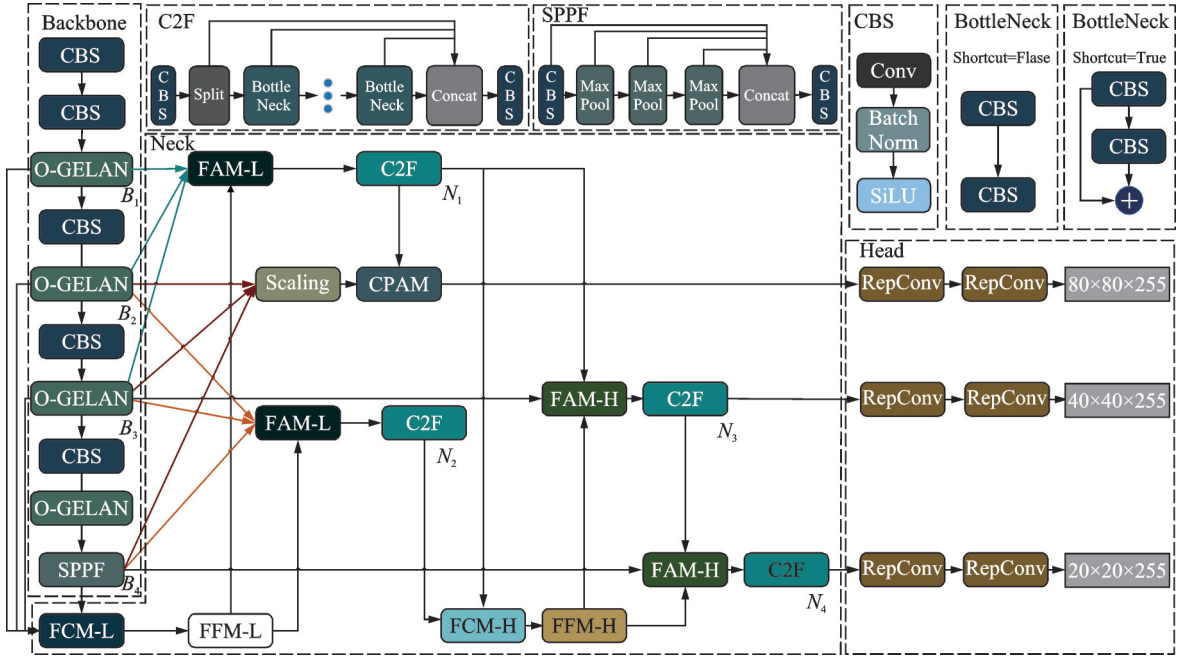


图1 改进 YOLOv8 模型结构图

Fig.1 Structure diagram of improved YOLOv8 model

计算重参数化,以较小的计算量实现更加丰富的跨层特征融合。同时,考虑到原始 Neck 的特征融合往往需要中间层进行过渡,容易造成特征融合不充分等问题,以 YOLOv8 为基础,设计了基于 GOLD-YOLO 的特征聚集-分发机制,通过 CFA 结构实现特征融合处理。考虑到电能表流水线中存在较多小目标密集排列场景,在特征融合阶段引入 CPAM,以聚焦目标空间位置,实现准确识别。此外,为进一步减少模型参数,设计了 RepConv 结构作为检测头,在降低模型计算需求的同时,提高模型在检测阶段对图像特征的解析能力。改进 YOLOv8 模型的训练过程如算法 1 所示。

算法 1 改进 YOLOv8 模型训练

输入: 由待检测图片 I_m 和真实标签 I_{GT} 组成的数据集及其迭代器 DataLoader, 训练轮次 N , O-GELAN 特征提取子模块个数 M

输出: 可用于目标检测的收敛模型 Model

- (1) for epoch in $1, 2, \dots, N$
- (2) for batch in DataLoader do
- (3) $I_m, I_{GT} = \text{batch}$
- (4) #在 BackBone 中
- (5) $B_0 = \text{CBS}(I_m)$ #初始化特征提取
- (6) for i in $1, 2, \dots, M$:
- (7) $B_i = \text{CBS}(\text{O-GELAN}(B_{i-1}))$ #经过 O-GELAN 进行特征提取
- (8) end for
- (9) $B_i = \text{SPPF}(B_i)$ #最后一层需经过 SPPF
- (10) #在 Neck 中
- (11) $N_1, N_2 = \text{C2F}(\text{CFA}(B_1, B_2, B_3, B_4))$ #第一阶段 CFA

```

(12)       $N_3, N_4 = \text{C2F}(\text{CFA}(B_4, N_1, N_2))$  #第二阶段 CFA
(13)       $F_{\text{scale}} = \text{Scaling}(B_2, B_3, B_4)$  #CPAM 前整合基础特征
(14)       $F_{\text{CPAM}} = \text{CPAM}(F_{\text{scale}}, N_1)$  #CPAM 计算
(15)      #在 Head 中
(16)      Box, Class = Head( $F_{\text{CPAM}}, N_3, N_4$ ) #得到定位框 Box 和类别 Class
(17)      Class_Loss = CE(Class, GT.class) #计算分类损失
(18)      Box_Loss = CIoU(Box, GT.box) #计算边框损失
(19)      #反向传播
(20)      Class_Loss.backpropagate()
(21)      Box_Loss.backpropagate()
(22)      End for
(23)      #更新权重
(24)      Update_Weights()
(25)      end for
(26)      return Model

```

2.1 O-GELAN 模块

在原始 YOLOv8 模型中,用于提取深层特征的 C2F 模块结构单一,通过反复堆叠瓶颈结构提取深层特征,对不同层次之间的特征差异与信息融合考虑较少,从而影响模型在小目标以及被遮挡目标上的表现。因此,提出了一种 O-GELAN 模块,其结构如图 2 所示,用于替换骨干网络部分的 C2F 模块以进一步丰富特征信息的层次。其中, GELAN 在 ELAN^[34] 的基础上,利用带有 RepConv 的 CSPNet^[35] 替换原始 ELAN 中的卷积,实现对原特征图不同层次特征的提取及信息融合。而本文则在此基础上结合 OREPA 方法,提出一种轻量化的计算方案,以简化 GELAN 中的并行计算,提高模型计算效率。

