

FEV-YOLOv8n: 轻量化安全帽佩戴检测方法

韩博^{1,2,3}, 张婧婧^{1,2,3}, 鲁子翱^{1,2,3}

- (1. 新疆农业大学 计算机与信息工程学院, 乌鲁木齐 830052;
2. 智能农业教育部工程研究中心, 乌鲁木齐 830052;
3. 新疆农业信息化工程技术研究中心, 乌鲁木齐 830052)

摘要: 针对基线 YOLOv8n 检测算法结构较复杂以及现有的安全帽佩戴检测算法参数量和计算量较大, 难以在终端部署等问题, 提出一种基于 FEV-YOLOv8n 的轻量化检测模型; 设计一种轻量级的 FasterC2f 模块改进 YOLOv8n 的骨干网络, 实现网络的参数量和计算量的降低; 在 FasterC2f 模块中引入 EMA 注意力机制, 融合空间依赖和位置信息, 建立长短期的依赖关系, 增强对目标表征的关注, 以提高模型检测的精度; 使用 VoVGSCSP 模块改进颈部网络, 提高遮挡目标以及小目标的辨识度; 实验结果表明, 改进 YOLOv8n 模型 map 值为 92.5%, 相较于 YOLOv8n 算法, 模型大小减少 20%, 计算量降低 18.5%, 参数量降低 15.7%, 为安全帽佩戴检测的轻量化研究提供理论参考。

关键词: 目标检测; 安全帽; FasterC2f; 轻量化; Efficient Multi-Scale Attention; VoVGSCSP

FEV-YOLOv8n: Detection Methods for Wearing Lightweight Safety Helmets

HAN Bo^{1,2,3}, ZHANG Jingjing^{1,2,3}, LU Ziao^{1,2,3}

- (1. College of Computer and Information Engineering, Xinjiang Agricultural University Urumqi 830052, China;
2. Engineering Research Center of Intelligent Agriculture, Ministry of Education Urumqi 830052, China;
3. Xinjiang Agricultural Informatization Engineering Technology Research Center, Urumqi 830052, China)

Abstract: Aiming at the problems of more complicated structure for the baseline YOLOv8n detection algorithm and a large number of parameters and computational complexity for existing helmet wearing detection algorithms, it is difficult to be deployed at terminals, a lightweight detection model based on FEV-YOLOv8n is proposed. A lightweight FasterC2f module is designed to improve the backbone network of YOLOv8n, realizing the reduction of parameters and computation complexity of the network; the EMA attention mechanism is introduced into the FasterC2f module, fuses spatial dependence and positional information, establishes long and short-term dependence relationships, and enhances the attention to the target's representation, so as to improve the accuracy of the model's detection; and the VoVGSCSP is used to improve the neck network, and improve the recognition of occluded targets as well as small targets; Experimental results show that compared with the YOLOv8n algorithm, the improved YOLOv8n model reaches the map value by 92.5%, reduces the size by 20%, the computation by 18.5%, and the parameter quantity by 15.7%, which provides a theoretical reference for the lightweight of safety helmet wearing detection.

Keywords: target detection; safety helmet; FasterC2f; lightweight; efficient multi-scale attention; VoVGSCSP

0 引言

目标检测属于计算机视觉的重要研究领域之一, 其中安全帽佩戴检测在智能监控、智能交通领域中具有广泛的应用。安全帽佩戴检测旨在通过计算机视觉检测图

像或视频中工作人员在工作场所是否佩戴安全帽。由于人体姿态多变、尺寸大小不一、遮挡等因素影响, 安全帽在不同目标检测模型中的鲁棒性仍有待提升。

传统的目标检测方法使用滑动窗口手动提取特征进行分类识别, 主要以 Haar、HOG、Hu moment、SIF、

收稿日期: 2023-11-17; 修回日期: 2024-01-03。

基金项目: 新疆维吾尔自治区自然科学基金资助项目(2022D01A202); 新疆维吾尔自治区高校科研计划项目(XJEDU2020Y020); 新疆维吾尔自治区研究生科研创新项目(XJ2024G124)。

作者简介: 韩博(2001-), 男, 硕士研究生。

通讯作者: 张婧婧(1981-), 女, 硕士, 副教授。

引用格式: 韩博, 张婧婧, 鲁子翱. FEV-YOLOv8n: 轻量化安全帽佩戴检测方法[J]. 计算机测量与控制, 2025, 33(1): 69-77, 84.

SURF 和 DPM 为代表, 其中 Haar 特征是一种基于图像亮度变化的特征描述方法, 它通过在图像上滑动不同大小的矩形滤波器并计算各个区域内像素值的差异来提取特征。Haar 特征的优点在于它们对光照变化具有一定的鲁棒性, 但在复杂场景下的性能可能受到限制。方向梯度直方图 (HOG, histogram of oriented gradients) 是一种用于目标检测的特征描述方法, 它通过计算图像局部区域的梯度方向直方图来表示图像特征。HOG 特征在行人检测等任务中取得了很好的效果, 但对于光照和视角变化较大的情况可能表现不佳。Hu 矩是一组旋转不变的图像特征, 用于形状分析和目标识别。它们基于图像的像素值计算, 并且对图像的旋转、缩放和平移具有不变性。Hu 矩在理论上具有很好的不变性, 但在实际应用中可能对噪声和图像变换敏感。S 尺度不变特征变换 (SIFT, scale-invariant feature transform) 和加速稳健特征 (SURF, speeded-up robust features) 是用于图像特征提取和匹配的经典算法, 对尺度和旋转具有很好的不变性。DPM (Deformable Parts Model) 是一种基于部件的目标检测方法, 它将目标表示为由多个部件组成的结构, 并且可以对目标的变形和姿态进行建模, 在多尺度和姿态变化较大的情况下表现优异。这些方法在过去的几十年中被广泛应用, 但随着深度学习的发展, 许多基于深度学习的目标检测方法已经取代了这些传统方法, 并在许多任务中取得了更好的性能。随着计算机视觉、深度学习等技术的应用发展, 新的目标检测算法被一一提出。单阶段检测算法属于端到端的识别检测, 如: YOLO 系列^[1]、RetinaNet 系列^[3]、SSD 系列^[5]等。目前多数目标检测算法通过改进网络结构提升模型的准确率, 同时导致其复杂度变高, 因此目标检测轻量化的研究逐渐被关注。文献 [7] 使用基于深度可分离卷积的 MobileNetV3 模块作为 YOLOv5s 的骨干网络, 降低骨干网络的参数量, 结合 CBAM 注意力机制, 同时引入知识蒸馏算法, 最终在安全帽检测任务上的平均精度为 89.5%, 模型大小为 8.4 MB。文献 [8] 使用轻量级的 GhostNet 网络作为 YOLOv5 的骨干网络, 降低模型的参数量, 实现模型的轻量化, 添加 CA 注意力机制, 使用一种四尺度融合模块, 提出了 YOLOv5-Q 检测安全帽检测模型, 在安全帽检测任务上的平均精度为 93.7%, 模型大小为 26.47 MB。文献 [9] 提出基于 Ghost 模块重新设计并裁剪的 YOLOv4 的骨干网络, 实现了安全帽佩戴检测模型的轻量化, 颈部网络使用通道加权相加的轻量化 BiFPNs3 模块替换原始 FPN+PAN 的颈部网络, 改进后的 YOLOv4 网络检测的平均精度为 91.1%, GLOPS 为 8.20。

