

基于机器学习方法的员工安全帽佩戴检测

张硕羲¹ 任佳亮² 陈 峰¹ 曹存盼² 易 杨¹

ZHANG Shuoxi REN Jialiang CHEN Feng CAO Cunpan YI Yang

摘 要

针对当前化工行业中员工安全帽佩戴检测方法效率低、实时性差等问题,为减少员工由于未佩戴安全帽进入生产装置而发生意外情况,提高安全生产效率,文章提出了一种基于改进 YOLOv5 算法的员工安全帽佩戴检测方法。首先,采用多尺度 Retinex (MSR) 图像处理算法对生产现场视频图像进行预处理,提高视频图像的信噪比,降低噪声干扰;其次,结合 SKNet 动态注意力机制对原 YOLOv5 算法模型进行改进,提高视频特征信息提取能力,优化算法的检测精度;最后,对改进的 Im-YOLOv5 模型的检测准确率进行实验验证。实验结果表明,改进的 Im-YOLOv5 模型的 mAP (mean average precision) 达到了 93.5%,其准确率相较于 YOLOv5 得到了提升,具有较好的检测效果。

关键词

Retinex (MSR); 预处理; 信噪比; 动态注意力机制; Im-YOLOv5

doi: 10.3969/j.issn.1672-9528.2025.01.023

0 引言

近年来,化工行业安全事故的发生率备受关注。相关报道表明,化工安全事故的发生多因工作人员安全意识不足,未及时佩戴安全帽导致^[1]。如沧州炼油厂员工给原油换热器更换垫片过程中,由于照明灯突然自灭从而造成从脚手架坠落,因其未佩戴安全帽,造成坠落后重度脑挫裂伤,抢救无效死亡。

而在实际的化工生产中,由于作业人员流动性大,负责安全监管人员监督不到位等因素的影响,作业人员未及时佩戴安全帽的行为时常发生,带来极大安全隐患。因此,研究高效的员工安全帽佩戴检测方法具有重要的意义。丁文龙等人^[2]提出一种基于 YOLOv3 的安全帽检测方法,该方法采用了 K-means++ 聚类算法,对先验框的尺寸进行优化,并且引入注意力机制对特征提取网络进行改进,实验结果表明,该方法的平均准确率达到 88.16%;李帅等人^[3]提出一种基于改进 YOLOv4 算法的安全帽检测方法。该方法引入 K-means 聚类优化锚框尺寸,同时采用空洞卷积避免模型过拟合,提升模型对中小物体的检测效率,实验结果表明该方法检测效果较好,但实时性不高;Hayat 等人^[4]提出了一种基于 YOLO 算法的建筑工地安全帽自动检测方法,实验准确率达到 92.44%,因此在检测安全帽方面显示出优异的结果。本文基于图像处理及目标检测算法,提出一种改进的 Im-YOLOv5 员工安全帽佩戴检测方法,降低化工厂意外事件发生率。

1 Retinex 方法介绍

由于化工生产环境复杂,往往受天气、粉尘及灯光等因素的影响,造成监控视频图像失真,使得图像的特征信息丢失。为了增加视频图像的特征信息,提高员工安全帽佩戴检测的有效性,首先采用 Retinex 算法对视频监控的图像进行预处理。

Retinex 方法^[5]是基于视网膜皮层理论被提出的一种图像处理算法,该方法采用的处理思想即视野中的原始图像,其像素值的变化范围是由外界的光照强度所决定的,而原始图像所具有的反射系数则决定了其固有属性,如图 1 所示。因此 Retinex 算法的核心处理步骤就是去除外界光照对物体自身的影响,从而保留物体的固有属性,如图 2 所示。

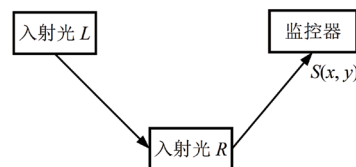


图 1 Retinex 结构图

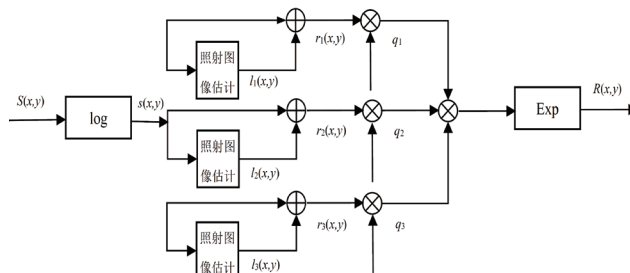


图 2 Retinex 算法流程图

1. 中油(长汀)催化剂有限公司 福建龙岩 361000

2. 中国石油兰州石化公司 甘肃兰州 730000

具体步骤为:

$$S(x, y) = L(x, y) \cdot R(x, y) \quad (1)$$

$$r(x, y) = \sum_{k=1}^K q_k \{\log S(x, y) - \log [G_k(x, y) * S(x, y)]\} \quad (2)$$

