

队伍编号	MCB2302586
题号	A

基于 DeepLab v3+ 的坑洼道路图像语义分割

摘要

本文构建一个基于深度学习的语义分割模型 DeepLab v3+，融合空洞卷积、空间特征金字塔池化等模块，自行构建并标注数据集，采用数据增广方法扩充训练样本以驱动模型训练。

对于问题一，由于传统边缘特征提取算法与目标检测算法对于坑洼道路边缘提取任务均无法适用，故考虑采用**基于深度学习的语义分割方法**来对坑洼道路图像进行语义分割。由于题目仅公布测试数据，本文从公开数据集中提取 **400 张**坑洼道路图像，使用 **Lable me** 工具自行对数据集进行**像素级语义信息**标注。然后，使用**垂直翻转、水平翻转、高斯模糊**等数据增广手段丰富训练数据空间特征多样性，扩充训练样本总数达 1600 张。本文采用基于**编码器-解码器**结构的 DeepLab v3+ 模型对坑洼道路图像执行语义分割任务，使分割后图像与原图像保持尺寸一致。在编码器结构中，使用 **ResNet-101** 作为特征提取骨干网络用于提取坑洼特征。同时，该模型使用**空洞卷积**来扩大感受野范围，抑制语义信息损失，使用**空间特征金字塔池化**模块来提取坑洼道路图像中的**多尺度特征**，赋予模型更强的特征提取能力与泛化能力。在解码器结构中，模型采用**双线性插值法**对特征图进行上采样来扩大特征图尺寸，将输出图像恢复到与原始图像一致的尺寸。为了使 DeepLab v3+ 模型在坑洼道路数据集上更快更稳定地收敛，首先对模型进行**迁移学习预训练**，然后在坑洼道路数据集上进行微调。在训练过程中，选取**随机梯度下降**作为参数优化器，并对每轮次的学习率进行动态调整。本文选取 **PA、CPA、MIoU、FWIoU** 四个指标来评测模型的语义分割性能，选取 **DRN、MobileNet、ResNet-101** 来验证特征提取骨干网络的性能，选取 **U-Net、SegNet、DeepLab v3+** 来验证语义分割模型的性能。最后，DeepLab v3+ 语义分割模型在自主构建的测试数据集上的 AP、CAP、MIoU、FWIoU 分别达到 **0.9563、0.9353、0.8816、0.9177**。

对于问题二，为简化语义分割计算量，首先调用初赛的**基于伪样本生成增广与通道注意力增强的坑洼道路图像分类模型**对坑洼道路图像题目所给数据进行分类，从中分类出坑洼道路图像。然后使用问题一建立的坑 DeepLab v3+ 语义分割模型对筛选出的图像进行语义分割，输出**精确到像素点的语义信息图**。最后根据 RGB 图像通道矩阵中的非零数值统计坑洼像素点个数，从而计算坑洼面积占比。

关键字：边缘检测 计算机视觉 图像分割

目录

一、 问题重述.....	1
1.1 问题背景.....	1
1.2 问题描述.....	1
二、 问题分析.....	1
三、 模型假设.....	2
四、 符号说明.....	2
五、 模型建立与求解.....	2
5.1 问题一：建立坑洼道路语义分割模型.....	2
5.1.1 传统边缘检测算法与目标检测算法的缺陷	3
5.1.2 坑洼道路图像语义分割	5
5.1.3 获取训练数据	9
5.1.4 数据预处理	10
5.1.5 DeepLab v3+	12
5.1.6 迁移学习	16
5.1.7 优化器设置	17
5.1.8 语义分割网络参数优化过程	18
5.1.9 深度学习语义分割实验效果	19
5.1.10 坑洼道路难样本处理效果	22
5.2 问题二：计算坑洼面积占比.....	23
5.2.1 坑洼道路图像分类	23
5.2.2 面积占比计算	24
六、 模型的评价与推广.....	24
6.1 模型优点.....	24
6.2 模型缺点.....	25
6.3 模型推广.....	25
参考文献.....	25
A 附录 坑洼道路图像语义分割效果展示.....	27

B 附录 源代码	29
2.1 问题一源代码	29
2.1.1 训练与测试源代码	29
2.1.2 构建数据集源代码	38
2.1.3 数据增广源代码	46
2.1.4 DeepLab v3+ 网络结构源代码	49
2.1.5 UNet 源代码	54
2.1.6 SegNet 源代码	63
2.1.7 边缘检测与目标检测	71
2.2 问题二源代码	81
2.2.1 坑洼图像筛选源代码	81
2.2.2 填写结果文件源代码	83