上述方法针对安全帽检测任务都做了不同的轻量化改进, 但是基础模型参数量和复杂度较高, 导致轻量化

后的检测模型仍然存在计算量较大的问题, 不利于终端部署的任务, 对密集型任务检测和遮挡任务性能关注度较小, 针对上述问题, 本文提出的改进 YOLOv8s 目标检测算法, 兼顾了安全帽检测任务的精度和轻量化的要求, 具体表现如下:

1) 针对 YOLOv8n 骨干网络的参数量相对较大, 基于轻量级图像分类网络 FasterNet^[10] 中的 PConv 卷积, 构造出一种 FasterBlock 结构, 使用 FasterBlock 模块替换 C2f 模块中的 BottleNeck 模块, 提出了一种轻量级的 FasterC2f 模块, 同时保留其他位置的普通卷积块, 首先替换骨干网络中的 C2f 模块, 改进后的模型的参数量降低了 11.7%, GFLOPS 降低了 15.7%, 模型的平均精度值降低了 0.2%; 然后替换骨干和颈部网络中 C2f 模块, 改进后的模型的参数量降低了 23.3%, GFLOPS 降低了 22.2%, 模型的平均精度值降低了 0.4%。

2) 针对 YOLOv8n 骨干网络较深层次中丢失信息较严重的问题, 在不增加计算开销的前提下, 在 FasterC2f 模块引入 Efficient Multi-Scale Attention^[11] 注意力机制, 设计了一种 FasterC2fEMA 结构, 保持通道维数的情况下, 通过跨空间学习方法融合空间依赖和位置信息, 增强关注对象的表示, 提高对密集型目标的检测能力。将改进后 FasterC2fEMA 替换骨干网络中的 C2f 模块, 改进后的模型相较于 YOLOv8n 模型, 参数量降低了 11.7%, 模型的平均精度值降低了 0.1%。

3) 针对 YOLOv8n 的颈部网络特征信息融合不充分的问题, 使用一次性聚合的跨级部分网路模块 (VOVGSCSP^[12]) 替换 C2f 模块, 使用残差连接将浅层次的特征图信息与深层次的特征图信息进行合并, 提高模型对远景小目标检测能力。VOVGSCSP 模块分为两种设计模型: 堆叠 1 次、堆叠 3 次, 分别替换颈部 C2f 模块, 堆叠 1 次的 VOVGSCSP 设计模式效果最好, 在使用 FasterC2fEMA 模块替换骨干网络的 C2f 模块的基础上, 使用堆叠 1 次的 VOVGSCSP 替换颈部网络的 C2f, 改进后的模型相较于 YOLOv8n 模型, 参数量降低了 21%, GFLOPS 下降了 18.5%, 模型的平均精度值提升了 0.1%。

1 YOLOv8 模型

YOLOv8 是一种先进的 SOTA (state of the art) 模型, 在 YOLOv5 的基础上进行改进, 无论是在检测精度或速度方面均处于领先地位。YOLOv8 模型由三部分组成: 骨干网络 (Backbone)、颈部网络 (Neck) 和预测头 (Head)。在骨干网络中, YOLOv8 依然采用 CSP 架构, 将特征图分为两部分。第一部分使用卷积运算, 第二部分与第一部分的数据直接连接, 将 YOLOv5 的

C3 模块结合 YOLOv7 的 ELAN 设计了 C2f, 使得模型获得丰富的梯度流信息。在颈部网络中, 加深网络的层次以获取更为丰富的特征信息, 采用特征金字塔网络 (FPN) 和路径聚合网络 (PAN) 结合实现多尺度特征融合 (FPN-PAN), 在密集检测任务中效果较好, 但是较深的网络往往带来较多的卷积操作, 会导致物体的位置信息和特征信息丢失较为严重。在头部结构中, 采用 YOLOX^[13] 的解耦头, 分类头和回归头不再共享参数, 只保留了分类和回归的分支; 同时采用 Anchor-Free^[14] 的检测方法, 通过确定检测物体的中心区域并估计和边界框的距离。在损失函数的计算中, 采用了 TaskAlignedAssigner^[15] 正负样本分配策略, 使用了 DFL (Distribution Focal Loss) 交叉熵损失函数。DFL 交叉熵损失函数的计算方法如公式 (1) 所示:

$$DFL(S_i, S_{i+1}) = -((y_{i+1} - y) \log(S_i) + (y - y_i) \log(S_{i+1})) \quad (1)$$

其中: S_i, S_{i+1} 表示 y_i, y_{i+1} 的一般分布。

2 轻量化方法和注意力机制的相关理论

2.1 PConv 轻量化卷积块

FasterNet 卷积神经网络提出了一种新的 Partial 卷积块 (PConv), 通过同时减少冗余计算和内存的访问, 更为有效地提取空间特征。PConv 卷积将一套规律的特征提取方案应用于空间特征提取的部分输入通道, 对于其余通道不做处理, 同时保持输入和输出特征图的通道数量相同; 而标准卷积则是卷积核通道数与输入特征图通道数保持一致的特征提取方式。下文引入 PConv 卷积与标准卷积性能的对比, 其结构如图 1 所示。