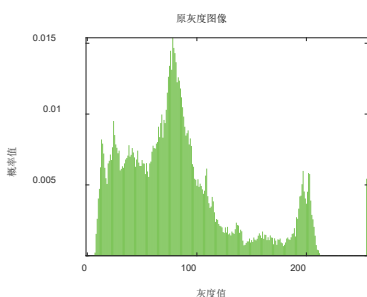
$$G_k(x, y) = \lambda_k \exp \left[-\frac{x^2 + y^2}{2\sigma_k^2} \right] \quad (3)$$

式中: $G_k(x, y)$ 为高斯函数; K 表示尺度的个数; q_k 代表每一个尺度的权值。

图 3 为处理前的夜间视频监控图像, 由图像可以看出, 由于光线不足造成图像色彩失真, 且像素分布不均匀。图 4 为处理后视频监控图像, 可以看出, 处理后的图像色彩得到提升, 细节信息更加完整且像素分布更加均匀, 提高了特征信息提取效果。



(a) 夜间视频监控图像

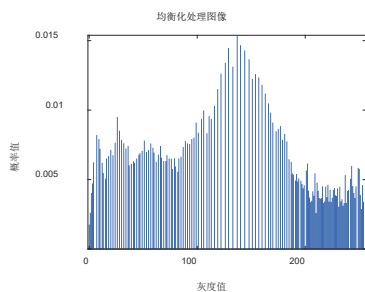


(b) 像素值分布

图 3 处理前图像



(a) 处理后视频图像



(b) 像素值分布

图 4 处理后视频监控图像

2 算法介绍

2.1 动态注意力机制

针对 YOLOv5 算法用于员工安全帽佩戴检测的准确率, 本文结合 SKNet 动态注意力机制对 YOLOv5 算法模型进行改进, 进一步提高安全帽佩戴的检测精度。在原 YOLOv5 算法模型中, 采用标准卷积神经网络进行特征提取, 由于标准卷积神经网络中不同层的人工神经元感受野的尺寸是相同的, 所以在接收到不同的刺激信号时, 感受野的大小无法做出相应的调节。为了弥补标准卷积的缺陷, 原 YOLOv5 模型中添加动态选择的机制, 使得神经元在接收到不同尺度的信息时, 感受野能够自动调整为相匹配的尺寸。

动态选择注意力机制的关键是选择单元核 (SK), 选择单元核将多个不同尺寸核的通道融合起来, 将多个选择单元核堆叠在一起就形成了选择注意力机制 (SKNet)。该机制能够自适应调节感受野的尺寸并对不同大小的目标进行捕获, 具有很好的性能。SKNet 结构图如图 5 所示。

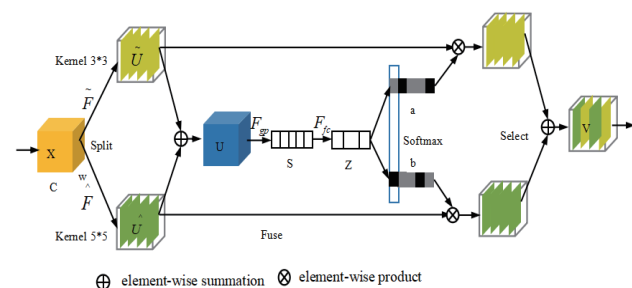


图 5 SKNet 结构图

2.2 YOLOv5 算法

YOLOv5 算法是一种从端到端的检测过程, 其核心是卷积神经网络的特征提取结构, 可识别输入目标类别及输出位置, 是图像分类与定位相结合的算法^[6-8]。YOLOv5 算法模型结构主要包括 Backbone (包括 Focus 模块、CBL 模块、CSP 模块以及 SPP 模块)、Neck (包括 CSP 模块、CBL 模块以及 CONV 模块) 及 Prediction。火焰目标检测过程如下:

(1) 将视频图像划分为 $N \times N$ 个单元格, 每个单元格针对大中小不同尺度的目标生成先验框, 识别目标的中心落在某个网格中, 则由该网格的先验框负责跟踪识别该目标。

$$c = P \times H \quad (4)$$

式中: c 为置信度, 表示该边框中目标的分类概率及匹配目标的性能; P 为预测框内的目标概率, 当边界框内无目标出现时, P 值取 0, 相反取 1; H 为预测框与真实框的交并比。

(2) 对划分的视频图像归一化处理, 送入下层特征提取网络, 对视频图像目标进行特征提取。

(3) 通过聚类设置标准框, 分为不同大小的框, 针对不同检测目标, 计算标准框的位置, 即中心点坐标。

(4) 依据预测坐标的偏移值，计算目标的中心点位置及标注框宽高。

(5) 输出识别分类结果。

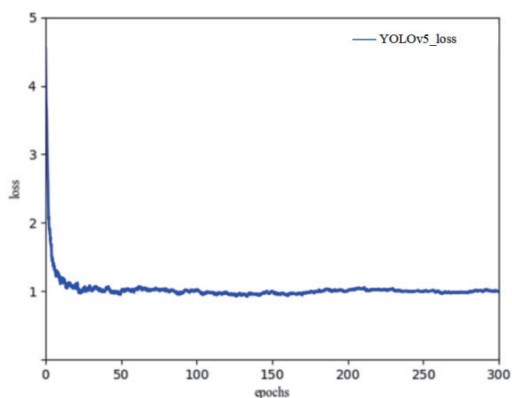
3 实验结果

为进一步验证本文提出的 Im-YOLOv5 算法模型对员工安全帽佩戴检测的可行性，该节分别对 YOLOv5 及本文提出的 Im-YOLOv5 进行实验验证，实验中采用的数据图片为某生产装置现场作业人员图像，具体实验环境配置参数如表 1 所示。

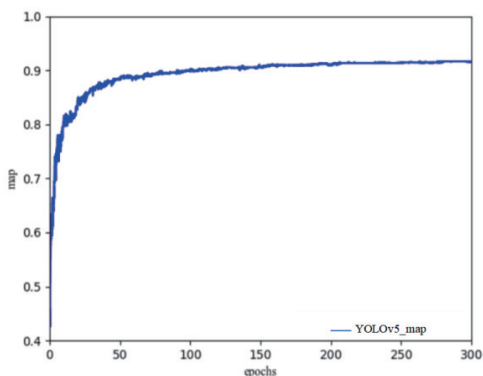
表 1 环境配置参数

参数设置	设定值
迭代次数	300
动量因子	0.7
学习率	0.002
训练批次	40

由图 6 (a) 可以得到，YOLOv5 模型初始损失值约为 4.5，训练过程中随着迭代次数的增加，在迭代次数至 40 轮左右达到平稳，平稳值为 0.96 左右。从图 6 (b) 中可以得到，YOLOv5 模型的 mAP 值为 91% 左右。

