

基于改进YOLOv13模型的光伏板缺陷检测方法研究

王旭涛*, 邓 豪, 李凌霄#

重庆理工大学数学科学学院, 重庆

收稿日期: 2025年9月11日; 录用日期: 2025年10月4日; 发布日期: 2025年10月11日

摘 要

针对目前光伏板缺陷检测中仍存在小目标检测率低、检测速度慢以及适应性差等问题, 提出了一种基于YOLOv13n的改进光伏板缺陷检测模型YOLOv13n-PV。在YOLOv13的Backbone网络结构中引入了CA注意力机制, 这种机制使得算法能够更加专注于目标的位置信息和类别判定, 从而有效地提升对于小目标和密集目标的特征提取能力。同时, 对于损失函数进行优化, 采用类别加权和小目标加重的双重加权机制, 增强模型对缺陷的敏感度和加强对小缺陷的检测效果。通过各项试验结果表明, 本文提出的算法模型在测试数据集下的平均准确率、召回率分别为94.1%和93.6%。分别优于原始YOLOv13n模型算法的93.2%和91.7%。且模型的计算量没有过多增加, 在提升检测性能的同时能够兼顾算法计算效率, 因此可以快速、准确地实现光伏板的缺陷检测, 为新能源系统中的光伏法发电提供技术支持。

关键词

光伏板, YOLOv13, CA注意力机制, 深度学习, 损失函数, 缺陷检测

Research on a Photovoltaic Panel Defect Detection Method Based on an Improved YOLOv13

Xutao Wang*, Hao Deng, Lingxiao Li#

School of Mathematical Sciences, Chongqing University of Technology, Chongqing

Received: September 11, 2025; accepted: October 4, 2025; published: October 11, 2025

*第一作者。

#通讯作者。

Abstract

To address the persistent issues in photovoltaic panel defect detection—namely the low detection rate for small targets, slow inference speed, and poor adaptability—we propose an improved defect detection model, YOLOv13n-PV, based on YOLOv13n. In the YOLOv13 backbone, we incorporate the Coordinate Attention (CA) mechanism, which enables the algorithm to focus more effectively on spatial localization and category discrimination, thereby enhancing feature extraction for small and densely distributed targets. Meanwhile, the loss function is optimized by adopting a dual-weighting mechanism that combines class weighting and small-object weighting. This approach enhances the model's sensitivity to defects and strengthens the detection performance for small defects. The results of various experiments show that the average precision and recall of the algorithm model proposed in this paper on the test dataset are 94.1% and 93.6%, respectively. The values are respectively higher than 93.2% and 91.7% of the original YOLOv13n model algorithm. While improving detection performance, it can also balance the computational efficiency of the algorithm. Therefore, it can realize fast and accurate defect detection of photovoltaic panels, providing technical support for photovoltaic power generation in new energy systems.

Keywords

Photovoltaic Panel, YOLOv13, CA Attention Mechanism, Deep Learning, Loss Function, Defect Detection

Copyright © 2025 by author(s) and Hans Publishers Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

1. 引言

在实现“双碳”目标的新形势下，太阳能技术创新和应用正迎来前所未有的发展机遇[1]。太阳能应用的主要形式是光伏发电。其中 2024 年上半年新增并网容量 102.48 GW，其中集中式光伏电站 49.60 GW，分布式光伏 52.88 GW [2]。在发电系统中，光伏板将光能转化为电能。但是在此过程中，由于光伏板受到紫外线、腐蚀以及恶劣天气下受潮等影响，可能会引发断栅、隐裂等问题，这些缺陷会对光伏发电转化效率造成严重的影响，进而影响整个发电系统的稳定性[3]。

因此，如何及时有效地对光伏板进行有效的缺陷检测非常重要。早期的光伏板缺陷检测主要是通过人工目测、光伏板的红外热成像检测以及一些传统的机器视觉。但早期的人工检测方法成本高、效率低以及误检率高等不足，并且接触光伏板时可能会造成二次损坏，使得目前检测难以满足现代光伏板生产效率高、性能好的检测要求。

基于以上背景，研究目的是将深度学习技术拓展到光伏板缺陷检测领域，从而实现一个具有实时检测且精度高的光伏板缺陷检测的算法。

缺陷检测随着计算机技术的发展，深度学习在缺陷检测领域被广泛应用且取得了显著的成果。基于深度学习的目标检测算法在步骤上包括双阶段(two-stage)和单阶段(one-stage)两种目标检测框架[4]。在双阶段的检测算法中，R-CNN [5]是一种基于区域的卷积神经网络，通过使用选择性搜索算法来对输入图像提取潜在的候选区域。然后，对每个候选区域进行特征提取，并使用支持向量机(SVM)进行分类。这样就将目标检测问题转化为一个候选区域分类的任务。而 Faster R-CNN [6]在 R-CNN 的基础上进行了改进，

引入了区域建议网络(Region Proposal Network, RPN),使得整个目标检测系统实现端到端地进行训练。但由于候选区域生成过程可能存在一些限制,导致漏检或定位不准确。

而在单阶段算法中,SSD (Single Shot MultiBox Detector) [7]采用了基于卷积神经网络(CNN)的特征提取器,并在多个不同尺度的特征图上进行目标检测。它通过在不同层级的特征图上应用不同大小的卷积核来检测不同尺寸的目标。这种多尺度的检测策略使得 SSD 能够有效地检测不同大小的目标。而 YOLO (You Only Look Once)也是一种单阶段的目标检测算法。与 SSD 不同的是, YOLO 是将目标检测问题转化为一个回归问题。它将输入图像分成一个固定大小的网格,并在每个网格单元中预测目标的边界框和类别。

YOLO [8]算法凭借检测速度快,且只需要一次前向传播就可以得到所有目标的检测结果的特点,实现了对目标的实时检测。YOLO 作为单阶段模型的代表之一,相较于更早提出的两阶段目标检测算法,不仅拥有更快的预测速度;对于背景图像(非物体)中的部分被包含在候选框的情况误检率更低,还拥有更好的算法通用性。以上这些特性,都使 YOLO 系列模型成为工业目标检测场景首选的算法。