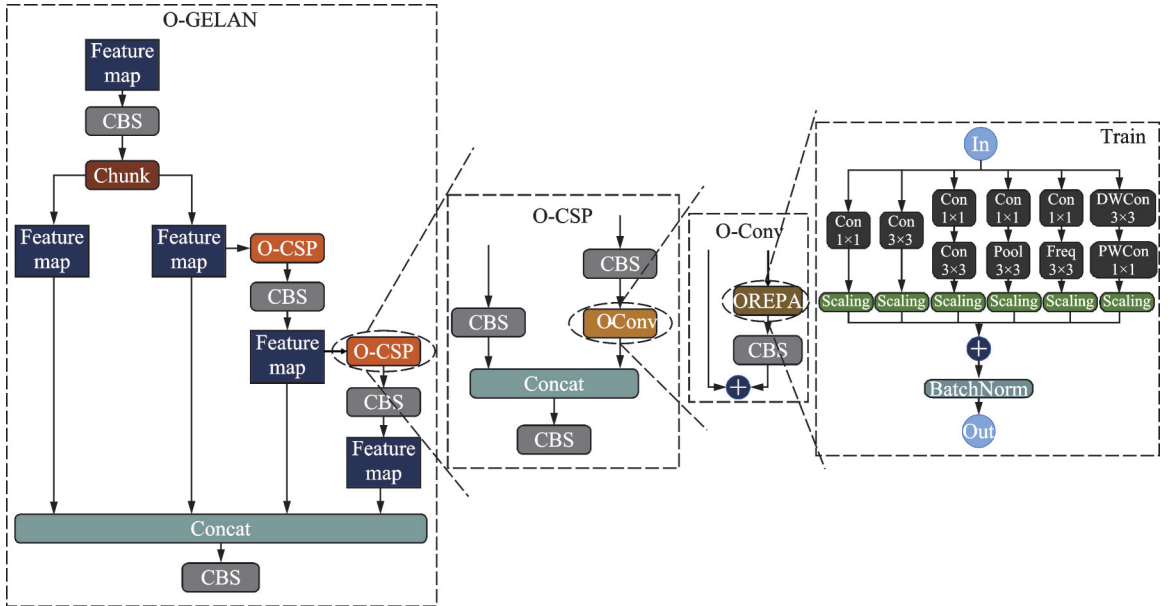


图2 O-GELAN 模块结构图

Fig.2 Structure diagram of O-GELAN module

如图2所示,O-GELAN模块中,原始特征经过 1×1 尺寸的卷积后,按通道被拆分为两个子图。其中的一个子特征图将经由多次O-CSP结构进行深度特征提取,得到多个包含不同层次信息的特征图。每个O-CSP结构将接收到的特征信息分别经过两个卷积核尺寸为 1×1 的卷积进行升维,并使其一经过多层带有OREPA模块的O-Conv结构获得尺度更为丰富的特征用于特征的拼接整合。最终将原始特征图拆解的子图作为先验信息,与多个特征图进行特征拼接,获得信息更为丰富的特征结果。

如图2最右侧所示,OREPA模块包含6种并行结构的卷积,这些卷积模块在训练过程中为模型提供了更强的表征能力,相较于固定感受野的卷积,这样的设计能够使模型在检测过程中更好地注意到独立目标的检测边界,避免了大尺度下边界模糊或小尺度下边界锯齿等问题。然而,这也产生了更多的计算资源消耗。因此,利用在线重参数化技术,在模型推理阶段将6个并行分支的参数合并到一起,实现高效快速推理。

2.2 颈部结构

原始的颈部结构只能直接融合相邻特征层的特征信息,在面对不相邻特征层时,往往需要借助中间层进行特征信息传递,特征的融合效率和信息整合程度受到影响。考虑到骨干网络中的O-GELAN模块提取到的特征层次更加丰富,这一问题势必愈发明显。基于上述分析,设计采用一种结合CPAM的CFA结构作为Neck实现特征融合。

2.2.1 收集-融合-分发机制

结合图1中的标注,对于模型骨干结构以及颈部结构中所获得的特征图做出如下假设。首先,经过Backbone进行特征提取后,模型将生成尺度不同的4个基础特征图。该阶段的信息尚未进行特征融合,其包含的特征层次之间相对独立,故将其作为初始阶段的特征,分别记为 B_1 、 B_2 、 B_3 和 B_4 。其次,对于经过一次信息融合处理的特征,将其记为 N_1 和 N_2 ;这类特征实现了一定程度的特征交互,但仍需进一步处理。最后,对于信息融合程度更高的特征,将其记为 N_3 和 N_4 ,该类特征可直接输入检测头进行解析。

特征收集-融合-分发结构通过采集各个特征图的初始信息直接进行跨层信息聚合,将融合后的信息针对性地分发以进行后续处理。具体而言,CFA结构主要包括特征收集模块(Feature collection module, FCM)、特征融合模块(Feature fusion module, FFM)以及特征分发模块(Feature allocation module, FAM)三部分,每个部分针对原始特征和初级融合特征采用特异化的处理,其中,FCM模块的结构如图3所示。根据处理信息的

层次不同,将FCM分为针对原始特征的FCM-L和针对初级融合特征的FCM-H两种处理模式。FCM-L是对Backbone中不同尺度的原始特征图 B_1 、 B_2 、 B_3 和 B_4 进行尺度放缩并按通道拼接,以实现不相邻特征层的直接融合。以 $F_{\text{FCM-L}}$ 代表FCM-L模块的输出结果,该过程可表示为

$$F_{\text{FCM-L}} = \text{Concat}(\text{resize}(B_1), B_2, \text{resize}(B_3), \text{resize}(B_4)) \quad (1)$$

式中: $\text{resize}(\bullet)$ 表示池化或上采样操作,以进行适当的尺度缩放,Concat($\bullet$)代表按通道拼接操作。而考虑到特征层次更深的 B_4 特征图包含更多与类别相关的语义特征,FCM-H则将经过一次信息融合的特征 N_1 、 N_2 与原始特征 B_4 进行尺度缩放与拼接,以进一步融合高级别信息促进定位后的分类。以 $F_{\text{FCM-H}}$ 代表FCM-H模块的输出结果,则该过程可表示为

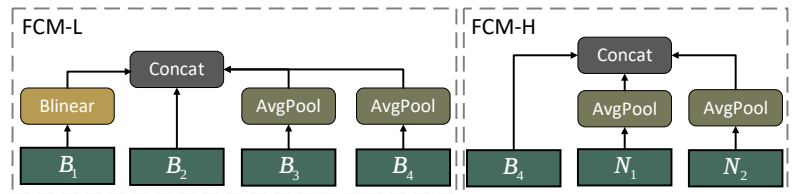


图3 特征收集模块结构图

Fig.3 Structure diagram of feature collection module

$$F_{FCM-H} = \text{Concat}(B_4, \text{resize}(N_1), \text{resize}(N_2)) \quad (2)$$

此后,由FCM收集的特征将经过FFM进行特征信息融合,FFM的结构如图4所示。针对不同阶段的特征,FFM采用了不同的特征融合方式。FFM同样可分为FFM-L和FFM-H两部分。针对原始特征,FFM接收FCM拼接的特征后,经由多个RepConv实现进一步的特征交互,以析出更丰富的语义特征,并拆分成 F_{FFM-L_1} 和 F_{FFM-L_2} ,分发至不同的通路以实现针对性融合。该过程可表示为

$$F_{FFM-L_1}, F_{FFM-L_2} = \text{Split}(\text{RepBlock}(F_{FCM-L})) \quad (3)$$