一、问题重述

1.1 问题背景

通过数字图像分析技术辨识出存在坑洼的道路，不仅有助于保障交通安全，提高驾驶条件，还可以促进自动驾驶汽车智能化发展。除此之外，这一机器视觉任务对于地质勘探、航天科学和自然灾害等多个领域的研究和应用具有重要意义。

传统的分类算法通常难以识别出坑洼图像复杂多变的特征，因此在特征提取和分类任务上表现有限。近年来，深度学习技术在计算机视觉领域崭露头角，它具备卓越的特征提取和表示能力，可以自动捕捉图像的最关键信息，为解决坑洼道路检测问题提供了新途径。

在坑洼道路检测和识别任务中，研究者通常将道路图像分为两类：正常和坑洼。使用深度学习技术可以有效捕捉坑洼图像中轮廓、纹理和形态等重要特征，并将它们转化为更适合分类的表征，从而显著提升坑洼图像的识别性能。

1.2 问题描述

问题一：为完成道路避障、道路修复等工作，需要更精准的识别坑洼的边缘、面积等特性。请建立模型提取道路中坑洼的边缘。

问题二：根据识别的坑洼边缘，估计坑洼的面积，并将附件 3：test_data.zip 中的图片计算每个图片中坑洼面积占整个图像面积的比例填写到 test_result2.csv 中，如果是正常道路占比填写 0，坑洼道路填写百分比整数部分。

二、问题分析

对于问题一，需要提取图像中的坑洼边缘，传统的边缘检测算法如 Canny 算子、Sobel 算子等需要进行手工设计，且鲁棒性较差，难以准确地从整张图像中检测出坑洼边缘特征。因此，可以考虑使用抗干扰性强且特征提取能力强大的深度学习方法来检测坑洼边缘。然而单纯的目标检测方法只能在图像中检测到坑洼位置，而无法检测到具体的坑洼边缘，因此，应该对坑洼道路数据执行图像分割任务，从而精确地检测坑洼边缘。而基于深度学习的语义分割可以使用标注了语义信息的坑洼道路图像进行驱动训练，进行过训练的模型对图像中的坑洼边缘特征有较为精确的提取能力。

对于问题二，语义分割模型将会对输入图像中的每个像素点进行坑洼图案与正常背景进行二分类，每一类会标注不同的颜色。要计算整张图像中坑洼图案的面积占比，只需统计图像 RGB 通道中相应被标注为坑洼图案的像素个数，其与整张图像像素总数之比即为坑洼图案占比。

三、模型假设

- 假设人工进行的像素级标注均为坑洼图像真实的语义情况。
- 本文已经固定模型参数随机种子, 对比实验时不存在模型初始化参数不一致的情况。

四、符号说明

符号	符号解释
$\mathcal{D}_{train}$	训练数据集
$\mathcal{D}_{test}$	测试数据集
f_{ϕ}	语义分割网络参数
f_C	初赛问题中的分类模型
η	学习率
$\mathbf{x}$	图像
$\mathbf{x}_{seg}$	分割后图像
$S_{pothole}$	坑洼面积占比

五、模型建立与求解

5.1 问题一：建立坑洼道路语义分割模型

在问题一中，为训练语义分割模型，本文收集公开数据并进行语义标注，构造 DeepLab v3+ 语义分割模型来对坑洼道路图像提取坑洼边缘特征。同时，为了验证本文构造模型的良好性能，进行了充分的对比实验。具体建模流程如图 1 所示。

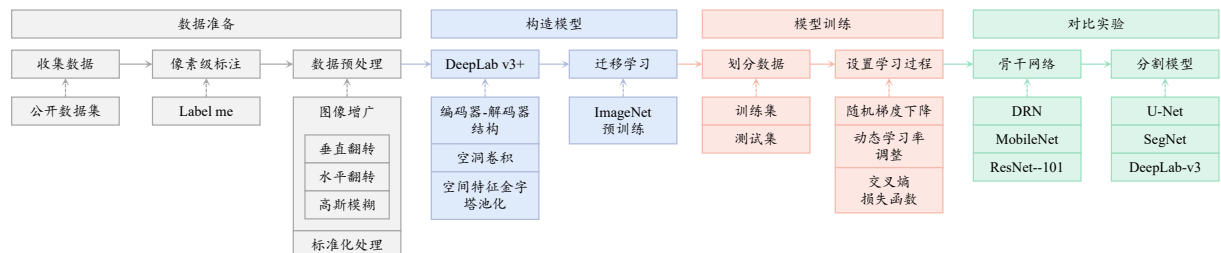


图 1 问题一建模流程

5.1.1 传统边缘检测算法与目标检测算法的缺陷

传统边缘检测算法 边缘检测是计算机视觉和图像处理中的基本任务之一，其目标是识别图像中物体的边界。边缘通常表示图像中灰度级别或颜色发生显著变化的地方，这些变化可能是物体之间的界限或图像中的重要特征。