2. YOLOv13 算法

YOLOv13 [9]算法于 2025 年 6 月提出,与之前的 YOLO 系列算法相比, YOLOv13 在减少了计算量的同时还提高了检测的精度能够更好地识别和定位目标物体,对小目标的检测效果也较好。YOLOv13 的具体创新点具体包括引入了 HyperACE 机制来用于捕捉高阶多对多语义关系,提出了 FullAD 架构,将特征增强贯穿于 Backbone、Neck、Head 整个网络结构以及采用 DSC3k2 模块来构建轻量模块,降低模型的复杂度。

YOLOv13 网络结构

YOLOv13 模型的网络结构依然是包括 Backbone、Neck 以及 Head 三部分构成,其具体结构见下图 1。

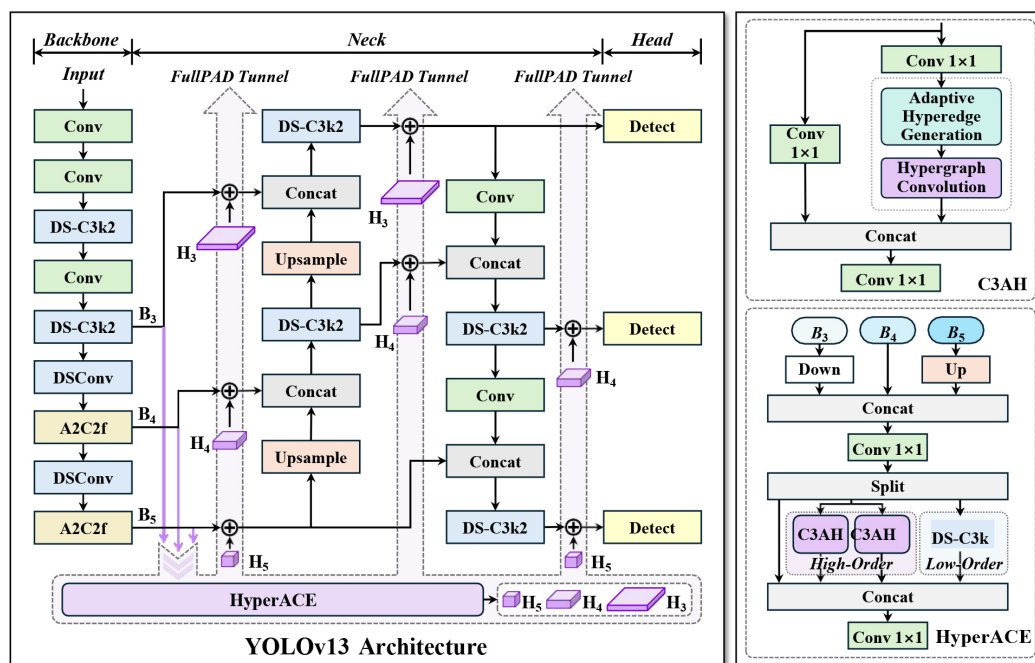


Figure 1. Network structure of the YOLOv13 mode

图 1. YOLOv13 网络结构图

在卷积神经网络，它通过卷积核与输入图像进行卷积运算，图像中的像素信息转化为网络能够理解和利用的特征图这一过程不仅高效地提取了图像的关键特征，还为后续的网络层提供了有力的支持，进一步增强了整个网络在图像识别等任务中的性能。卷积的具体操作步骤如下：首先将卷积核放在图像矩阵的最左上角，便覆盖了这一角的对应数值，将重叠块的数值进行乘积，最后将得到的所有乘积值累加便得到了一个输出值，置于新生成矩阵的左上角，接下来卷积核右移，重复上述过程得到对应位置新的输出，通过滑窗的方式遍历整个原图像矩阵后，新的特征矩阵便生成完毕。卷积层就是由多个卷积滤波器组成，卷积操作就是基于卷积滤波器来进行计算，其计算过程见图 2。

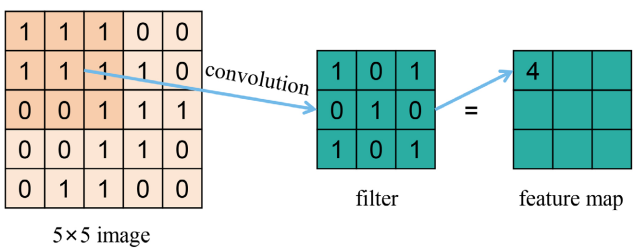


Figure 2. The process of convolution calculation
图 2. 卷积计算过程

YOLOv13 网络结构的 Backbone 部分使用了 Conv、A2C2f、DSConv 以及 DSC3k2 提取多尺度特征图 B_1 、 B_2 、 B_3 、 B_4 和 B_5 ，其中 Conv 作为基础特征提取单元来构建底层特征，而 DSConv 和 DSC3k2 通过轻量化的设计来降低计算负担，最后 A2C2f 模块再弥补轻量化带来的精度损失。

HyperACE 是 YOLOv13 提出的核心模块(自适应超图计算见图 3)，它介于 Backbone 与 Neck 网络结构之间，基于 C3AH 模块的全局高阶感知分支和基于 DS-C3k 块的局部低阶感知分支。C3AH 模块通过自适应超图计算对高阶视觉关联进行线性复杂度建模，保留了 CSP bottleneck 分支分裂机制，同时集成了自适应超图计算模块，实现了跨空间位置的全局高阶语义聚合。解决了之前的模型只能建模局部两元关系的问题。