式中:RepBlock($\bullet$)表示包括CBS和RepConv的特征交互过程,Split($\bullet$)表示按通道切分操作。

对于初级融合的特征,FFM接受FCM拼接的特征后,经由多组Transformer编码器进行进一步的特征编码,利用自注意力机制对全局特征的高处理能力实现不同特征层的精细融合,并拆分为 F_{FFM-H_1} 和 F_{FFM-H_2} 分发至不同的通路。该过程可表示为

$$F_{FFM-H_1}, F_{FFM-H_2} = \text{Split}(\text{TE}(F_{FCM-H})) \quad (4)$$

式中: F_{FFM-H_1} 和 F_{FFM-H_2} 分别代表切分后的两组特征结果,TE($\bullet$)代表Transformer编码器过程。

经过FFM融合的特征结果按如图1所示进行拆分,分发至不同的串行通路,特征分发模块FAM的结构如图5所示。针对不同阶段的特征,FAM采用不同的特征预处理策略。考虑Backbone产出的特征直接来自原始图像,特征处理阶段产生的噪声和臆测信息相对较少。因此,对于低级别原始特征,FAM-L对其相邻3层进行放缩和拼接,使之更多地参与特征表示。以 F_{FAM-FL} 表示该部分计算结果,该过程可表示为

$$F_{FAM-FL} = \text{Concat}(\text{resize}(B_{i-1}), \text{CBS}(B_i), \text{resize}(B_{i+1})) \quad (5)$$

与之相反,FAM-H在高级别特征阶段更侧重于对图像语义信息的表示。因此,同时采用了原始特征与经过一次融合的特征,由串联方式累积融合特征,以期在检测头解析中能够提供抽象程度更高的深层信息。以 F_{FAM-FH} 表示该部分计算结果,该过程可表示为

$$F_{FAM-FH} = \text{Concat}(\text{resize}(B), N) \quad (6)$$

令 F_{FAM-F} 作为 F_{FAM-FL} 和 F_{FAM-FH} 的统一表示,则FAM需要对FFM产生的融合特征进行降维和尺

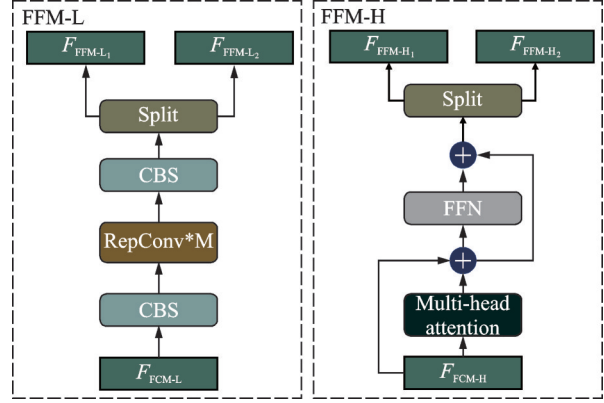


图4 特征融合模块结构图

Fig.4 Structure diagram of feature fusion module

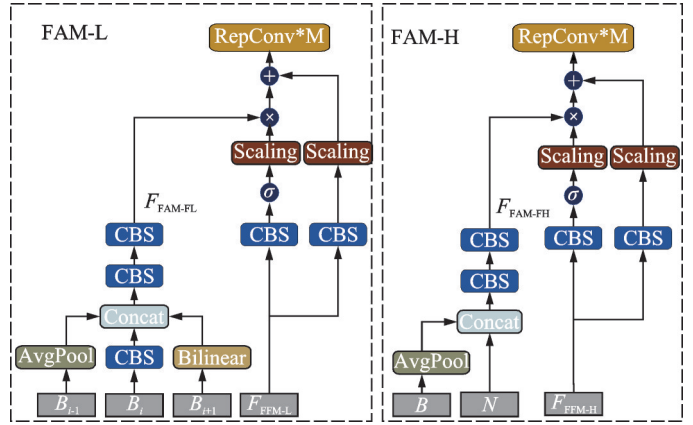


图5 特征分发模块结构图

Fig.5 Structure diagram of feature allocation module

度变换操作,以将其与 F_{FAM-F} 进行有效融合。一方面,FFM 所产生的融合特征可作为跨层特征信息的补充,记为 F_{FAM-pd} ;另一方面,该融合特征也可作为权重,对原始特征或初步融合的特征进行加权,记为 F_{FAM-wt} ,加强在后续操作中对多维特征所反映的重点区域的关注。二者的计算过程可分别表示为

$$F_{FAM-pd} = \text{resize}(\text{CBS}(F_{FFM})) \quad (7)$$

$$F_{FAM-wt} = \text{resize}(\sigma(\text{CBS}(F_{FFM}))) \quad (8)$$

式中: F_{FFM} 代表 FFM 的输出特征, $\text{CBS}(\cdot)$ 表示 CBS 操作, $\sigma(\cdot)$ 表示 Sigmoid 函数。

综上所述,令 F_{FAM} 为 FAM 的输出结果,则结合图 5, FAM 模块的处理过程可表示为

$$F_{FAM} = F_{FAM-F} \otimes F_{FAM-wt} + F_{FAM-pd} \quad (9)$$

2.2.2 通道位置注意力机制

在 CFA 结构的基础上,考虑到电能表流水线场景下存在较多的小目标密集排列情况,模型需要更加精准的空间位置信息,而对比原始 Neck, CFA 结构已得到融合程度足够高的特征信息。相比进一步获取大量的空间位置信息,本文的优化方向在于从融合特征中获取位置相关的特征信息。为此,改进 YOLOv8 在 CFA 中引入了 CPAM,其结构如图 6 所示。

结合图 1 和图 6 可知,在本研究中, CPAM 的处理对象主要包括两个部分:一是对 B_1 、 B_2 以及 B_3 特征图进行缩放拼接后的特征图,记为 F_{scale} ;二是 B_2 、 B_3 和 B_4 进行初级融合后的特征图 N_1 。

在通道注意力部分,即图 6 上方。 F_{scale} 通过平均池化层和一维卷积聚焦通道上的信息,并将其作为权重与原始特征相乘,由此完成通道空间到平面空间的信息引入,得到特征图 F_{CA} 。在位置注意力部分,即图 6 下方。 F_{CA} 与 N_1 相加的结果将按二维空间上的宽和高两个方向分别进行平均池化,并对该结果拼接降维后切分。此时,切分结果中分别包含了二维空间中正交方向的独立位置信息,将其转换为权重后对求和特征加权,并进一步交互得到 F_{CPAM} 。