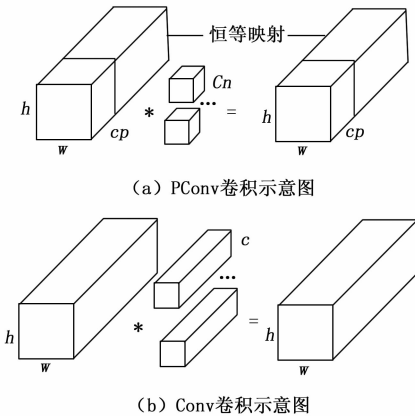


图 1 PConv 与 Conv 卷积的结构对比

对于一个输入 $Input \in R^{c \times h \times w}$, PConv 使用 c_n 个滤波器, 普通卷积 (Conv) 使用 c 个滤波器, 一个 PConv 的内存使用量如下:

$$Pconv_{mem} = h \times w \times 2c_n + k^2 \times c_n^2 \approx h \times w \times 2c_n \quad (2)$$

FLOPs 为:

$$F_{Pconv} = h \times w \times k^2 \times c_n^2 \quad (3)$$

一个 Conv 卷积的内存使用量如下:

$$Conv_{mem} = h \times w \times 2c + k^2 \times c^2 \approx h \times w \times 2c \quad (4)$$

FLOPs 为:

$$F_{conv} = h \times w \times k^2 \times c^2 \quad (5)$$

由上述公式 (2) ~ (5) 可得:

$$\frac{Pconv_{mem}}{Conv_{mem}} = \frac{c_n}{c} \quad (6)$$

$$\frac{F_{Pconv}}{F_{conv}} = \left(\frac{c_n}{c}\right)^2 \quad (7)$$

采用 Pconv 替换标准卷积, 仅对 c_n 个通道做卷积, 并不意味着抛弃其余通道, 而是通过使用恒等映射由后续的标准卷积进行特征提取, 因此减少了内存的使用量, 降低浮点数的运算, 从而在一定程度上实现模型的轻量化, 同时 Pconv 通过使用恒等映射保留原始的特征图信息, 结合标准卷积提取特征, 对特征图的信息进行有效提取。

2.2 EMA 注意力机制

EMA (Efficient Multi-Scale Attention) 注意力机制可以降低卷积运算中的通道维数, 同时生成像素级注意力的高级特征图。EMA 模块沿用 CoordinateAttention^[16] 注意力机制的设计策略, 将位置信息嵌入到通道注意力中, 将通道注意力分解为沿两个不同方向聚合特征的一维特征编码过程, 分别沿水平方向和垂直方向做一维全局平均池化, 不同于通道注意力将输入使用二维的全局池化转化为单个特征向量。EMA 注意力机制将通道注意力分解为两个一维向量的特征编码, 垂直方向捕获长距离的依赖性, 水平方向保留精确的位置信息, 同时增加一个卷积核大小为 3×3 的并行分支, 聚合多尺度的空间结构信息, 它们可以互补地应用于输入的特征图, 有效地建立长短期地依赖关系, 增强对感兴趣物体的表征地关注, EMA 注意力机制模块结构如图 2 所示。

具体地说, 主要分为以下 3 部分:

1) 特征分组, 给定尺寸为 $Input \in R^{C \times H \times W}$, EMA 将 $Input$ 跨通道维度的方向分为 G 个子特征, 学习不同子特征中的感兴趣区域, 加强特征表示。

2) 并行子网, EMA 模块采用 3 个并行路线来提取分组特征图, 在信息编码阶段, 第一个分支为垂直坐标方向使用尺寸大小 $(1, W)$ 的池化核对每个通道进行编码, 具体计算如公式 (8) 所示:

$$z_c^w = \frac{1}{H} \sum_{0 \leq j < H} x_c(j, w) \quad (8)$$

其中: z_c^w 表示宽度为 w 的第 c 个通道的输出, 仅对垂直方向做处理, 水平方向则不做变化。

第二个分支为水平方向使用尺寸大消息为 $(H, 1)$ 的池化核对每个通道进行编码, 具体计算如公式 (9)

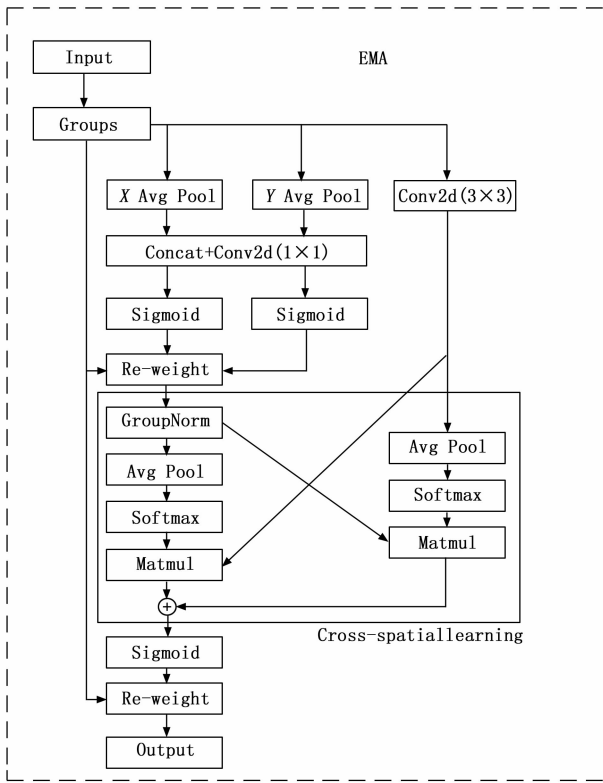


图 2 EMA 注意力机制

所示：

$$z_c^h = \frac{1}{H} \sum_{0 \leq j < H} x_c(h, i) \quad (9)$$

其中： z_c^h 表示高度为 h 的第 c 个通道的输出，仅对水平方向做处理，垂直方向则不做变化。

将垂直方向和水平方向的特征图进行融合，并使用卷积核大小为 1×1 的卷积进行跨通道信息交流，具体计算如公式 (10) 如下：

$$f = P([z^h, z^w]) \quad (10)$$

f 表示垂直方向和水平方向融合后的特征图，其中 $f \in R^{C//G \times (H+W)}$ ， G 表示划分比例， P 表示卷积核大小为 1×1 的卷积，用于变换通道数。