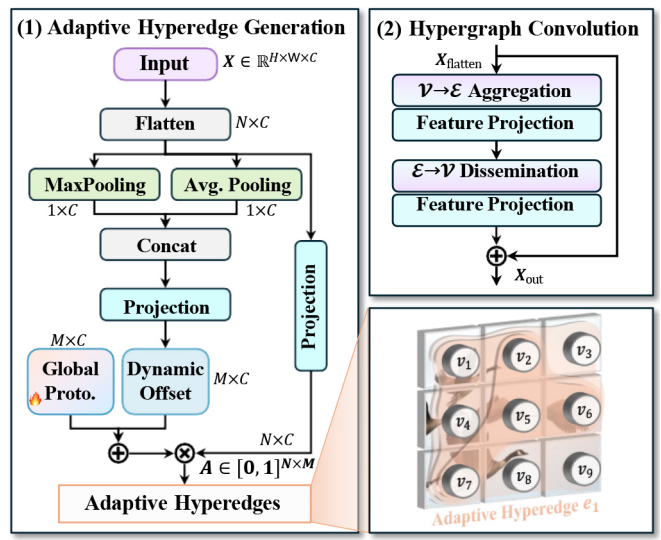


Figure 3. Details of the adaptive hypergraph construction and convolution
图 3. 自适应超图计算原理

YOLOv13 网络的 Neck 结构在设计上聚焦于特征融合的性能提升：一方面，通过嵌入 DSConv 与 A2C2f 等高效模块，从结构层面优化了融合流程、提升了运算效率，同时强化了特征的区分度与表征能力；另一方面，针对 Backbone 输出的不同尺度特征图，该部分通过跨尺度融合机制实现特征信息的深度交互，让最终生成的每一个特征层既保留了底层特征的细节纹理信息，又融入了高层特征的语义类别信息，从而为网络的检测精度提升奠定基础。

YOLOv13 的 Head 部分最主要的作用是多尺度目标预测，其中的 Conv 与 DSC3k2 是通过对 Neck 输出的特征进行提取、优化以及增强区分度。Concat 利用局部跨尺度特征融合适配不同尺寸目标的特征表达。而 Detect 则是设置了三个检测分支来分别对应 Neck 的高、中、低分辨率特征来分别检测小、中、大目标最终输出坐标位置、类别以及置信度。

3. 基于 YOLOv13n 的改进光伏板缺陷检测算法

本文通过改进 YOLOv13n 模型的结构，来改进其在光伏板缺陷检测中针对裂纹、划痕以及断栅等小缺陷在检测中检测效果不足。针对于光伏板的缺陷检测，需要根据数据集的特征来对模型进行改进和优化。在 YOLOv13 的网络结构的 Backbone 中引入 CA 注意力机制，使得其更好的对小目标和密集目标的特征进行提取。同时在此基础上还对损失函数进行了改进，通过类别加权和小目标加权双重优化，使得对小目标缺陷的检测更准确。

3.1. CA 注意力机制

坐标注意力机制(coordinate attention, CA, 见图 4) [10]。是将位置信息加入到了通道注意力当中，使网络可以在更大区域上进行注意。为了缓解以前的注意力机制如 SENet, CBAM 等提出的二维全局池化造成的位置信息丢失、该注意力机制将通道注意分解为两个平行的一维特征编码的过程，分别在两个方向上聚合特征，一个方向得到准确的位置信息，另一个方向得到远程依赖关系，对生成的特征图进行编码以形成一对方向感知和位置敏感的特征。

CA 通过精准的位置信息编码通道关系和长程依赖关系，主要包括两个步骤：坐标信息嵌入和坐标注意力生成，一个 CA 模块可以看作是一个用来增强特征表示能力的计算单元它可以将任何中间张量 $X = [x_1, x_2, \dots, x_c] \in R^{C \times H \times W}$ 作为输入并输出一个有增强表示能力的同样尺寸的输出 $Y = [y_1, y_2, \dots, y_c]$ ，其中 C 为通道数， H 和 W 分别为输入图片的高和宽。

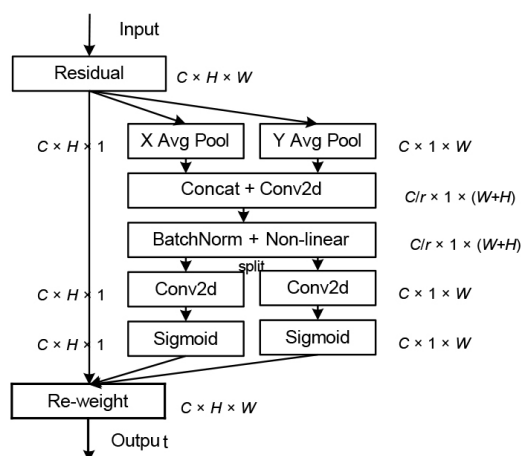


Figure 4. Structure of the CA module

图 4. CA 模块结构

对于坐标信息嵌入来说, 为了促进注意力模块可以获取精准位置信息的空间长距离依赖关系, CA 模块将全局池化分为一对一维特征编码操作。对于输入的特征图 X , 维度为 $C \times H \times W$, 先使用大小为 $(H, 1)$ 和 $(1, W)$ 的池化核分别沿水平方向坐标和竖直方向的坐标对每个通道进行编码, 也就是高度为 h 的第 c 个通道与宽度为 w 的第 c 个通道的输出, 输出公式如式下所示:

$$Z_c^h(h) = \frac{1}{W} \sum_{0 \leq l < W} X_c(l, h) \quad (1)$$

$$Z_c^w(w) = \frac{1}{H} \sum_{0 \leq j < H} X_c(j, w) \quad (2)$$

以上公式返回一对方向感知注意力特征 Z^h 和 Z^w , 使得网络对目标的定位更加精确。

对于坐标注意力生成来说, 级联之前模块先生成两个特征层, 然后使用一个共享的 1×1 卷积进行变换 F_1 。并激活, 其公式如下所示:

$$f = \delta\left(F_1\left(\begin{bmatrix} Z^h \\ Z^w \end{bmatrix}\right)\right) \quad (3)$$