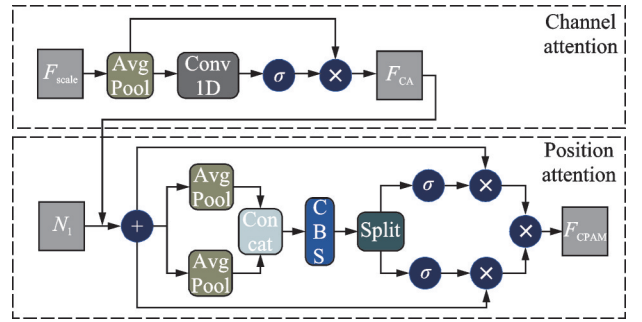


图 6 通道位置注意力机制结构图

Fig.6 Structure diagram of channel position attention mechanism

2.3 RepConv 检测头

原始 YOLOv8 的 Head 对检测定位与分类任务采用解耦设计,用两个独立的卷积分支分别进行检测与分类。与过去的技术相比,这一方案显著提升了检测与分类的准确性,但也带来更大的计算量。而电能表检定流水线环境对实时数据的检测速度有较高的要求,过大的计算量难以匹配工业场景的需求。基于上述分析,本文重新将 Head 中的检测与分类任务进行耦合设计,并引入 RepConv 替换卷积模块,在降低参数量的同时尽可能保持较高的检测精度。其结构如图 7 所示。

在训练阶段,如图 7 左侧所示, RepConv 是一个由 BatchNorm、 3×3 尺寸卷积以及 1×1 尺寸卷积组成的并行结构,合并后经 SiLU 层激活以获取多尺度特征信息,以此来保证其有足够的参数量和足够多元的感受野较强的特征表达能力。而在推理阶段,如图 7 右侧所示,

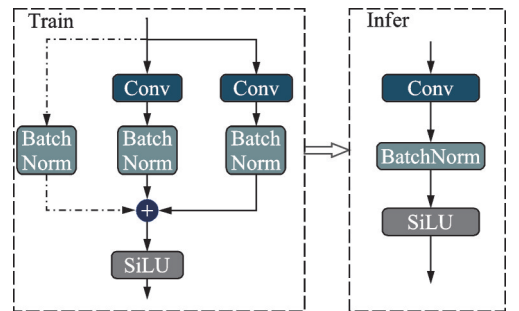


图 7 RepConv 结构图

Fig.7 Structure diagram of RepConv

RepConv采用结构化重参数技术,将3个并行分支的卷积结构和BN层的参数合并为一条推理线路。因此,模型本身依然保持了3条支路特有的表达能力,但推理过程的运算时间被大幅缩减,实际推理的部分仅需一个 3×3 尺寸卷积的CBS块,从而在保证精度的前提下有效提高了计算效率。

3 实验与分析

所做实验在Ubuntu20.04、CUDA 11.4环境下进行,CPU采用Intel(R) Xeon(R) Gold 5318Y CPU @ 2.10 GHz,显卡采用NVIDIA GeForce RTX 3090,显存约24 GB。以Pytorch为主要框架实现改进YOLOv8模型。实验中,批次大小设为32,所采用优化器为SGD优化器,初始学习率设为0.01,并采用余弦退火的学习率衰减策略。模型训练500轮次达到收敛。

3.1 数据集说明

所采用的数据采集自现场真实的智能电表检定流水线环境。在实际流水线环境下,就3种不同角度采集60帧率的视频,对其进行随机裁切并标注。最终在各角度下分别标注200张,共计600张图片,并以 640×640 大小输入模型训练。数据集中部分示例如图8所示。

每张图片所包含的检测目标包括:抓取机器人(Robot)、输送线上的空托盘(Tray)以及辊筒线上装满电能表的周转箱(Box)。数据集中各标签数量分布如表1所示,按照7:3的比例随机划分为训练集和测试集。



图8 数据集中各个角度示例图

Fig.8 Example images from various angles in dataset

表1 数据集中各标签的分布明细表
Table 1 Distribution details of each label in dataset

检测目标	角度1	角度2	角度3	总计
Robot	601	397	400	1 398
Tray	688	0	763	1 451
Box	58	843	425	1 326

3.2 评价指标

采用目标检测中的常见评价指标,主要包括精确率(Precision)、召回率(Recall)、平均精度(mean Average precision, mAP)。除此之外,引入评价模型效率的指标,主要包括参数量(Parameters)、浮点运算数(Floating point operations, FLOPs)以及每秒处理帧数(Frames per second, FPS)。

精确率表示检测结果中检测正确的概率,通常被表述为正确预测个数与被预测总数之比;召回率则表示对于所有含有目标的图像中,真正存在目标实体的概率,通常被表述为正确预测个数与图像中所含目标总数之比,其计算公式分别可表示为

$$\text{Precision} = \frac{\text{TP}}{\text{FP} + \text{TP}} \quad (10)$$

$$\text{Recall} = \frac{\text{TP}}{\text{FN} + \text{TP}} \quad (11)$$

式中:TP表示预测正确的样本数,FP表示将非目标预测为目标的样本数,FN表示对真实目标预测错误的样本数。

此外,根据不同阈值下的精确率和召回率,可分别以精确率和召回率为横纵坐标构建一条曲线,平均精度AP即为该曲线的积分结果。对各个类别下的AP求均值,即可得到mAP,其计算公式分别为

$$\text{AP} = \int_0^1 \text{Precision}(r) dr \quad (12)$$

$$\text{mAP} = \frac{1}{n} \sum_{i=1}^n \text{AP}_i$$

(13)

就 mAP 而言,采用两种不同阈值进行评价。mAP(0.5)代表在 IoU=0.5 阈值下的 mAP,mAP(0.5:0.95)表示在以 0.05 为步长,IoU 取值范围在 0.5 到 0.95 下 mAP 的平均值。

3.3 对比实验

为了验证所改进 YOLOv8 模型的有效性,设计两组对比实验。一是与 YOLOv8 系列下,不同体量不同检测宽度的模型进行对比,以验证模型在低计算量的同时具有较好的性能。二是与当前主流的深度学习模型进行对比,旨在验证模型在针对电能表流水线目标尺度差异大、目标遮挡与密集排列问题上的有效性。

3.3.1 与 YOLOv8 系列模型的对比

在同 YOLOv8 系列模型的对比中,采用基线模型包括:YOLOv8-n、YOLOv8-s、YOLOv8-m、YOLOv8-l 以及 YOLOv8-x。对比实验结果如表 2 所示。

表 2 与 YOLOv8 系列模型对比实验结果
Table 2 Comparison of experimental results with YOLOv8 series models