题目要求准确地从图像中提取出坑洼边缘。然而，传统的基于模板的边缘检测算法（如：Sobel 算子、Canny 算子和 Laplacian 算子等）虽然能够从图像中提取出明显的边缘特征，但这些算子并不能精准地识别出哪些是坑洼边缘，哪些是其他物体的边缘，只能笼统地提取出图像中的所有边缘特征，其检测效果如图 2 所示。

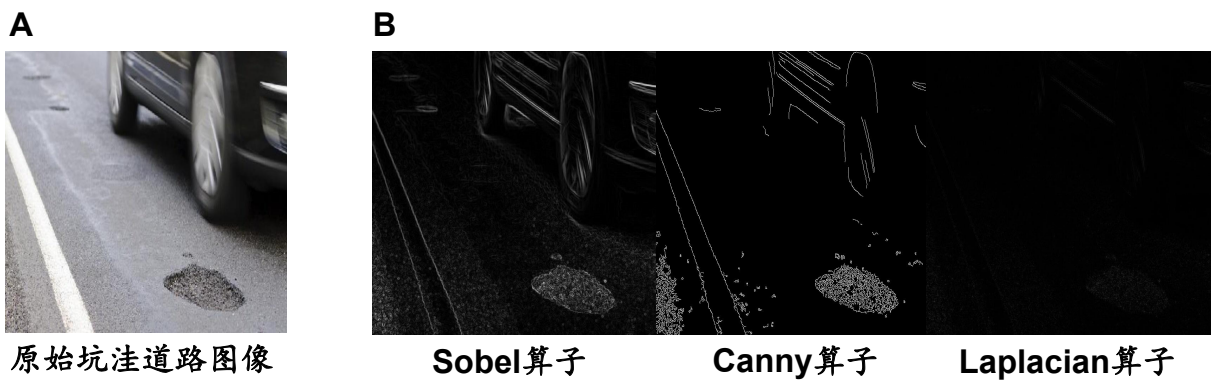


图 2 使用传统边缘检测算法进行坑洼道路图像边缘检测。(A) 原始坑洼道路图像。(B) 传统边缘检测算法提取出的图像边缘，从左到右的提取算法依次为 Canny 算子、Canny 算子、Laplacian 算子。

从图 2 中可以看出，Laplacian 算子提取的边缘非常模糊，Sobel 与 Canny 算子提取的边缘虽然更加清晰，但无法辨别哪些边缘属于坑洼，无法准确地刻画坑洼边缘特征。且当坑洼内部本身纹理特征较为复杂时，传统地边缘检测算法会将坑洼内部地纹理特征看作是边缘特征一同提取，无法做到仅提取边缘特征。

为进一步更加详实地检验传统边缘检测算法对坑洼道路图像中坑洼边缘的提取效果，本文从踢给测试数据中随机抽取 16 张图像，分别使用边缘检测的常用算子 Sobel 算子、Canny 算子、Laplacian 算子进行边缘检测，实验结果如图 3 所示。

从图 3 中可以看出，边缘检测算法只能从图像中提取不同物体间的边缘，无法精确地从众多复杂的边缘细节中提取出具体的坑洼边缘。故对于路况复杂、情况多种多样的坑洼道路图像来说，传统边缘检测算法无法准确地提取出坑洼边缘。

目标检测算法的缺陷 目标检测旨在识别图像中的特定物体并确定其位置，在坑洼道路图像中的具体应用效果如图 4 所示。此处使用开源的坑洼道路检测 YOLO 模型进行目标检测。