其中 $f \in R^{C/r \times (H+W)}$ 是对空间信息在水平方向和竖直方向的中间特征图, r 表示下采样比例。 $[a, b]$ 表示沿空间维度的连接操作, δ 表示非线性激活函数。然后, 沿空间维度将 f 切分成两个单独的张量 $f^h \in R^{C/r \times H}$ 和 $f^w \in R^{C/r \times W}$ 再利用两个 1×1 卷积 F_h 和 F_w 将特征图 f^h 和 f^w 变换到和书如 X 同样的通道数, 得到结果如公式 所示:

$$g^h = \sigma\left(F_h\left(J^h\right)\right) \quad (4)$$

$$g^w = \sigma\left(F_w\left(J^w\right)\right) \quad (5)$$

最后, 对 g^h 和 g^w 进行拓展, 作为注意力权重, CA 模块的最终输出可以表示如下所示:

$$y_c(i, j) = x_c(i, j) \times g_c^h(i) \times g_c^w(j) \quad (6)$$

Coordinate Attention 作为一种简单、灵活且易于集成的模块, 展现了一种创新的移动网络注意力机制。其独特之处在于, 无需增加额外的参数, 即可有效提升网络的精度。

3.2. 损失函数

YOLOv13 中目标检测的损失函数包括 CLS、IoU 以及 DFL 三部分, 由这三部分加权求和得到。

3.2.1. CLS

在目标检测中, 通常情况下背景的样本量通常比目标的样本量要多, 如果使用普通的交叉熵损失, 那么模型可能会被大量的易分类的背景样本主导, 从而导致对目标的样本学习不足。YOLOv13 的分类损失中通常采用 Focal Loss, 这种分类损失针对类别不平衡的问题, 通过 $(1 - p_t)^\gamma$ 来调节因子, 从而降低易分类负样本的权重, 聚焦难以分类的样本, 其计算原理如下:

$$L_{\text{cls}} = -\alpha \cdot (1 - p_t)^\gamma \cdot y \cdot \log(p_t) - (1 - \alpha) \cdot p_t^\gamma \cdot (1 - y) \cdot \log(1 - p_t) \quad (7)$$

其中, y 的取值为 0 (背景) 或 1 (目标), p_t 表示预测的概率, α 是平衡因子, 用于调节正负样本的权重, γ 是聚焦参数, 其值越大代表对样本越关注。

3.2.2. CIoU 损失

目标检测中的边界框回归损失函数是通过学习预测边界框的位置, 使得模型尽可能地接近真实的边界框。进而提供检测目标的精确定位和区域的关键信息。而 YOLOv13 的边界框损失中使用了基于 CIoU

损失, 增强缺陷类别的小目标权重, CIoU 是对于 IoU 损失的改进, 其计算方法如下:

$$L_{iou} = 1 - \text{CIoU}(b, b_{gt}) \quad (8)$$

$$\text{CIoU} = \text{IoU} - \frac{\rho^2(b, b_{gt})}{c^2} - \alpha \cdot v \quad (9)$$

$$\text{IoU} = \frac{|b \cap b_{gt}|}{|b \cup b_{gt}|} \quad (10)$$

$$v = \frac{4}{\pi^2} \cdot \left(\arctan \frac{w_{gt}}{h_{gt}} - \arctan \frac{w}{h} \right)^2 \quad (11)$$

$$\alpha = \frac{v}{(1 - \text{IoU}) + v} \quad (12)$$

其中 c 表示包围两框的最小外接矩形对角线长度, $\rho^2(b, b_{gt})$ 表示预测框 b 与真实框 b_{gt} 的中心欧氏距离平方。

3.2.3. 分布焦点损失

DFL 是为了解决传统连续值回归的局限性, 通过将坐标建模为离散的概率的分布, 从而来提升边界框定位的精度, 在小目标和边缘目标的检测中起到了重要的作用。其对于边界框坐标的离散分布损失, 同样采用强关键样本权重, 计算方式为:

$$L_{df} = \frac{1}{S_{\text{target}}} \sum_{fg} \left[\text{DFL}(\hat{d}, d) \cdot w_{\text{class}} \cdot w_{\text{small}} \right] \quad (13)$$

其中 $\text{DFL}(\hat{d}, d)$ 代表分布焦点损失, 将连续坐标值建模为离散分布, 计算并预测分布 $\hat{d}$ 与目标分布 d 的交叉熵, 即:

$$\text{DFL}(\hat{d}, d) = \text{CE}(\hat{d}, \bar{d}) \cdot (\bar{d} - d) + \text{CE}(\hat{d}, \underline{d}) \cdot (d - \underline{d}) \quad (14)$$

其中 $\underline{d}$ 为 d 的向下取整, $\bar{d}$ 为向上取整, CE 代表交叉熵损失。

3.3. 改进损失函数

改进之后的损失函数加入了类别加权机制、小目标加权机制以及动态计算损失。

首先是类别加权机制, 在实际的场景中, 光伏板中正常区域的样本数量要远远多于缺陷区域, 如果没有类别加权, 则模型会更倾向于降低整体损失而导致部分缺陷被漏检。而加权损失则是对于不同类别的缺陷赋予其不同的权重, 以此来提高对该类别损失的关注度。通过损失加权的方式, 可以使模型在训练中对于缺陷类分配更多的注意力, 最终实现漏检率更低, 检测更准确的检测目标。对于不同类别的缺陷, 其类别加权的表达式为:

$$L_{cls} = \frac{1}{N_{\text{pos}}} \sum_{i=1}^N \sum_{c=1}^C w_c \cdot y_{i,c} \cdot \text{BCE}(\hat{y}_{i,c}) \quad (15)$$

其中 N_{pos} 为正样本总数, N 为总的样本数, C 为缺陷类别的总数, w_c 表示第 c 类的权重, $y_{i,c}$ 和 $\hat{y}_{i,c}$ 分别表示第 i 个样本是否属于第 c 类标签以及属于 c 类别的预测概率。对于小目标缺陷来增大其 w_c 的值来实现对小目标类别赋予更高的权重。