模型	Precision	Recall	mAP(0.5)	mAP(0.5:0.95)	Parameters/10 ⁶	FLOPs/10 ⁹	FPS
YOLOv8-n	0.974	0.974	0.990	0.815	3.0	8.1	240.9
YOLOv8-s	0.965	0.979	0.988	0.825	11.1	28.4	162.9
YOLOv8-m	0.977	0.975	0.990	0.825	25.8	78.7	109.8
YOLOv8-l	0.971	0.978	0.988	0.826	43.6	164.8	65.8
YOLOv8-x	0.969	0.980	0.989	0.825	68.1	257.4	48.4
本文方法	0.968	0.980	0.990	0.827	6.3	8.7	111.5

由表 2 可知,所提出的方法在 mAP(0.5)和 mAP(0.5:0.95)上的表现最为突出,分别达到 0.990 和 0.827,综合表现最优。对比原始 YOLOv8-n,本文方法在 Recall 和 mAP(0.5:0.95)上分别提升了 0.6% 和 1.5%。对比表现较好的 YOLOv8-m,所提出方法也在 Recall 和 mAP(0.5:0.95)上分别提高了 0.005 与 0.002。尽管在 Precision 上,改进 YOLOv8 模型降低了 0.9%,但其参数数量和 FLOPs 却分别达到其 24.4% 和 11.1%,模型得到了显著优化。同 YOLOv8 系列模型的对比实验结果如图 9 所示。

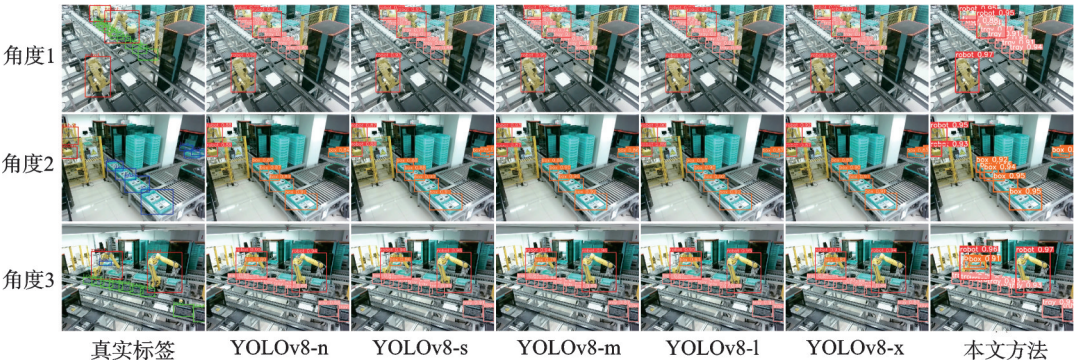


图 9 与 YOLOv8 系列模型实验结果对比

Fig.9 Comparison of experimental results with YOLOv8 series models

3.3.2 与其他深度学习方法对比

在同当前主流的深度学习方法对比实验中,用于对比的模型包括:双级检测的Faster-RCNN^[16]与Cascade-RCNN^[36],YOLO系模型DAMO-YOLO^[37]和GOLD-YOLO^[8],利用自注意力的新兴目标检测方案DETR^[38]与RT-DETR^[39]。对比实验结果如表3所示。

表3 与其他深度学习方法对比实验结果

Table 3 Comparison of experimental results with other deep learning methods

模型	mAP(0.5)	mAP(0.5:0.95)	Parameters/ 10^6	FLOPs/ 10^9	FPS
Faster-RCNN	0.947	0.628	41.3	209.1	41.4
Cascade-RCNN	0.975	0.759	69.1	389.3	24.0
DAMO-YOLO	0.952	0.444	8.5	18.1	112.6
GOLD-YOLO	0.984	0.788	5.6	12.1	119.3
DETR	0.980	0.707	36.7	91.3	44.6
RT-DETR	0.983	0.803	31.9	103.4	106.7
本文方法	0.990	0.827	6.3	8.7	111.5

由表3可知,本文方法对比目前主流深度学习方法在mAP上具有最佳表现,在电能表流水线目标检测任务中,本文方法在mAP(0.5)上对比GOLD-YOLO提高了0.6%,而在mAP(0.5:0.95)上对比RT-DETR则提高了3.0%,检测精度提升明显。而在模型体量方面,所提模型具有最低的FLOPs,参数量仅次于GOLD-YOLO,较之仅增加了0.7 M,FPS也保持在110以上,整体性能表现较好。各个模型的对比实验结果如图10所示。

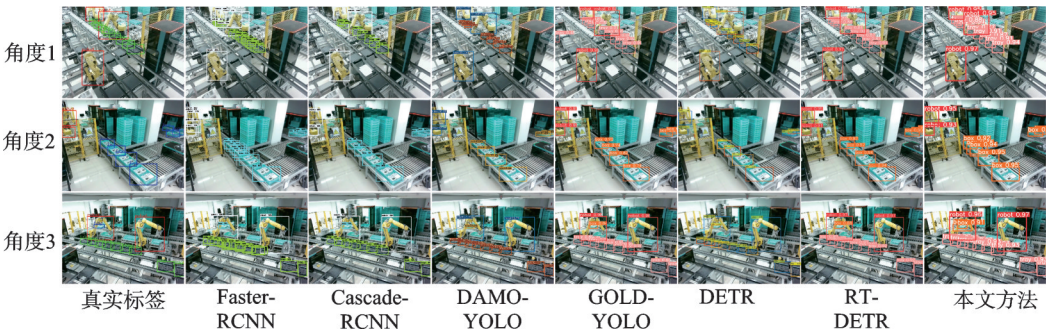


图10 与其他深度学习方法实验结果对比

Fig.10 Comparison of experimental results with other deep learning methods

3.4 消融实验

为验证所提出的各个模块有效性,针对各部分设计了消融实验。为方便表述,将含有不同模块的实验模型分别记为A、B、C模型,各个模型所具备的设计如表4所示。记A模型为在YOLOv8n的基础上,仅在Backbone中添加O-GELAN替换C2F结构的对照模型;B模型为对A模型修改Neck实现CFA结构的对照模型;C模型为在B模型基础上添加CPAM的对照模型;本文方法则在C模型中修改了Head结构。实验在相同参数下进行,实验结果如表5所示。由表5可看出,Tray目标的各项指标精度较均值大多略有下降。考虑这是由以下原因造成:首先,该对象数量较其他两类更多;其次,该类目标体积更小,检测难度更大;最后,在