使用目标检测算法虽然能够在图像中确定坑洼具体位置，但只能在原始图像的基础上标注出矩形的回归框，与实际的坑洼边缘存在较大误差。故目标检测算法虽然能在图

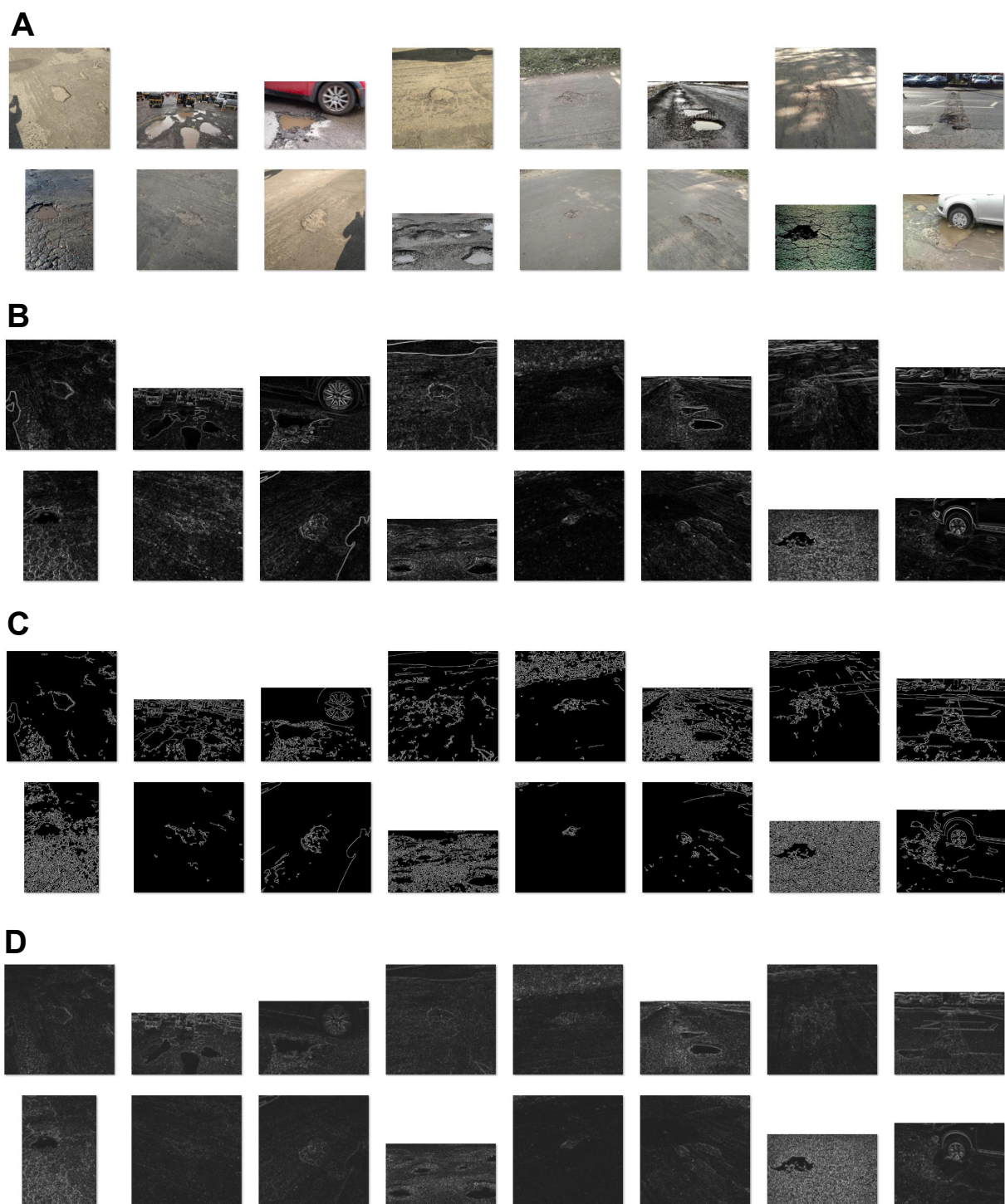


图3 传统边缘检测算法实验效果。(A) 原始坑洼图像。(B) Sobel 算子。(C) Canny 算子。(D) Laplacian 算子。

像中确定坑洼道路位置，大致估算坑洼面积，但面对非规则的复杂坑洼图像时无法准确提取坑洼边缘，只能使用矩形的回归框对坑洼位置进行拟合，计算不够精确。

并且坑洼道路图像场景复杂多样，有许特征空间较为复杂地难样本，如模糊、坑洼数量多、裂痕狭小等，为目标检测任务带来极大困扰。目标检测算法在难样本上的效果如图 5 所示。