进一步地沿着通道的维度将融合后的特征图 f 分解成 2 个独立的张量，垂直方向为 $f^h \in R^{C//G \times H \times 1}$ ，水平方向为 $f^w \in R^{C//G \times 1 \times W}$ ，然后使用 sigmoid 激活函数，计算过程如公式 (11) 和 (12) 所示：

$$g^h = \sigma(f^h) \quad (11)$$

$$g^w = \sigma(f^w) \quad (12)$$

其中： σ 表示 sigmoid 激活函数，将 g^w 和 g^h 的值作为 EMA 注意力的中间权重。

第三个分支是一个卷积核大小为 3×3 的普通卷积。

3) 跨空间学习阶段，使用 2 维的全局平均池化作用于处理后的水平方向和垂直方向的特征图信息，在特征通道的联合激活之前变换最小分支的输出的维度与之

相匹配。2 维全局池化的操作如公式 (13) 所示，输入尺寸是 $H \times W \times C$ ，输出尺寸为 $1 \times 1 \times C$ 的过程：

$$z_c = \frac{1}{H \times W} \sum_{j=1}^W \sum_{i=1}^H x_c(i, j) \quad (13)$$

其中： z_c 表述第 c 个通道相关的输出， H 和 W 分别表示图像的高度和宽度， $x_c(i, j)$ 表示高度是 i 、宽度为 j 的第 c 个通道的输入。在 2 维全局平均池化后采用了 Softmax 的非线性激活函数适应线性变换，通过将并行处理的输出与 Softmax 的输出相乘，得到一张空间注意力图，融合了同一阶段不同尺度的空间信息。同样的，对卷积核大小为 3×3 的普通卷积所提取使用 2 维全局平均池化和 Softmax 的非线性激活函数，生成另外一张空间注意力图，精确地记录空间位置信息，将生成的空间注意力权重值的信息进行聚合，最后使用 Sigmoid 非线性激活函数。

2.3 VoVGSCSP 细颈型模块

VoVGSCSP 模块是由 GSBottleneck 模块为构成，GSBottleneck 模块由 GSConv 基础模块构成，GSConv 以较低的时间复杂度，最大限度地保留通道间的联系，冗余重复的特征图信息较少。GSConv 模块由标准卷积 (Conv)，深度可分离卷积 (DWConv) 和特征图洗牌 (shuffle) 组成，DWConv 和 shuffle 操作增强非线性的表达能力。GSBottleneck 模块在 GSConv 的基础上，结合了残差思想。VoVGSCSP 模块的设计基于 GSBottleneck 模块，使用两条分支进行特征的提取，聚合分支的特征图，对聚合后的多通道特征图使用卷积提取信息，获得更加丰富的特征信息，结合类残差的跨阶段操作，同时提升颈部网络的非线性表达能力。GSConv、GSBottleneck、VoVGSCSP 模块结构如图 3 所示。

3 YOLOv8 的改进

3.1 融合 EMA 注意力机制的轻量化骨干网络

3.1.1 FasterC2f 模块

YOLOv8n 中的 C2f 模块中的 Bottleneck 是由两个连续的 CBS 模块组成，CBS 模块则由普通卷积、批处理归一化、SiLU 激活函数 3 部分构成，CBS 存在参数量较大，计算复杂度较高的问题，本文提出一种低复杂度、轻量级的 FasterBlock 模块，FasterBlock 模块是由一个 PConv 模块、两个卷积核大小为 1×1 的 Conv 组成，FasterBlock 模块使用 MobileNetV2 种的倒残差思想，扩展中间层的通道数量，中间层使用归一化层，将 ReLU 激活函数替换为 SiLU^[17] 激活函数，FasterBlock 模块通过减少卷积操作的数量实现网络参数量的大幅度减少，加快推理速度，基于 FasterBlock 模块构造出 FasterC2f 模块，将原始的 C2f 模块中的 Bottleneck 替换为 FasterBlock 模块，改进的 FasterBlock 网络结构图如图 4。