其次是小目标加权机制。在训练过程中小目标因为在图像中占据的像素小而导致其表达能力不足而

导致特征提取困难，同时因为目标小而导致边界框微小偏移对 IoU 的影响要大于大目标而造成对定位的误差更敏感。而小目标加权机制则通过计算目标缺陷的面积，对于面积小于指定阈值的目标判定为小目标，对于小目标对其损失进行加权放大以此来动态调整缺陷的权重，使小缺陷的损失在总损失中占比更加合理。其计算方法为：

$$L_{box} = \frac{1}{N_{pos}} \sum_{i=1}^{N_{pos}} [1 + w_s \cdot I(S_i < S_t)] \cdot w_{c_i} \cdot L_{iou}(b_i, \hat{b}_i) \quad (16)$$

其中 S_i 表示缺陷目标的面积， S_t 表示小目标面积的阈值， I 为指示函数，当判断条件成立时为 1，否则为 0。当缺陷目标面积小于小目标面积的阈值时则判定为小目标，对于小目标则额外加入额外权重 w_s 。

最后是多损失组件权重融合，如果只在单一损失中应用加权可能会导致模型优化失衡，而多损失组件加权融合则在分类、CIoU 以及 DFL 三个损失中同步应用缺陷和小目标权重，从而实现模型在检测的准确度和定位的精度上都可以得到优化。

3.4. 改进后的 YOLOv13 网络结构设计

根据以上内容，本文将在 YOLOv13 原有的网络结构的 Backbone 部分引入 CA 注意力机制，并对损失函数进行优化，改进后 YOLOv13 的 Backbone 部分的网络结构见图 5。

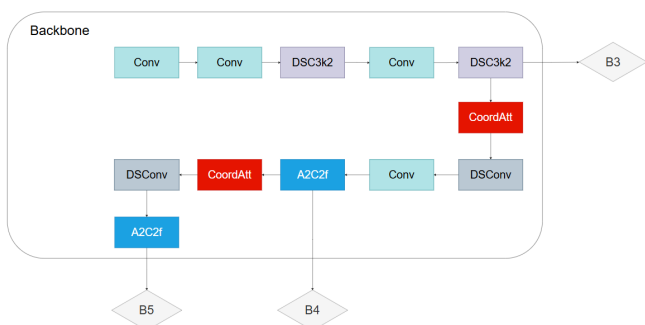


Figure 5. The improved Backbone structure
图 5. 改进后的 Backbone 结构

4. 实验

4.1. 实验数数据集的准备

实验数据集来自北京航空航天大学与河北工业大学发布的光伏板缺陷数据集，其中包括了 12 个不同类别的异常缺陷一共 3684 张图像。数据集示例见图 6、图 7。

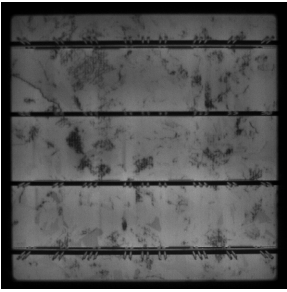


Figure 6. Examples of the dataset
图 6. 数据集示例

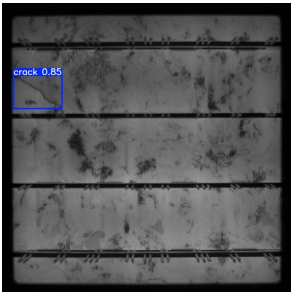


Figure 7. Demonstration of dataset defects
图 7. 数据集缺陷展示

4.1.1. 数据集的划分

根据含缺陷的光伏板数据集的大小，将训练集、验证集、测试集划分比例设置为 7:2:1。即 2578 张图片设置为训练集，736 张设置为验证集，361 张设置为测试集。为了实验的可靠性有一定的保证，对数据集进行划分时，采用随机打乱顺序划分，这样划分可以消除数据的顺序性和相关性，减少模型受到数据排列顺序的影响。

4.1.2. 数据集的转换

由于光伏板数据集的标签采用的是 XML 形式标注的，而 YOLO 只能识别 TXT 形式的标签。因此，在进行训练之前需要将数据集的标签进行转化，使得模型可以正确读取和识别到标签中的信息。

4.2. 实验平台

4.2.1. 实验配置

实验配置见表 1。

Table 1. Configuration of the experimental environment
图 1. 实验环境配置

类型	参数
CPU	12th Gen Intel(R) Core(TM) i5-12600K
GPU	NVIDIA GeForce RTX 3080 Ti
操作系统	Linux
显存	12GB
Python 版本	3.9.5
Pytorch 版本	2.4.1
开发环境	VScode
加速环境	Cuda11.8

4.2.2. 实验评价指标

我们可以通过下面这些指标来对我们算法的性能进行衡量：
在这些常用的指标中有这样一些参数如下表 2 所示：

Table 2. Common parameters
表 2. 常用参数

	正样本	负样本
预测正样本	<i>TP</i>	<i>FP</i>
预测负样本	<i>FN</i>	<i>TN</i>

其中： TP 表示正样本中被预测为正样本的个数； FP 表示负样本的被预测为正样本的个数； FN 表示正样本中被预测为负样本的个数； TN 表示负样本中被预测为负样本的个数。