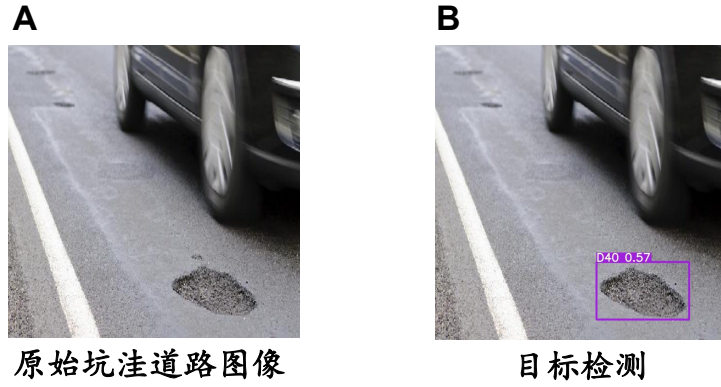


图 4 坑洼道路图像中的目标检测示意。(A) 原始坑洼道路图像。(B) 对图像进行目标检测，检测坑洼位置。

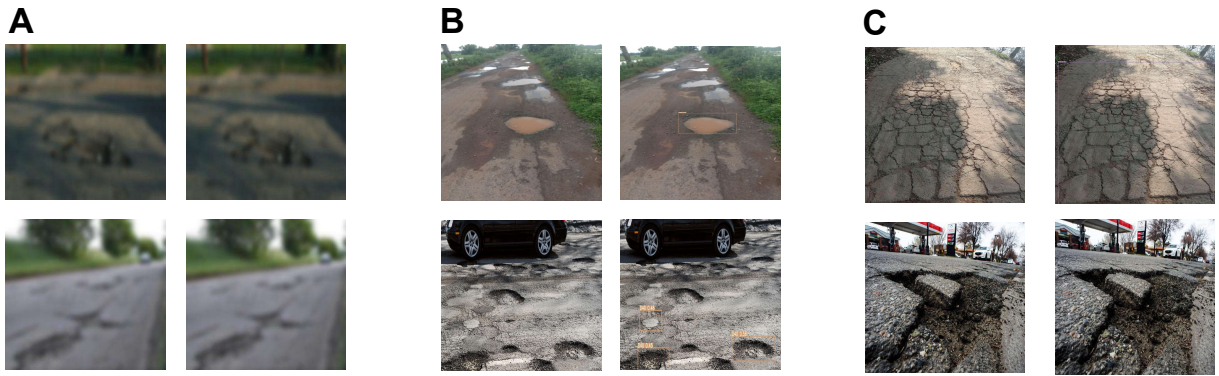


图 5 目标检测算法在坑洼道路图像难样本中的检测效果

可以明显看出，目标检测算法对模糊坑洼图像失去识别能力，对多坑图像仅能检测出图像中少量坑洼，无法识别出图像中所有坑洼，对于坑洼纹理特征较复杂的特征，输出整个回归框的结果误差太大，且对较细小的裂缝型坑洼识别能力有所下降。

5.1.2 坑洼道路图像语义分割

传统边缘检测算法能识别不规则物体边缘，但无法定位到图像中的坑洼，目标检测算法只能使用回归框标定的矩形区域，计算误差过大，无法精确识别纹理复杂、不规则的坑洼边缘。因此，为了准确的识别到路面坑洼区域并提取出坑洼边缘，可使用计算机视觉领域的语义分割算法对坑洼道路图像进行处理。

语义分割的目标是将图像中的每个像素分配到对应的语义类别，从而达到对图像进行像素级别的理解和分类。与目标检测任务不同，语义分割关注的是每个像素的语义，而不仅仅是物体的边界和位置，其基本工作流程框架如式 (1) 所示。

$$\mathbf{X}_{region} \longrightarrow M_s \Rightarrow \mathbf{X}_{mark} \quad (1)$$

其中 $\mathbf{X}_{region}$ 为语义分割模型的输入，表示原始图像； M_s 表示语义分割模型； $\mathbf{X}_{mark}$

为语义分割模型的输出，表示进行过像素级标记的图像。因此，与目标检测相比，对坑洼道路图像进行语义分割，不仅能识别出图像当中坑洼的位置，而且能更加精确的在图像中刻画坑洼的不规则边缘。

对坑洼道路进行语义分割的效果示意如图 6 所示。

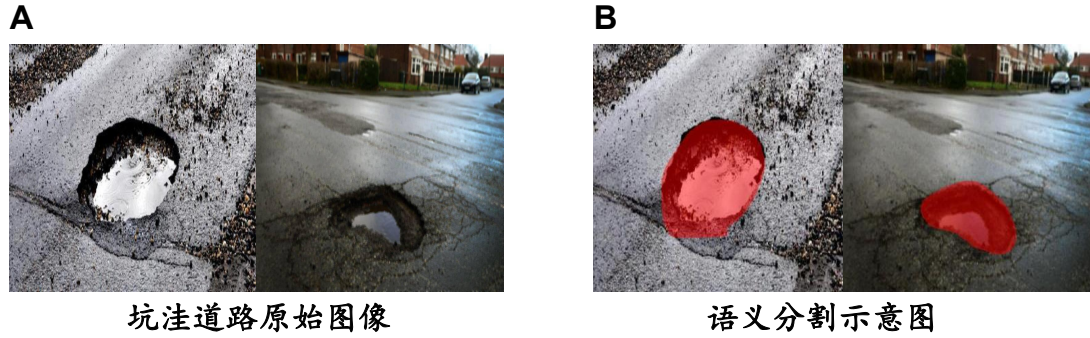


图 6 坑洼道路语义分割示意图。(A) 原始坑洼道路图像。(B) 对坑洼道路图像进行语义分割

通过观察图 6，可以明显看出语义分割技术能够精准的检测图像中不规则坑洼的形状边缘。与仅能输出矩形框的目标检测算法相比，语义分割在边缘检测地准确率上实现了显著的提升。