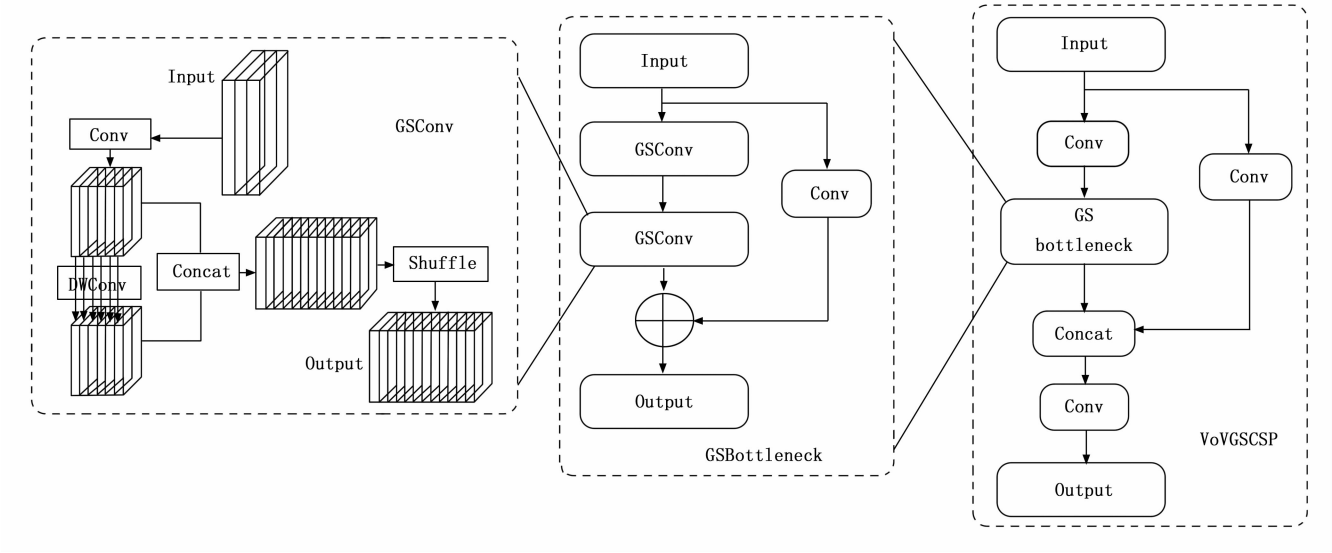


图 3 VoVGSCSP 模块

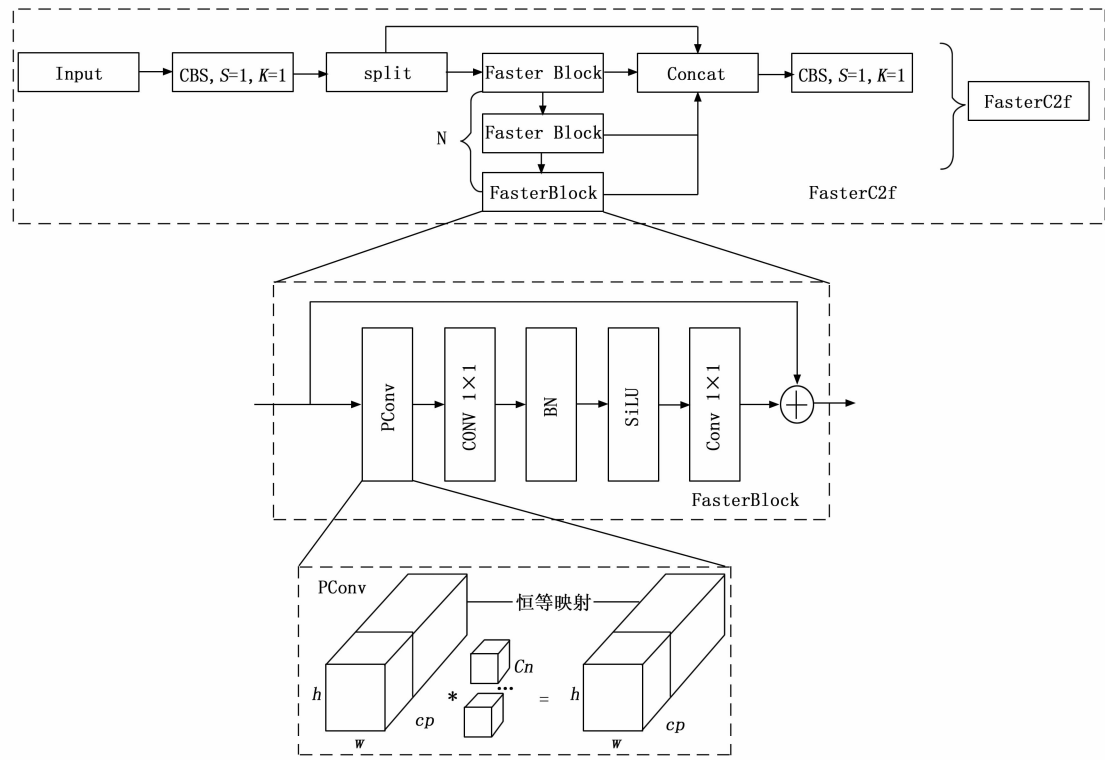


图 4 基于 FasterBlock 的 FasterC2f 模块

3.1.2 FasterC2fEMA 模块

在真实的工地、隧道等施工场景下，环境背景较为复杂，在隧道内往往环境较暗，在工地上往往背景较亮，使得安全帽佩戴的检测容易受到复杂背景的干扰，从而导致检测任务出现错检、漏检等情况，因此通过引入注意力机制，弱化复杂背景的干扰，使得检测模型更加关注是否佩戴安全帽检测这项任务。将 EMA 注意力机制增加在 FasterC2f 模块的尾部，构造出 Fast-

erC2fEMA 模块，通过跨空间的学习生成空间注意力权重值，融合多尺度空间信息和位置信息。EMA 注意力机制可以捕获特征图的像素级成对关系并突出显示像素的全局上下文信息。在不增加参数量和计算量的情况下，实现特征图信息的充分融合，FasterC2fEMA 模块结构如图 5 所示。

将 YOLOv8n 中 Backbone 的 C2f 模块替换为 FasterC2fEMA 模块，减少骨干网络的参数量，实现了模型

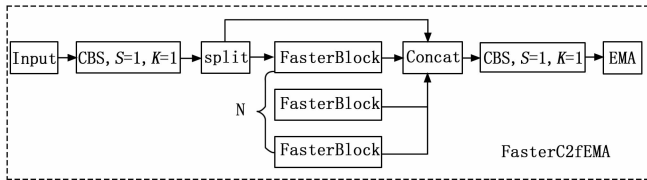


图 5 嵌入 EMA 注意力机制的 FasterC2f 模块

轻量化的目标,改进后的骨干网络如图 6 所示。

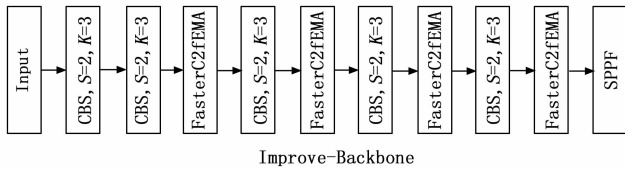


图 6 改进后的 YOLOv8 骨干网络

3.2 引入 VoVGSCSP 模块的颈部网络

针对安全帽检测任务中存在远景小目标难以检测的问题,在颈部网络中引入一次聚合法跨阶段局部网络模块 (VoVGSCSP) 充分融合骨干网络提取的特征图信息,将 VoVGSCSP 模块替换 YOLOv8n 中的颈部模块中的 C2f 模块,实现骨干网络提取特征图信息的充分融合,增强对小目标的检测能力,改进后的颈部的网络结构图如图 7 所示。

基于 YOLOv8n 模型,采用上述的改进方法,使用 FasterC2fEMA 模块替换骨干网络中 C2f 模块,在颈

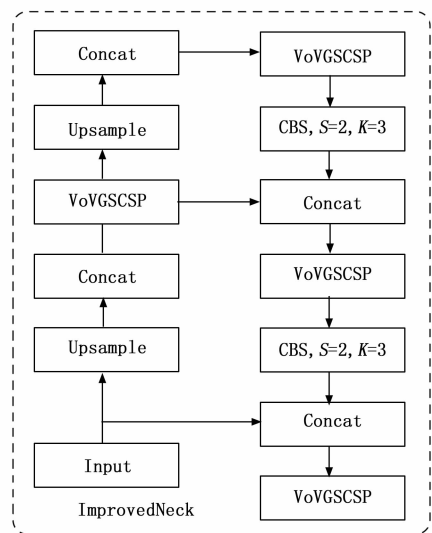


图 7 基于 VoVGSCSP 的改进颈部模块

部网络中使用堆叠次数为 1 的 VOVGSCSP 模块替换 C2f 模块,提出了轻量化的安全帽检测算法 FEV-YOLOv8n (FasterC2f-EMA-VOVGSCSP-YOLOv8n),结构如图 8 所示。