表4 改进YOLOv8的消融实验设计说明表

Table 4 Description table of ablation experimental design for improved YOLOv8

模型	O-GELAN	CFA 结构	CPAM	RepConv
YOLOv8n				
A	✓			
B	✓	✓		
C	✓	✓	✓	
本文方法	✓	✓	✓	✓

表 5 改进YOLOv8的消融实验结果

Table 5 Results of ablation experiments for improved YOLOv8

模型	Class	Precision	Recall	mAP(0.5)	mAP(0.5:0.95)	Paramters/M	FLOPs/G	FPS
YOLOv8-n	All	0.974	0.974	0.990	0.815	3.0	8.1	240.9
	Robot	0.982	0.969	0.989	0.863			
	Tray	0.948	0.984	0.986	0.762			
	Box	0.991	0.970	0.994	0.821			
A	All	0.966	0.972	0.988	0.816	2.5	6.8	184.3
	Robot	0.987	0.948	0.983	0.868			
	Tray	0.929	0.990	0.988	0.764			
	Box	0.980	0.978	0.993	0.815			
B	All	0.963	0.980	0.990	0.818	5.6	9.0	118.0
	Robot	0.972	0.973	0.990	0.867			
	Tray	0.938	0.993	0.986	0.766			
	Box	0.980	0.975	0.994	0.822			
C	All	0.970	0.976	0.989	0.824	5.6	9.2	106.8
	Robot	0.979	0.970	0.991	0.880			
	Tray	0.951	0.987	0.984	0.757			
	Box	0.981	0.970	0.993	0.834			
本文方法	All	0.968	0.980	0.990	0.827	6.3	8.7	111.5
	Robot	0.979	0.973	0.988	0.883			
	Tray	0.945	0.987	0.986	0.769			
	Box	0.980	0.981	0.994	0.828			

检定流水线环境中,Tray的密集排列问题和遮挡问题相对更加严重,检测精度受限。消融实验中,各个模型的推理结果如图 11 所示。

结合图 11 进行分析,不同模块的作用方向并不完全一致,各个模块均发挥了一定的效用。对于 O-GELAN 模块来说,模型在 mAP(0.5:0.95)上取得了 0.001 的提高,但其在 Tray 目标上的 Recall、mAP(0.5)和 mAP(0.5:0.95)指标上分别提升了 0.6%、0.2% 和 0.3%。考虑这是得益于 O-GELAN 更宽的结

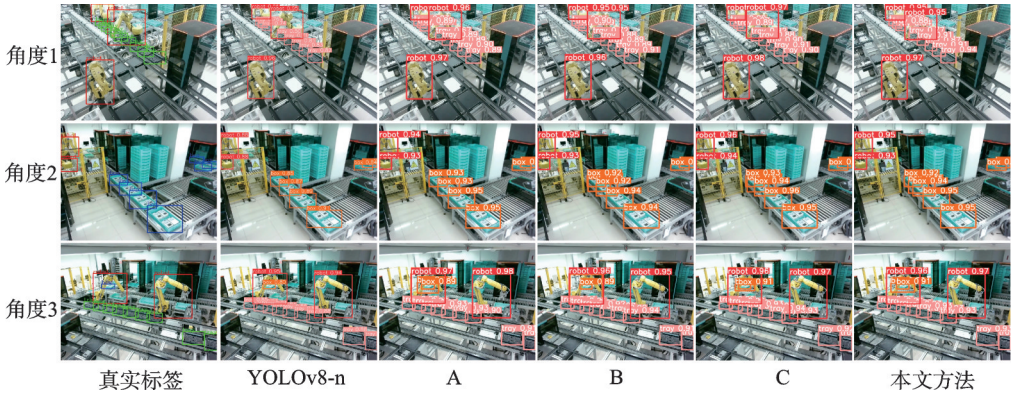


图 11 改进YOLOv8的消融实验结果对比

Fig.11 Comparison of ablation experimental results for improved YOLOv8

构从而获得了尺度更加多样的特征信息,使之在小目标上的表现有所提高。由于 OREPA 的重参数化计算,O-GELAN 的引入非但没有增加计算量,反而较之在参数量和 FLOPs 上分别下降了 16.7% 和 16.0%。

特征收集-融合-分发结构则对各类目标均有一定程度的增益。整体上,CFA 在 $mAP(0.5)$ 和 $mAP(0.5:0.95)$ 上均带来了 0.2% 的提升。然而,对于大目标 Robot,模型在 Precision 上却下降了 1.5%;而在相对较大的目标 Box 上,模型则在 Precision 上基本持平;在小目标 Tray 上 Precision 则提升了 1.0%。考虑由于 CFA 对浅层特征信息的融合处理较少,导致其在大尺度目标上的决策相对激进,而在小尺度目标上的决策愈加保守。CPAM 的引入在一定程度上缓解了该问题。整体上,CPAM 在 $mAP(0.5:0.95)$ 上提高了 0.7%,在 Precision 上,Robot、Tray 和 Box 较原先分别提高了 0.7%、1.4% 和 0.1%,但其在 Recall 值上也有一定的牺牲。

此外,引入 RepConv 作为检测头后,模型的 FLOPs 和 FPS 得到了改善。FLOPs 下降为原先的 94.6%,FPS 则提高了 4.4%。精度方面,考虑由于 RepConv 提供了尺度更多样的感受野,模型整体上 $mAP(0.5:0.95)$ 值提高了 0.4%,在 Robot 和 Tray 目标上分别具有 0.3% 和 1.6% 的提升。

在电能表检定流水线所存在的密集小目标问题、目标遮挡问题等情况,O-GELAN 和 CFA 模块也在其中发挥了重要作用。针对密集小目标情况,O-GELAN 模块对比 CFA 具有更优秀的特征表示能力。C2F 模块以及 O-GELAN 模块在相同骨干网络层数下的 EigenGradCAM 可视化结果图如图 12 所示,图中的热量分布代表了该模块的注意力分布。从图 12 中可看出,在托盘密集分布的场景下,C2F 模块更加关注目标的中心区域,而 O-GELAN 模块更加关注目标的边缘区域。具体来说,C2F 关注的更多的是托盘的纹理变化,O-GELAN 则更倾向于托盘上的能够表示边界范围的特征。对于边界区分不明显的情况,O-GELAN 对比 C2F 具有明显优势。

而针对目标遮挡问题,一方面,O-GELAN 在特征提取时能够保持多分支并行的高特征表示能力,可获取不同尺度的感受野信息,对于被部分遮挡的目标而言,提供了识别检测的可能。另一方面,CFA 对于不同尺度的特征提供了更加深层的特征融合方式,针对不同

层次的特征实现了针对性处理。为了验证上述模块对于目标遮挡问题的有效性,图 13 给出了两个模块在相同颈部结构层数下的 EigenGradCAM 可视化结果,以表征其所关注的重心变化。由图 13 可知,对比 C2F 模块,在遮挡问题方面,O-GELAN 表现出了更加灵敏的关注能力。对于被遮挡的箱子,O-GELAN 的注意区域更为激进;对于被遮挡的托盘,O-GELAN 则能明显识别到该目标的特点。而对比常规的颈部结构,CFA 作为颈部结构在该问题上则更加聚焦。对于被遮挡的箱子,CFA 所带来的特征解析使模型能够更加聚焦在目标物体的部分结构上;对于被机器人遮挡的托盘,CFA 虽然没有准确得识别出,但边框已经不再局限于整个机器人,关注区域发生了明显的漂移。