语义分割之所以能够从坑洼道路图像中准确地分割出完整的坑洼图案，得益于语义分割对图像执行像素级的分类，具体过程如算法 1 所示。

Algorithm 1: 针对坑洼道路图像的语义分割

Data: 原始坑洼道路图像数据集 $\mathcal{D}(X)$

Result: 语义分割后的坑洼道路图像数据集 $\mathcal{D}'(X)$

```

1 // 从数据集选取训练样本;
2 for  $i = 1$  to  $N(\mathcal{D})$  do
3   //  $\mathbf{x}_i$  代表  $\mathcal{D}(X)$  中的第  $i$  张图像;
4   for  $p_{xy}$  to  $\mathbf{x}_i$  do
5     对  $p_{xy}$  进行坑洼区域与正常区域的二分类;
6     if  $p_{xy} \in \text{坑洼区域}$  then
7       对  $p_{xy}$  进行特殊颜色标记;
8        $p_{xy}^{mark}$  替换  $\mathbf{x}_i$  中的  $p_{xy}$ ;
9    $\mathbf{x}_i \Rightarrow \mathbf{x}_i^{mark}$ ;
10   $\mathbf{x}_i^{mark} \rightarrow \mathcal{D}'(X)$ ;
11 return 语义分割后的坑洼道路图像数据集  $\mathcal{D}'(X)$ 

```

语义分割中特征提取网络的主要结构 语义分割可看作是对图像中各像素点的分类任务，因此，语义分割模型也是建立在分类模型基础上的，主要利用卷积神经网络来提取特征并进行分类。语义分割任务中用于特征提取语义信息的主要结构为卷积层、池化层、激活函数，如图 7 所示。

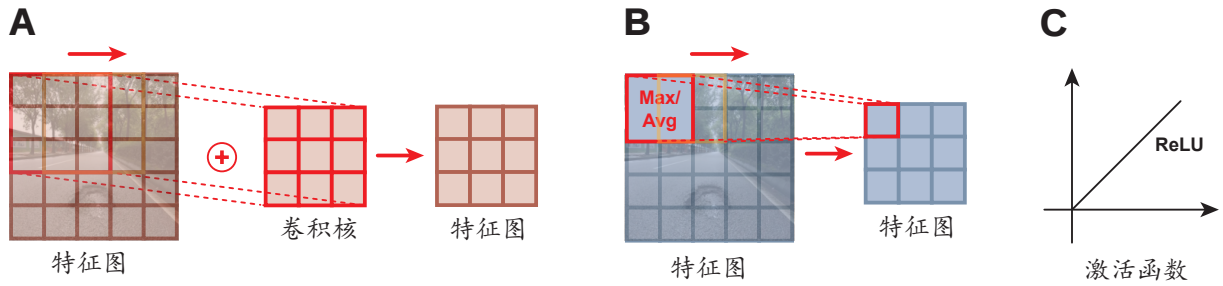


图 7 卷积神经网络工作示意图。(A) 卷积层工作原理；(B) 池化层工作原理；(C) 激活函数工作原理。

- **卷积层：**卷积神经网络的卷积层可以接收三维数据的输入，包括宽度、高度和 RGB 通道等信息，并且同样以的三维数据进行输出，有助于保留图像的形状特征和局部结构，能更好地提取图像的纹理特征。卷积层的工作原理如图 7A 所示，使用卷积操作来处理输入图像，如式 (2) 所示。其中 $\mathbf{W}$ 与 $\mathbf{X}$ 分别表示卷积核与完整输入图像， $*$ 表示卷积运算。

$$\mathbf{C} = \mathbf{W} * \mathbf{X} \quad (2)$$

具体过程如下，首先一个被称为卷积核的滤波器以一定的步长在输入图像上进行滑动，然后在每个像素点上，卷积核与输入图像的对应卷积核与输入图像的对应区域执行元素相乘并求和，产生一个单个的输出值。当卷积核滑动完整张输入图像后，得到的特征值依然是一个矩阵，这个输出矩阵便是卷积层提取出的特征图。具体的卷积操作如式 3 所示。

$$\mathbf{C}_{ij} = \sum_{u=1}^U \sum_{v=1}^V \mathbf{W}_{uv} \mathbf{X}_{i-u+1, j-v+1} \quad (3)$$

其中图像 $\mathbf{X} \in \mathbb{R}^{M \times N}$ ，卷积核 $\mathbf{W} \in \mathbb{R}^{U \times V}$ ， $\mathbf{W}_{uv}$ 表示卷积核矩阵中的数值， $\mathbf{X}_{i-u+1, j-v+1}$ 表示与卷积核进行卷积运算的原图像素点， $\mathbf{C}_{ij}$ 表示经过卷积操作后得到的特征图中的像素点。