4 实验结果分析

4.1 实验环境介绍

本次实验的环境配置如表 1 所示。

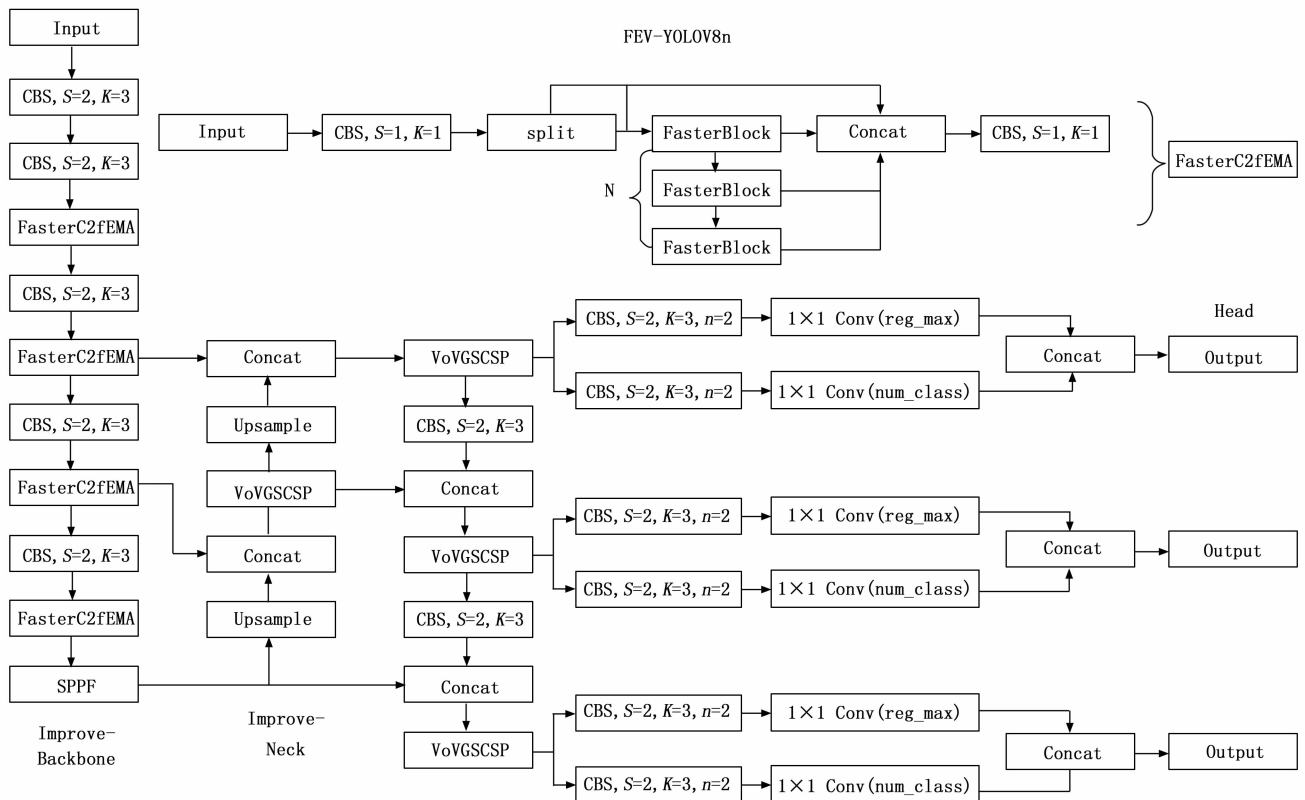


图 8 改进 YOLOv8n 结构

表 1 实验环境配置	
名称	环境配置
操作系统	CentOS 7.9.2009
显卡	NVIDIA GeForce RTX 3090
深度学习环境	Pytorch 1.11.0
集成开发环境	Anaconda, CUDA 11.3, CUDNN 8.6.0

4.2 数据集与参数配置

文中使用的数据集是基于公开的安全帽佩戴检测数据集: SHWD (Safety Helmet Wearing^[18]), 并对数据集进行预处理; 最终数据集一共包含是 7851 张图片和 TXT 标签文件。该数据集样本分为两大类: “hat” 表示佩戴安全帽, “person” 表示没有佩戴安全帽, 将安全帽数据集按照训练集 (70%)、验证集 (20%) 和测试集 (10%) 的比例划分。

- 模型训练时参数设置为:
- 1) 优化器: Adam
 - 2) 图像的分辨率: 640 * 640
 - 3) Batchsize 设置: 32
 - 4) 设置的初始学习率: 0.001
 - 5) 最终学习率: 0.000 03

后 10 轮关闭 Mosaic 图像数据增强算法, 模型改进后也按照此训练参数进行训练。

4.3 评价指标

训练好的模型需要采用评估指标进行评估、检验模型的准确性, 本文实验所使用的评估指标为: 召回率 (Recall)、准确率 (Precision)、AP (average precision) 和 mAP (mean average precison) 4 种。准确率计算方法, 如公式 (14) 所示:

$$Precision = \frac{TP}{TP + FP}$$
 (14)

召回率计算方法, 如公式 (15) 所示:

$$Recall = \frac{TP}{TP + FN}$$
 (15)

其中: TP 表示正确判定为正样本的数量, FP 表示错误判定为正样本的数量, FN 表示正确判断为负样本的数量。

AP 计算方法, 如公式 (16) 如下:

$$AP = \int_0^1 PR \mathrm{d}r$$
 (16)

AP 是计算某个类别的 P-R 曲线下的面积。
mAP 的计算方法, 如公式 (17) 所示:

$$mAP = \frac{1}{n} \sum_{i=0}^n AP_i$$
 (17)

mAP 是 AP 值所有类别下的均值。

4.4 轻量化卷积块的对比试验

为证明 PConv 卷积的轻量和高效, 设置对比试验, 在

相同的实验条件下, 分别对 Conv、Depth-Conv、GSConv、PConv、DCNV2 五种卷积在运行时间、FPS、参数量等评价指标上进行性能测试, 实验结果如表 2 所示。

表 2 多种卷积的性能对比情况

名称	总时间/s	平均时间/s	FPS	FLOPs	参数量
Conv	22.447	0.004 49	223	77.58 G	147.840 K
Depth-Conv	22.612	0.004 52	221	9.731 G	18.304 K
GSConv	24.796	0.004 96	202	39.762 G	75.712 K
DCNV2	98.214	0.019 64	51	16.576 G	31.387 K
PConv	11.438	0.002 29	437	5.100 G	9.472 K

从表 2 可以看出, PConv 卷积的总运行时间为 11.438 s 等, 表明在上述 5 种卷积中, PConv 实时性能最好, 参数量和计算量最低, 说明基于 PConv 卷积的 FasterC2f 适用于 YOLOV8n 的骨干网络的轻量化。