由以上几个参数可以得出以下三个评价指标：

Precision (精确率)：在预测为正样本的目标中，正样本的比例。其计算公式为：

$$Precision = \frac{TP}{TP + FP} \quad (17)$$

Accuracy (准确率)：在所有被检测的目标中，被正确预测的比例。其计算公式为：

$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN} \quad (18)$$

Recall (召回率)：在正样本中，被预测为正样本的比例。其计算公式为：

$$Recall = \frac{TP}{TP + FN} \quad (19)$$

通常情况下，我们很希望 **Precision**、**Recall** 都越高越好。但是 **Precision** (精确率)和 **Recall** (召回率)之间存在一种权衡关系，调整分类模型的阈值会对它们产生影响。所以我们需要构建 **Precision-Recall** 曲线来帮助分析。通过将 **Recall** 作为 X 轴，将 **Precision** 作为 Y 轴，在 **Precision-Recall** 空间中绘制所有可能的点，每个点对应于不同的预测阈值下模型的性能，连接这些点就形成了 **P-R** 曲线。通过分析 **P-R** 曲线，可以选择合适的阈值来平衡 **Precision** 和 **Recall**，以满足具体任务需求。

mAP (mean Average Precision)：即平均精度均值，在所有类别中，平均得到整个数据集的平均精确度，便可得到整个数据集的 **mAP** 值。其计算方式为：

$$mAP = \frac{1}{n} \sum_{i=1}^n \int_0^1 F_{P-R} dx \quad (20)$$

4.3. 实验结果

本实验实中 epoch 设置了 1024, 每个批次传入的图片数 batch_size 设置为 4, 初始学习率设置为 0.01, 用 GPU 设备进行训练, 输入尺寸设置为 640×640 , workers 设置为 0, 不启用混合精度进行训练。

为验证添加 CA 注意力机制以及改进损失函数的有效性, 对原始的 YOLOv13n 模型分别进行只加入 CA 注意力机制、只改进损失函数以及同时加入 CA 注意力机制与改进损失函数进行消融实验, 而 CA 注意力机制在 YOLOv13n 的 Backbone 部分进行添加。实验结果见下表 3。

Table 3. Experimental results

表 3. 实验结果

Model	Position	P	R	mAP50	mAP50-95
YOLOv13n		0.914	0.912	0.934	0.604
YOLOv13n + CA	Backbone	0.925	0.907	0.944	0.596
YOLOv13 + 改进 Loss	Loss	0.933	0.919	0.951	0.609
YOLOv13n-PV	Backbone	0.931	0.928	0.957	0.632

根据实验结果, 仅在 YOLOv13n 模型的 Backbone 部分添加 CA 注意力机制后, 其模型的准确率 P 与回归精度 mAP 分别提升了 0.011 与 0.01, 但是召回率 R 与 $mAP50-95$ 有所下降, 说明 CA 注意力机制在此过程中弱化了低响应特征区域从而导致了召回率与 $mAP50-95$ 的下降。而改进了损失函数后在检测精度与回归精度方面分别提升了 0.019 与 0.017, 说明了改进损失函数后模型对缺陷的检测更加准确。

与原 YOLOv13n 模型训练得到的结果相比较，在 Backbone 部分添加 CA 注意力机制并改进损失函数之后模型的准确率 P、召回率 R 和回归精度 mAP 都有所提升，准确率提升了 1.85%，召回率提升了 1.74%，平均精度也提升了 0.023。同时，与仅在 Backbone 部分增加 CA 注意力机制相比，模型的召回率 R 与 mAP50-95 也有所提升，这表明在增加了 CA 注意力机制且对损失函数改进后的 YOLOv13-PV 模型在光伏缺陷检测方面更准确、全面。

而增加的 CA 注意力机制为轻量化模块，参数量与计算量较低，且损失函数的修改不会改变模型的结构与推理过程。修改后的 YOLOv13n-PV 模型与原模型 YOLOv13n 在光伏板缺陷数据集上的参数比较见表 4。

Table 4. Comparison of model parameters
表 4. 模型参数对比

模型	参数量(M)	GFLOPs (G)	FPS
YOLOv13n	2.5	6.4	507.6
YOLOv13n-PV	2.55	6.43	505

由表可知，对于改进后的模型，参数量等方面仅有少量增加，推理速度也未明显下降，各方面变化幅度均小于 1%。因此，该模型在提升了检测精度的同时仍兼顾了计算效率。

同时为了验证模型的泛化能力与其可靠性，在测试集上对改进后的模型进行验证，并将各个指标进行对比，结果见图 8。

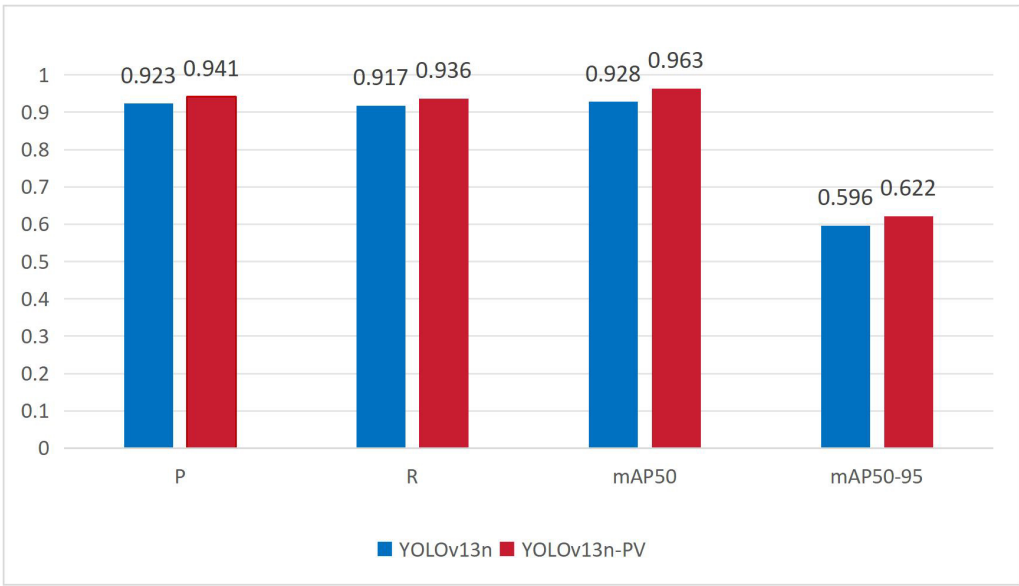


Figure 8. Model metrics on the test set
图 8. 测试集中模型指标

由图可知，改进后的 YOLOv13-PV 光伏板缺陷检测模型在测试集上的多个指标要更好，其中准确率比原始模型提高至 0.941，召回率提高至 0.936，且平均精度方面也有所提高。且改进后的 YOLOv13-PV 模型对于小目标缺陷有更好的检测效果，原始模型与改进后模型对于“裂纹”、“断角”以及“星型裂纹”三种缺陷检测效果对比见图 9。