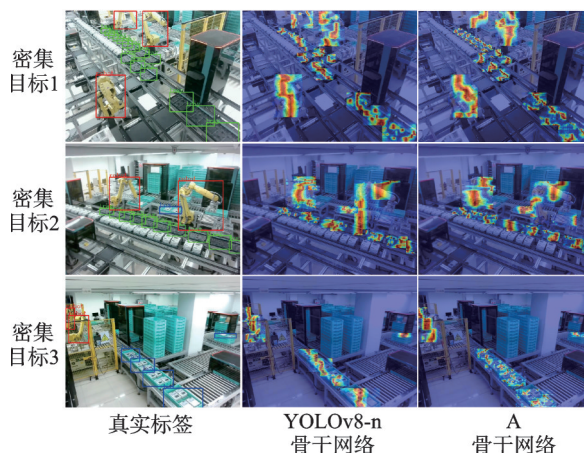


图 12 不同模块在密集目标上的 EigenGradCAM 可视化结果

Fig.12 EigenGradCAM visualization results for different modules on dense targets

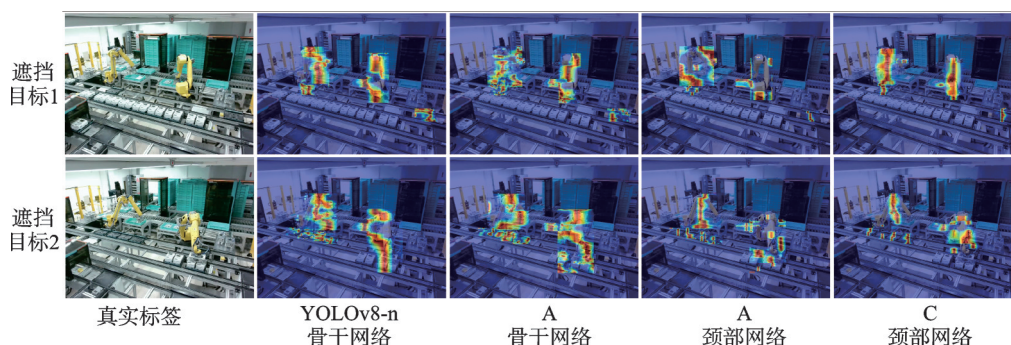


图 13 不同模块在遮挡目标上的EigenGradCAM可视化结果图

Fig.13 EigenGradCAM visualization results of different modules on occluded targets

4 结束语

为了对电能表流水线检定环境中的故障检测与消缺补充视觉信息,降低对传感器信息的依赖,进一步减少维护及人工成本。基于YOLOv8设计了一种针对电能表流水线检定环境的轻量化目标检测方案,以捕捉检定环境中的主要目标。首先,针对电能表检定流水线上小目标丰富以及目标的遮挡问题,提出了一种轻量化的O-GELAN模块,用于在骨干网络阶段进行更丰富的特征层次提取。其次,针对流水线上目标尺度差距大,密集排列场景多的问题,设计了CFA结构的瓶颈网络进行更直观的特征融合和深层交互,并引入CPAM以提供更精确的位置信息处理密集目标。最后,针对工业环境的实时性要求,对检测头中的分支进行耦合设计并采用RepConv代替卷积,以进一步实现轻量化部署。实验在电能表检定流水线的现场真实数据集上进行,结果表明,所提出的改进YOLOv8算法具有更高的准确率、更小的计算量和模型体量以及更快的检测速度。然而,该模型对于其他流水线环境的泛化效果还需进一步讨论验证。

参考文献:

- [1] KHELIFI H, LUO S, NOUR B, et al. An optimized proactive caching scheme based on mobility prediction for vehicular networks[C]//Proceedings of 2018 IEEE global communications conference (GLOBECOM). [S.l.]: IEEE, 2018: 1-6.
- [2] BLANCO-FILGUEIRA B, GARCIA-LESTA D, FERNÁNDEZ-SANJURJO M, et al. Deep learning-based multiple object visual tracking on embedded system for IoT and mobile edge computing applications[J]. IEEE Internet of Things Journal, 2019, 6(3): 5423-5431.
- [3] CHANDAKKAR P S, LI Y, DING P L K, et al. Strategies for re-training a pruned neural network in an edge computing paradigm[C]//Proceedings of 2017 IEEE International Conference on Edge Computing (EDGE). [S.l.]: IEEE, 2017: 244-247.
- [4] 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.
XIE Yuhong, XIE Yuan, CHEN Liang, et al. Object detection in real-world hazy scene[J]. Journal of Computer-Aided Design & Computer Graphics, 2021, 33(5): 733-745.
- [5] 侯鑫,曲国远,魏大洲,等.基于迭代稀疏训练的轻量化无人机目标检测算法[J].计算机研究与发展,2022, 59(4): 882-893.
HOU Xin, QU Guoyuan, WEI Dazhou, et al. A lightweight UAV object detection algorithm based on iterative sparse training [J]. Journal of Computer Research and Development, 2022, 59(4): 882-893.
- [6] WANG C Y, YEH I H, LIAO H Y M. YOLOv9: Learning what you want to learn using programmable gradient information [J]. arXiv preprint arXiv:2402.13616, 2024.
- [7] HU M, FENG J, HUA J, et al. Online convolutional re-parameterization[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. [S.l.]: IEEE, 2022: 568-577.
- [8] WANG C, HE W, NIE Y, et al. Gold-YOLO: Efficient object detector via gather-and-distribute mechanism[J]. Advances in