- **池化层：**池化层的作用是从输入的特征图上进行特征选择，去除特征图中的冗余信息，简化与特征图参数有关运算，实现特征降维，防止出现过拟合现象。如图 7B 所示，对于一张图像 $\mathbf{X} \in \mathbb{R}^{M \times N}$ ，可以分成若干个小块 $\mathbf{C}^i \in \mathbb{R}^{m \times n}$ ， $m < M, n < N$ 。然后在该区域中进行池化操作，选择出一个特征值来代表整个区域。对于图像 $\mathbf{X}$ 中的第 i 个小块区域 $\mathbf{C}_{m,n}^i$ 而言，常见的池化方法有最大池化与平均池

化。最大池化如式 (4) 所示，即选取区域内的最大值来代表整个区域。

$$\mathbf{Y}_{m,n}^i = \max_{j \in \mathbf{C}_{m,n}^i} x_j \quad (4)$$

平均池化如式 (5) 所示，即计算 $\mathbf{C}_{m,n}^i$ 内的平均特征值来代表整个区域。

$$\mathbf{Y}_{m,n}^i = \frac{1}{m \times n} \sum_{j \in \mathbf{C}_{m,n}^i} x_j \quad (5)$$

- **激活函数：**激活函数是卷积神经网络中的非线性变换，能让网络拥有捕捉更复杂特征的能力。激活函数同时能够通过抑制一部分神经元的活性来缓解网络训练中出现的梯度消失问题，也能学习到更加稳定的特征，提高网络的泛化能力。图 7C 所示的是卷积神经网络中最常用的激活函数 ReLU，其表达式如式 (6) 所示。

$$\text{ReLU}(x) = \begin{cases} x & x \geq 0 \\ 0 & x < 0 \end{cases} \quad (6)$$

相比其他复杂的激活函数，ReLU 函数只需要进行简单的比较与赋值操作，计算上更加高效，且具有良好的稀疏性质。

语义分割网络的特性 普通分类任务是基于图像全局信息的处理，只需最终预测一个全局概率，因此特征图维度降低，但是语义信息更高级。而语义分割需要的是精确到像素点级的处理，过小的特征图会损失很多信息。两种任务要求的特征向量处理如图 8 所示。

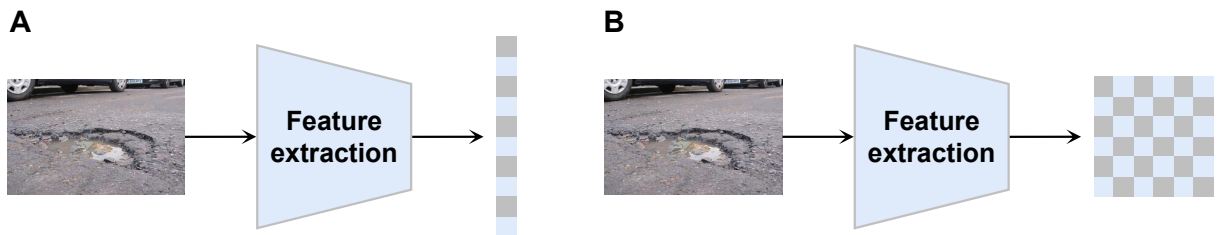


图 8 分类任务与语义分割任务要求的特征向量形式不同。(A) 分类任务需要输出维度小语义信息高维的特征向量。(B) 语义分割任务需要精准到像素级的语义信息，特征向量尺寸尽可能大。

如果在语义分割任务中对图像进行与分类任务相似的降维处理，将导致严重的语义信息损失，大大降低图像中语义分割的精度。但是如果一直让网络总的特征图保持原始大小，将会产生非常大的计算量，无法训练深层次的网络。而降维后的特征图虽然损失了语义信息，但同时也降低了计算量，使得深度网络的训练更加可行。因此，对于语义分割任务来说，平衡这两点尤为关键。

用于处理语义分割的全卷积神经网络 (Fully Convolution Network, FCN) [3] 将图像降维成 32×32 的特征图，虽然使特征图含有丰富的语义但是空间信息损失严重导致分