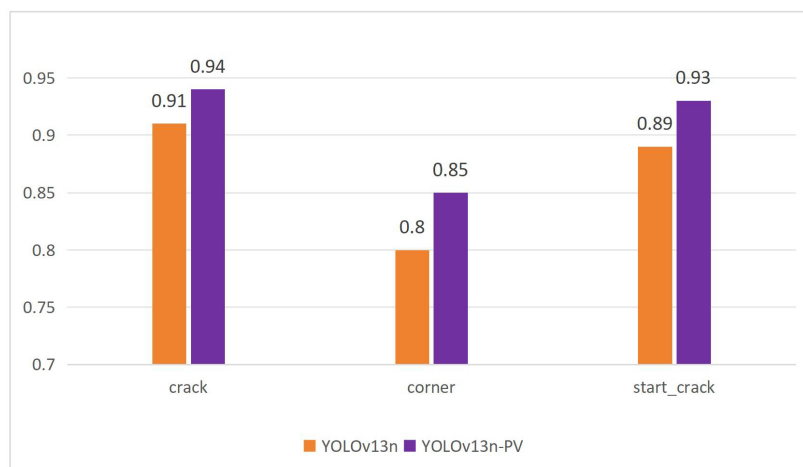


Figure 9. Comparison of defect detection effects among three types
图 9. 三种缺陷检测效果对比

改进后的缺陷检测效果对比图如下图 10~14 所示:

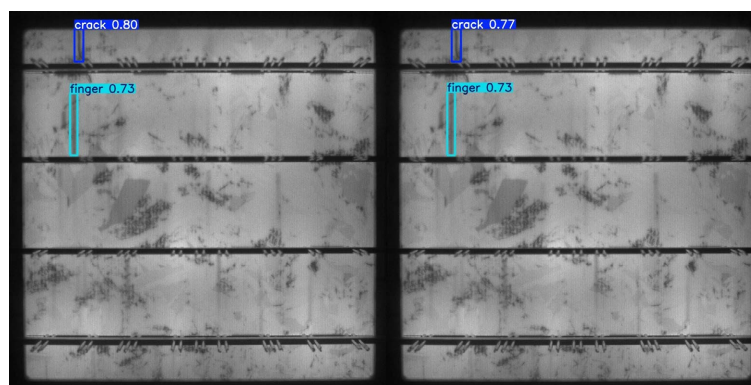


Figure 10. Comparison of prediction performance for crack defects
图 10. 裂纹缺陷预测效果对比

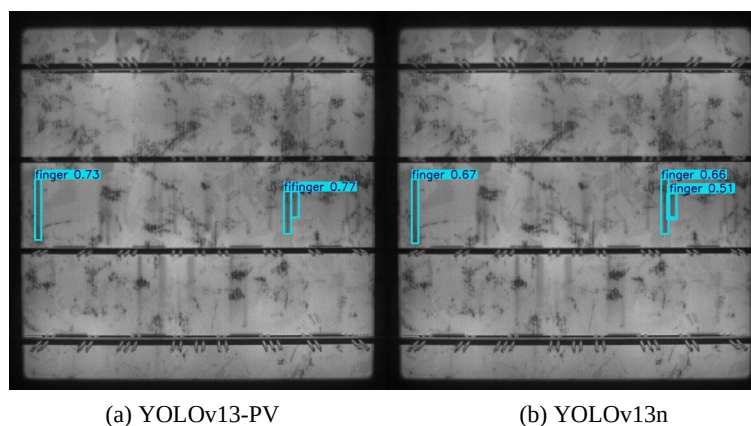


Figure 11. Comparison of prediction performance for crack defects
图 11. 断栅缺陷预测效果对比