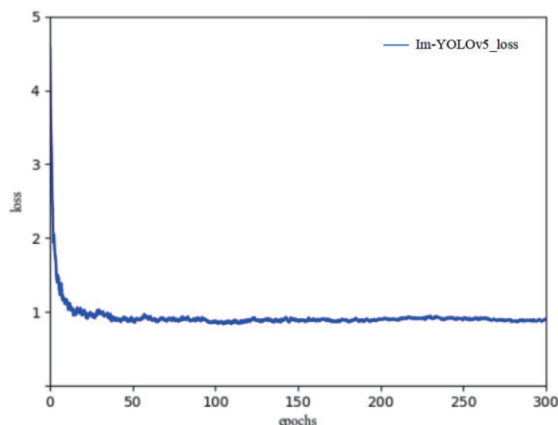


(a) YOLOv5 损失函数曲线

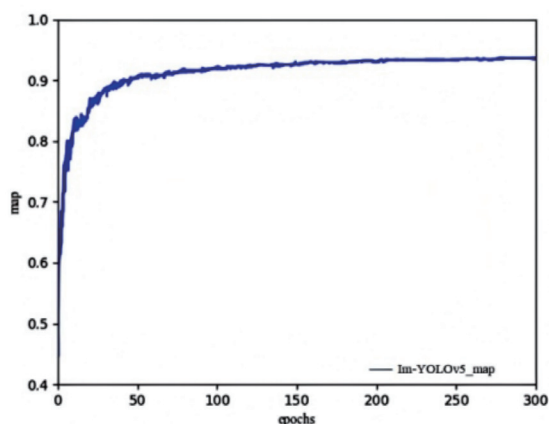


(b) YOLOv5 平均精度值曲线

图 6 YOLOv5 实验结果曲线



(a) Im-YOLOv5 损失函数曲线



(b) Im-YOLOv5 平均精度值曲线

图 7 Im-YOLOv5 实验结果曲线

由图 7 (a) 可以得到，Im-YOLOv5 模型初始损失值约为 4.5，在迭代次数至 30 轮左右达到平稳，平稳值为 0.92 左右。从图 7 (b) 中可以得到，Im-YOLOv5 模型的 mAP 值为 93% 左右。因此，相比于原始 YOLOv5 模型，Im-YOLOv5 模型对员工安全帽佩戴检测的精度更高。不同算法检测结果对比数据如表 2 所示。

表 2 不同算法检测结果对比

算法	迭代次数	训练批次	mAP/%
YOLOv3	300	64	88.1
YOLOv5	300	64	91.3
Im-YOLOv5	300	64	93.5

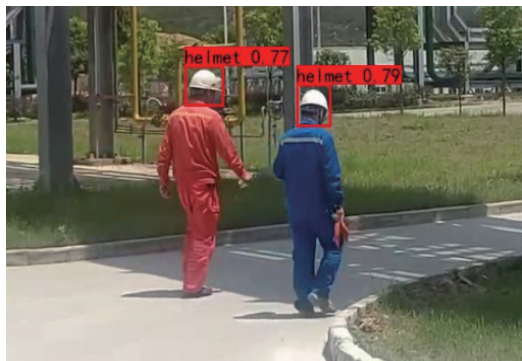
图 8 为员工安全帽佩戴的检测结果，其中标签 helmet 为佩戴安全帽，标签 head 为未佩戴安全帽。场景 1 为生产装置内工作人员；场景 2 为夜间园区行走人员；场景 3、场景 4 为白天园区行走人员。由结果可以看出，改进后的 Im-YOLOv5 模型检测准确率较高，具有较好的检测效果。



(a) 场景 1



(b) 场景 2



(c) 场景 3



(d) 场景 4

图 8 检测结果

4 结论

(1) 采用多尺度 Retinex (MSR) 图像处理算法对监控视频图像进行处理, 有效削弱了光照不足等因素对检测结果

的影响, 经过处理后视频图像的色彩得到了均衡处理, 特征信息更加丰富。

(2) 采用动态注意力机制 (SKNet) 对 YOLOv5 算法模型进行改进, 提高了特征信息的提取能力, 有效提高了员工安全帽佩戴检测的精度。

参考文献:

- [1] 王占舟, 许耀文, 赵霞, 等. 现代化化工企业事故分析及其应对策略 [J]. 化工管理, 2021(22):95-97.
- [2] 丁文龙, 费树珉. 基于改进 YOLOv3 的安全帽检测方法研究 [J]. 电子测试, 2022(11):84-86.
- [3] 李帅, 李丽宏, 王素刚, 等. 改进 YOLOv4 算法的安全帽检测 [J]. 现代电子技术, 2022, 45(3):103-110.
- [4] HAYAT A, MORGADO-DIAS F. Deep learning-based automatic safety helmet detection system for construction safety[EB/OL].(2022-08-18)[2024-04-29]. <https://www.mdpi.com/2076-3417/12/16/8268>.
- [5] LI W, LI S, WANG Y H, et al. Study on personnel detection based on retinex and YOLOv4 in building fire[EB/OL]. (2021-11-19)[2024-06-19].<https://iopscience.iop.org/article/10.1088/1742-6596/2185/1/012039>.
- [6] 李旺, 杨金宝, 孙婷, 等. 基于 Retinex 的多尺度单幅图像去雾网络 [J]. 青岛大学学报 (自然科学版), 2022, 35(4): 26-32.
- [7] LECCA M, GIANINI G, SERAPIONI R P. Mathematical insights into the original retinex algorithm for image enhancement[J].Journal of the optical society of america, 2022, 39(11): 2063-2072.
- [8] 余雅琪, 杨梦龙. 基于 Retinex 理论的全卷积网络低光图像增强方法 [J]. 现代信息科技, 2022, 6(17): 1-7.

【作者简介】

张硕羲 (1996—), 男, 甘肃庆阳人, 硕士, 工程师, 研究方向: 图像处理及目标识别。

(收稿日期: 2024-09-19)