- Neural Information Processing Systems, 2024, 36: 51094-51112.
- [9] KANG M, TING C M, TING F F, et al. ASF-YOLO: A novel YOLO model with attentional scale sequence fusion for cell instance segmentation[J]. arXiv preprint arXiv:2312.06458, 2023.
- [10] DING X, ZHANG X, MA N, et al. RepVGG: Making VGG-style convnets great again[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. [S.l.]: IEEE, 2021: 13733-13742.
- [11] HUANG L, YANG Y, DENG Y, et al. Densebox: Unifying landmark localization with end to end object detection[J]. arXiv preprint arXiv:1509.04874, 2015.
- [12] DALAL N, TRIGGS B. Histograms of oriented gradients for human detection[C]//Proceedings of 2005 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR'05). [S.l.]: IEEE, 2005: 886-893.
- [13] LOWE D G. Distinctive image features from scale-invariant keypoints[J]. International Journal of Computer Vision, 2004, 60: 91-110.
- [14] FREUND Y, SCHAPIRE R E. A decision-theoretic generalization of on-line learning and an application to boosting[J]. Journal of Computer and System Sciences, 1997, 55(1): 119-139.
- [15] WANG L. Theory of support vector machines: Theory and applications[M]. [S.l.]: springer, 2005.
- [16] REN S, HE K, GIRSHICK R, et al. Faster R-CNN: Towards real-time object detection with region proposal networks[J]. Advances in Neural Information Processing Systems, 2015, 28: 1440-1448.
- [17] JOSEPH R, 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. [S.l.]: IEEE, 2016: 779-788.
- [18] GIRSHICK R, DONAHUE J, DARRELL T, et al. Rich feature hierarchies for accurate object detection and semantic segmentation[C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. [S.l.]: IEEE, 2014: 580-587.
- [19] KRIZHEVSKY A, SUTSKEVER I, HINTON G E. ImageNet classification with deep convolutional neural networks[J]. Communications of the ACM, 2017, 60(6):84-90.
- [20] HE K, ZHANG X, REN S, et al. Spatial pyramid pooling in deep convolutional networks for visual recognition[J]. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2015, 37(9): 1904-1916.
- [21] GIRSHICK R. Fast R-CNN[C]//Proceedings of the IEEE International Conference on Computer Vision. [S.l.]: IEEE, 2015: 1440-1448.
- [22] LIU W, ANGUELOV D, ERHAN D, et al. SSD: Single shot multibox detector[C]//Proceedings of the 14th European Conference on Computer Vision. Amsterdam: Springer, 2016: 21-37.
- [23] REDMON J, FARHADI A. Yolov3: An incremental improvement[J]. arXiv preprint arXiv:1804.02767, 2018.
- [24] LI C, LI L, JIANG H, et al. YOLOv6: A single-stage object detection framework for industrial applications[J]. arXiv preprint arXiv:2209.02976, 2022.
- [25] WANG C Y, BOCHKOVSKIY A, LIAO H Y M. YOLOv7: Trainable bag-of-freebies sets new state-of-the-art for real-time object detectors[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. [S.l.]: IEEE, 2023: 7464-7475.
- [26] 胡一帆,秦岭,杨小健.改进的基于YOLOv3的人脸检测算法[J]. 数据采集与处理, 2023,38(5): 1092-1103.
HU Yifan, QIN Ling, YANG Xiaojian. Improved face detection algorithm based on YOLOv3[J]. Journal of Data Acquisition and Processing, 2023, 38(5): 1092-1103.
- [27] 贾子豪,王文青,刘光灿.改进YOLOv5的轻量化交通标志检测算法[J]. 数据采集与处理, 2023,38(6): 1434-1444.
JIA Zihao, WANG Wenqing, LIU Guangcan. Improved lightweight traffic sign detection algorithm of YOLOv5[J]. Journal of Data Acquisition and Processing, 2023, 38(6): 1434-1444.
- [28] 李生辉,李晓飞,宋璋晗,等.基于改进YOLOv5的船舶多尺度SAR图像检测算法[J]. 数据采集与处理, 2024,39(1): 120-131.
LI Shenghui, LI Xiaofei, SONG Zhanghan, et al. Multi-scale SAR image detection algorithm for ships based on improved YOLOv5[J]. Journal of Data Acquisition and Processing, 2024, 39(1): 120-131.
- [29] KOU X, LIU S, CHENG K, et al. Development of a YOLO-V3 based model for detecting defects on steel strip surface[J].

Measurement, 2021, 182: 109454.

- [30] XIE Y, HU W, XIE S, et al. Surface defect detection algorithm based on feature-enhanced YOLO[J]. Cognitive Computation, 2023, 15(2): 565-579.
- [31] WANG Y, WANG H, XIN Z. Efficient detection model of steel strip surface defects based on YOLO-V7[J]. IEEE Access, 2022, 10: 133936-133944.
- [32] 姚景丽,程光,万飞,等.改进 YOLOv8 的轻量化轴承缺陷检测算法[J/OL]. 计算机工程与应用. <https://link.cnki.net/urlid/11.2127.TP.20240815.1528.014>.
YAO Jingli, CHENG Guang, WAN Fei, et al. Improved lightweight bearing defect detection algorithm for YOLOv8[J/OL]. Computer Engineering and Applications. <https://link.cnki.net/urlid/11.2127.TP.20240815.1528.014>.
- [33] 冉庆东,郑力新.基于改进 YOLOv5 的锂电池极片缺陷检测方法[J]. 浙江大学学报(工学版), 2024, 58(9): 1811-1821.
RAN Qingdong, ZHENG Lixin. Defect detection method for lithium battery pole piece based on improved YOLOv5[J]. Journal of Zhejiang University (Engineering Edition), 2024, 58(9): 1811-1821.
- [34] HU J, SHEN L, SUN G. Squeeze-and-excitation networks[C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. [S.l.]: IEEE, 2018: 7132-7141.
- [35] WANG C Y, LIAO H Y M, WU Y H, et al. CSPNet: A new backbone that can enhance learning capability of CNN[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops. [S.l.]: IEEE, 2020: 390-391.
- [36] CAI Z, VASCONCELOS N. Cascade R-CNN: Delving into high quality object detection[C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. [S.l.]: IEEE, 2018: 6154-6162.
- [37] XU X, JIANG Y, CHEN W, et al. DAMO-YOLO: A report on real-time object detection design[J]. arXiv preprint arXiv: 2211.15444, 2022.
- [38] CARION N, MASSA F, SYNNAEVE G, et al. End-to-end object detection with transformers[C]//Proceedings of European Conference on Computer Vision. Cham: Springer, 2020: 213-229.
- [39] LV W, XU S, ZHAO Y, et al. DETRs beat YOLOs on real-time object detection[J]. arXiv preprint arXiv:2304.08069, 2023.

作者简介:



董贤光(1990-),通信作者,男,工程师,研究方向:电能计量与传感, E-mail: 1217468383@qq.com。



孙艳玲(1983-),女,副高级工程师,研究方向:电气自动化。



代燕杰(1983-),女,正高级工程师,研究方向:电能计量与传感。



邢宇(1991-),女,工程师,研究方向:电能计量与传感。



翟晓卉(1982-),女,副高级工程师,研究方向:电能计量。



孙凯(1990-),男,工程师,研究方向:电能计量与传感。



吕玉超(1998-),男,博士研究生,研究方向:人工智能、医学影像处理。



吴强(1979-),男,教授,研究方向:人工智能、医学影像处理、生物医学信号处理。



刘琚(1965-),男,教授,研究方向:盲信号处理、通信信号处理、多媒体通信、数字图像处理和DSP应用。

(编辑:夏道家)