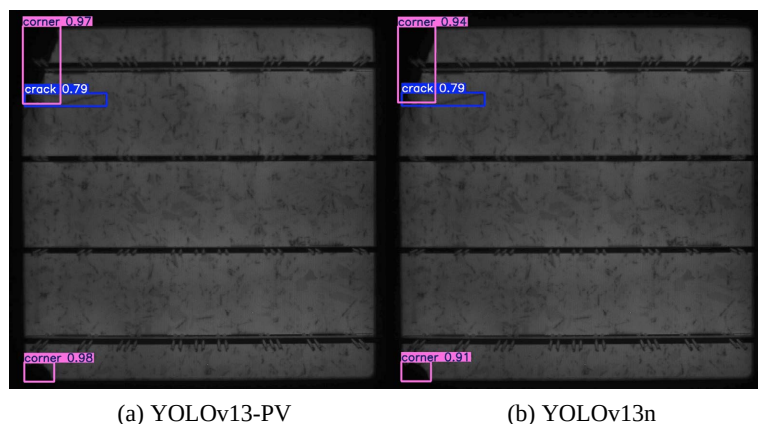


Figure 12. Comparison of detection and prediction performance for corner defects

图 12. 折痕缺陷检测预测效果对比

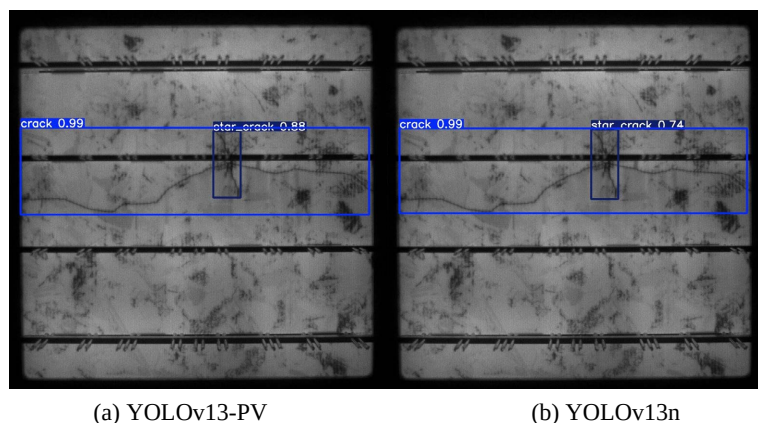


Figure 13. Comparison of detection and prediction performance for corner defects

图 13. 星状裂纹缺陷检测效果对比

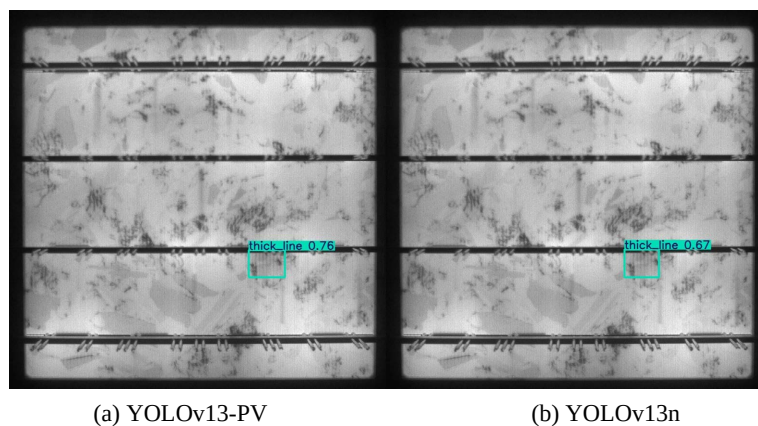


Figure 14. Comparison of detection performance for thick-line defects

图 14. 粗线缺陷检测效果对比

综上，改进后的光伏板缺陷检测模型在光伏板缺陷检测方面效果更好，通过引入了 CA 注意力机制以及改进损失函数之后使得模型在各方面上都得到了优化，在光伏板的缺陷检测方面有了一定的提升。

5. 展望

本文提出了一种以基于 YOLOv13n 模型为基础的光伏板缺陷检测模型 YOLOv13n-PV, 通过引入 CA 注意力机制以及改进损失函数来对模型进行优化。该模型可以快速准确地检测出光伏板的各种缺陷, 大大提升了对太阳能发电板缺陷的检测效率, 节约了人力和物力的成本, 提高了太阳能板发电的效率。其次, YOLOv13-PV 模型也有很好的泛化能力和识别能力, 可以适应不同的光伏板检测需求。由于光伏板种类较少, 因此该模型可以用来适应各种光伏板的缺陷检测需求。后续研究会进一步考虑对模型算法的复杂度进行优化, 使其在低算力设备上可以有更好的表现效果, 并且对于数据集进行算法的微调(如光伏板热成像数据集), 使其可以更好应用于各类情况。

基金项目

重庆理工大学本科科研项目(KLC24115)。

参考文献

- [1] 舒印彪, 赵勇, 赵良, 等. “双碳”目标下我国能源电力低碳转型路径[J]. 中国电机工程学报, 2023, 43(5): 1663-1672.
- [2] 国家能源局. 2024 年上半年光伏发电建设情况[J]. 电力科技与环保, 2021, 31(4): 46.
- [3] Waqar Akram, M., Li, G., Jin, Y. and Chen, X. (2022) Failures of Photovoltaic Modules and Their Detection: A Review. *Applied Energy*, **313**, Article 118822. <https://doi.org/10.1016/j.apenergy.2022.118822>
- [4] 宁健, 马淼, 柴立臣, 等. 深度学习的目标检测算法综述[J]. 信息记录材, 2022, 23(10): 1-4.
- [5] Girshick, R., Donahue, J., Darrell, T. and Malik, J. (2014) Rich Feature Hierarchies for Accurate Object Detection and Semantic Segmentation. 2014 *IEEE Conference on Computer Vision and Pattern Recognition*, Columbus, 23-28 June 2014, 580-587. <https://doi.org/10.1109/cvpr.2014.81>
- [6] Jiang, H. and Learned-Miller, E. (2017) Face Detection with the Faster R-CNN. 2017 *12th IEEE International Conference on Automatic Face & Gesture Recognition*, Washington, 30 May-3 June 2017, 650-657. <https://doi.org/10.1109/FG.2017.82>
- [7] Liu, W., Anguelov, D., Erhan, D., Szegedy, C., Reed, S., Fu, C., et al. (2016) SSD: Single Shot Multibox Detector. In: Leibe, B., Matas, J., Sebe, N. and Welling, M., Eds., *Lecture Notes in Computer Science*, Springer International Publishing, 21-37. https://doi.org/10.1007/978-3-319-46448-0_2
https://link.springer.com/chapter/10.1007/978-3-319-46448-0_2
- [8] Redmon, J., Divvala, S., Girshick, R. and Farhadi, A. (2016) You Only Look Once: Unified, Real-Time Object Detection. 2016 *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, Las Vegas, 27-30 June 2016, 779-788. <https://doi.org/10.1109/cvpr.2016.91>
<https://ieeexplore.ieee.org/document/7780460>
- [9] Lei, M., Li, S., Wu, Y., Hu, H., Zhou, Y., Zheng, X., Ding, G., Du, S., Wu, Z. and Gao, Y. (2025) YOLOv13: Real-Time Object Detection with Hypergraph-Enhanced Adaptive Visual Perception. arXiv: 2506.17733.
- [10] Hou, Q., Zhou, D. and Feng, J. (2021) Coordinate Attention for Efficient Mobile Network Design. 2021 *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, Nashville, 20-25 June 2021, 13708-13717. <https://doi.org/10.1109/cvpr46437.2021.01350>