4.5 改进过程的消融实验及分析

YOLOv8 网络根据网络宽度和深度的不同, 分为 n、s、m、l、x 五种不同网络结构模型, 本文以宽度和深度最小的 YOLOv8n 网络为基线, 对模块 FasterC2f、FasterC2fEMA 和 VoVGSCSP 模块的单独及组合改进任务上进行消融实验, 证明本文提出的改进方法的在安全帽检测任务上的有效性, 具体实验结果如表 3 所示。

表 3 改进过程的消融实验

算法	mAP@0.5/%	params/10 ⁶	GFLOPS	Size/MB
YOLOv8n	92.4	3.00	8.1	6
YOLOv8n-FasterC2f(Backbone)	92.2	2.65	7.0	5.5
YOLOv8n-FasterC2f(Backbone+Neck)	92.0	2.30	6.3	4.6
YOLOv8n-FasterC2fEMA(Backbone)	92.3	2.65	7.6	4.6
YOLOv8n-FasterC2f-VoVGSCSP(repeats=3)	92.0	2.53	6.3	5.1
YOLOv8-FasterC2fEMA-Backbone-VoVGSCSP(repeats=3)	92.1	2.53	6.8	5.1
YOLOv8-FasterC2fEMA-Backbone-VoVGSCSP(repeats=1)	92.5	2.37	6.6	4.8

从表 3 中可以发现, 使用 FasterC2f 替换 YOLOv8n 模型的 C2f 模块后, 参数量和 GFLOPS 显著下降, 实现了轻量化的目标, 但是模型的精度也出现降低的情况, 分析后发现由于标准卷积数量的减少, 导致特征丢失的情况, 针对特征丢失的情况, 仅将 FasterC2f 应用于骨干网络中, 同时在 FasterC2f 模块引入了 EMA 注意力机制去替换骨干网络的 C2f 模块, 模型的平均精度上升 0.1%; 针对小目标检测任务, 在颈部网络中使用 VoVGSCSP 模块替换颈部网络中的 C2f 模块, VoVGSCSP 模块重的堆叠层数为 1 时精度最高, 模型的参数

量最小, 仅为 2.37 MB。本文提出的 FEV-YOLOv8n 模型相较于 YOLOv8n 模型, 精度略微上升, 模型的参数量降低了 21%, 模型的计算量降低了 18.5%, 模型大小减小了 20%, 实现了安全帽检测模型的轻量化。

FEV-YOLOv8n 与 YOLOv8n 在安全帽数据集上的平均精度如图 9 所示。

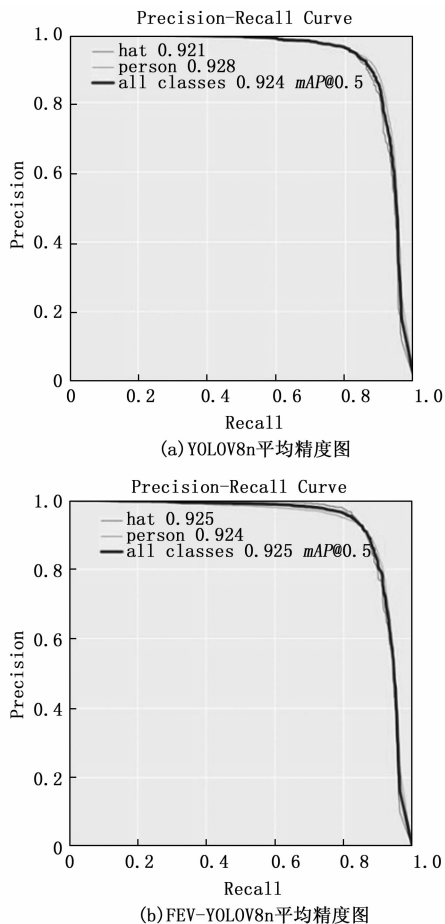


图 9 FEV-YOLOv8n 和 YOLOv8n 的平均精度图

4.6 检测效果对比分析

4.6.1 目标遮挡及远景小目标检测

为了更加直观地呈现改进模型与基线 YOLOv8n 检测效果的区别, 选取了遮挡类型的目标和远景小目标的图片进行检测效果的对比, 结果如图 10、图 11 所示。

通过观察检测效果, 可以发现 FEV-YOLOv8n 在带有遮挡的安全帽的检测任务上效果好于 YOLOv8n, 检测结果如图 10 所示。

在远景小目标检测上, YOLOv8n 出现错误检测的情况, FEV-YOLOv8n 能够正确识别, 检测结果如图 11 所示。

4.6.2 密集型目标检测

在密集型目标检测任务中, 由于图像分辨率的原因, FEV-YOLOv8n 和 YOLOv8n 都存在不同程度的漏

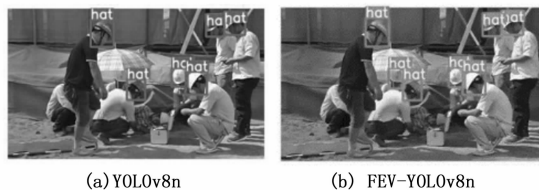


图 10 遮挡类型检测效果对比

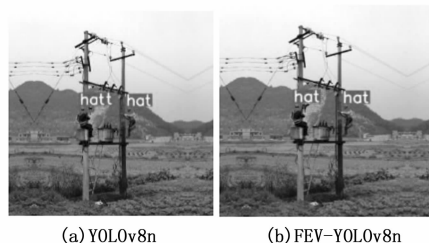


图 11 远景小目标类型检测效果对比

检情况。但是 YOLOv8n 出现错检、漏检率高于 FEV-YOLOv8n, 检测结果如图 12 所示。

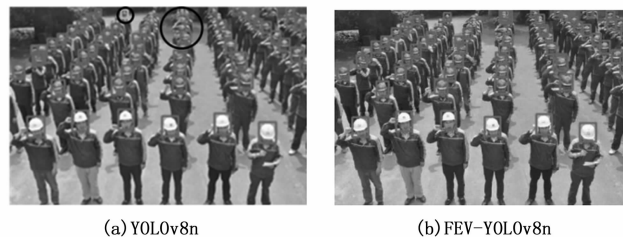


图 12 密集型目标检测效果对比

由不同安全帽检测场景下的 FEV-YOLOv8n 和 YOLOv8 的检测效果对比可知, 改进后的 FEV-YOLOv8n 在降低参数量、计算量的前提下, 实现轻量化目标的前提, 保持相对较高的精度, 同时适应不同的复杂场景, 实现预期检测效果。

4.7 FEV-YOLOv8n 与其他算法的比较

为了验证 FEV-YOLOv8n 算法的在轻量化研究方面的优越性和可行性, 与目前主流在安全帽检测任务上的 YOLO 系列算法进行结果对比, 进一步 FEV-YOLOv8n 的性能, 实验结果如表 4 所示。

表 4 FEV-YOLOv8n 与其他算法对比实验结果

算法	mAP@0.5/%	params/ 10^6	size/MB
YOLOv5-Q ^[8]	93.7	12.69	26.47
YOLO-M3 ^[9]	89.5	4.20	8.4
ShuffleNetV2-YOLOv5s ^[9]	63.5	1.40	2.8
MS-YOLO ^[19]	88.41	11.59	47.6
Improve-YOLOX ^[19]	89.8	10.32	41.6
YOLOv8n	92.4	3.00	8.1
FEV-YOLOv8n	92.5	2.37	4.8

根据表 4 的实验结果可知, 与 YOLO-M3 相比, FEV-YOLOv8n 模型的参数量减少, 平均精度提升了 3

个百分点; 与 MS-YOLO 模型和 Improve-YOLOX 模型相比, FEV-YOLOv8n 模型的大小和参数量大幅度地减小; 与 ShuffleNetV2-YOLOv5s 相比, FEV-YOLOv8n 虽然参数量和模型大小稍大, 但是平均精度比 ShuffleNetV2-YOLOv5s 高了 29 个百分点。FEV-YOLOv8n 相较于以往的轻量化模型和检测算法, 具有更为出色的检测性能, 参数量和模型的大小均满足终端部署的需求, 同时保持较高水平的平均精度。为了更加直观本文的 FEV-YOLOv8n 与其他算法相比较的优越性, 在 $mAP@0.5$ 、params 和 size 三个模型性能评价指标上进行可视化对比分析, 如图 13 所示。

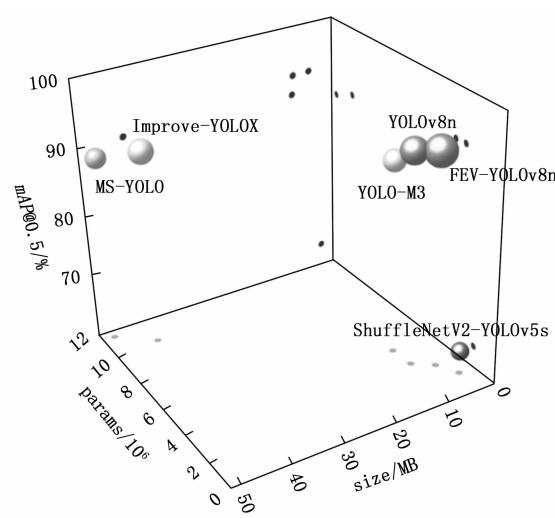


图 13 FEV-YOLOv8n 与其他轻量化算法的对比

5 结束语

在以往的安全帽佩戴检测模型中, 部分模型为保持较高的平均精度, 采用参数量大、结构复杂的网络结构, 难以在嵌入式终端设备中部署; 另一部分模型为追求极致的轻量化, 降低其平均精度, 在安全帽佩戴检测任务中表现不佳, 也难以满足安全帽佩戴检测任务的需求。本文提出一种 FasterC2f 模块的轻量化检测模型 YOLOv8n 的骨干网络, 同时在 FasterC2f 模块中引入 EMA 注意力机制, 通过增强对目标表征的关注, 提高了检测的检测精度。此外, 颈部网络中还引入 VoVG-SCSP 模块替换 C2f 模块, 提高对远景小目标任务和密集任务的检测能力。改进的 FEV-YOLOv8n 模型具备容量小、精度高、泛化能力强的优势, 易于在终端部署, 满足白天一般光照场景下安全帽佩戴检测的需求。但在夜晚或光线不足的情况下, 其模型检测效果还有待进一步研究。

参考文献:

[1] LIU Y, CHU H, SONG L, et al. An improved Tuna-YOLO model based on YOLO v3 for real-time tuna detection

considering lightweight deployment [J]. Journal of Marine Science and Engineering, 2023, 11 (3): 542.

[2] GUO Y, CHEN S, ZHAN R, et al. LMSD-YOLO: A lightweight YOLO algorithm for multi-scale SAR ship detection [J]. Remote Sensing, 2022, 14 (19): 4801.

[3] HUANG L, WANG Z, FU X. Pedestrian detection using RetinaNet with multi-branch structure and double pooling attention mechanism [J]. Multimedia Tools and Applications, 2023; 1-25.

[4] CHEN Y, ZHENG L, PENG H. Assessing pineapple maturity in complex scenarios using an improved retinanet algorithm [J]. Engenharia Agricola, 2023, 43: e20220180.

[5] ZHANG X, ZHANG Y, GAO T, et al. A novel SSD-Based detection algorithm suitable for small object [J]. IEEE TRANSACTIONS on Information and Systems, 2023, 106 (5): 625-634.

[6] JIANG T. Application of SSD network algorithm in panoramic video image vehicle detection system [J]. Open Computer Science, 2023, 13 (1): 20220270.

[7] 杨永波, 李 栋. 改进 YOLOv5 的轻量级安全帽佩戴检测算法 [J]. 计算机工程与应用, 2022, 58 (9): 201-207.

[8] 刘雪纯, 刘大铭, 刘若晨. 改进的轻量级安全帽佩戴检测算法 [J]. 液晶与显示, 2023, 38 (7): 964-974.

[9] 胡文骏, 杨莉琼, 肖宇峰, 等. 识别安全帽佩戴的轻量化网络模型 [J]. 计算机工程与应用, 2023, 59 (13): 149-155.

[10] CHEN J, KAO S, HE H, et al. Run, Don't walk: chasing higher FLOPS for faster neural networks [C] //Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, 2023: 12021-12031.

[11] OUYANG D, HE S, ZHANG G, et al. Efficient multi-scale attention module with cross-spatial learning [C] //ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP). IEEE, 2023: 1-5.

[12] YU M, WAN Q, TIAN S, et al. Equipment identification and localization method based on improved YOLOv5s model for production line [J]. Sensors, 2022, 22 (24): 10011.

[13] ZHANG Y, ZHANG W, YU J, et al. Complete and accurate holly fruits counting using YOLOX object detection [J]. Computers and Electronics in Agriculture, 2022, 198: 107062.

[14] CHENG G, WANG J, LI K, et al. Anchor-free oriented proposal generator for object detection [J]. IEEE Transactions on Geoscience and Remote Sensing, 2022, 60: 1-11.

(下转第 84 页)