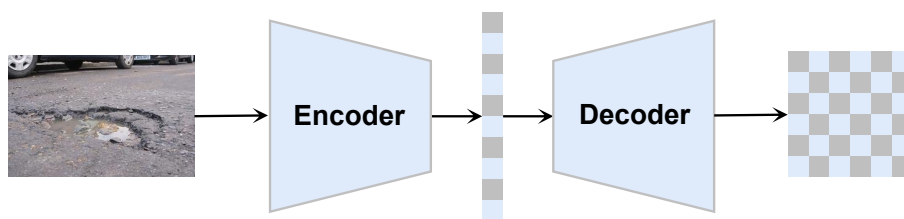


图 9 语义分割中的编码器-解码器结构

割不准确。语义分割的另一种解决方案是采用编码器-解码器（Encoder-Decoder）结构，如图 9 所示。

其中编码器是下采样模块，负责从原始图像中提取高维语义信息，输出一个形状尺寸较小的特征向量。解码器复杂将较小的特征向量恢复成原始图像大小，避免特征降维所引起的语义信息损失。

本文使用基于编码器-解码器结构的 DeepLab v3+ 模型对坑洼道路图像进行语义分割。

5.1.3 获取训练数据

获取原始坑洼道路图像 题目要求对 4942 张坑洼图像进行边缘检测，然而对于语义分割任务而言，要训练图像分割模型，需要对应原始图像的语义标记图像，如图 10 所示。



图 10 坑洼道路语义分割中的原始图像与语义标签

而题中仅给出 4942 张测试数据，且缺少语义标签，为使模型训练结果公平公正，体现模型在测试数据上的泛化能力，本文将提给 4942 张图像全部作为测试数据 $\mathcal{D}_{test}$ ，完全不接触训练过程。本文从来自阿里云天池实验室的坑洼道路公开数据集“Pothole Detection 坑洼检测”中选取样本作为训练集 $\mathcal{D}_{train}$ 。

阿里云天池实验室“Pothole Detection 坑洼检测”是用于坑洼道路目标检测任务的公开数据集，共包含 665 张坑洼道路图像，并且每张图像带有 PASCAL VOC 格式的边

界框注释标签。然而对于语义分割任务而言，需要带有像素级语义信息的标签才能驱动模型训练，故本文从公开数据集中选取 400 张坑洼道路图像，并自行进行数据标注。选取的数据集中部分图像如图 11 所示。



图 11 坑洼道路公开数据集中的部分图像示例

标注数据 为了对原始坑洼道路图像构造语义信息标签，本文使用 Labelme 工具对选取的 400 张坑洼道路图像进行语义信息标注。Labelme 是开源的图像标注工具，常用做计算机视觉领域的目标检测、语义分割等任务的图像标注。使用 Labelme 为坑洼道路图像进行的语义标注如图 12 所示。

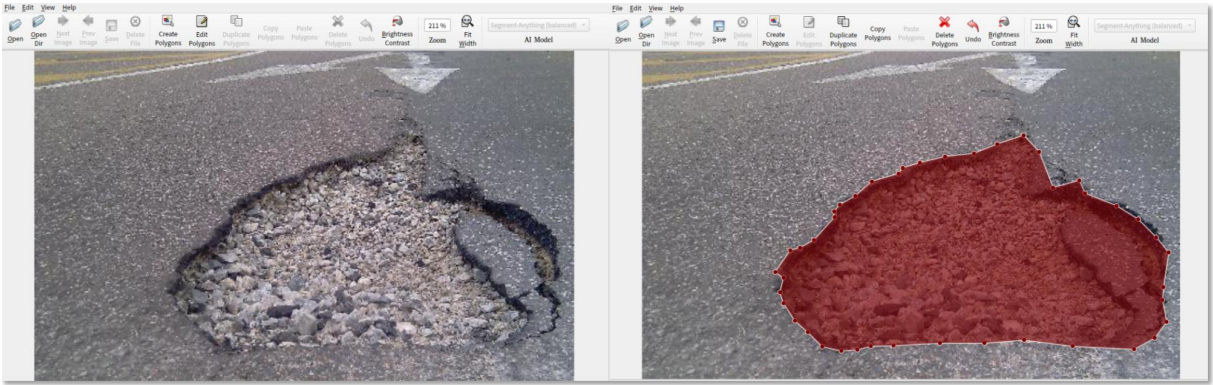


图 12 Labelme 语义标注示意

5.1.4 数据预处理

在坑洼道路图像语义分割任务中，适当地对坑洼道路图像数据进行预处理操作有助于提高模型的性能和鲁棒性，增强语义分割模型对复杂路面场景和变化的适应能力。在本文中，由于自行标注的数据集样本总数较少，故使用水平翻转、垂直翻转、高斯模糊等数据增广手段来扩充训练样本数量，并对所有训练样本进行标准化处理以提高模型稳定性。

数据增广 本文使用水平翻转、垂直翻转、高斯模糊三种数据增广方法来扩充原有数据样本，原有数据样本总数为 400 张，每种数据增广方法在所有原始样本进行操作处理，故每种增广手段扩充 400 张图像，增广后数据总数为 1600 张图像。三种数据增广方法效果如图 13 所示。

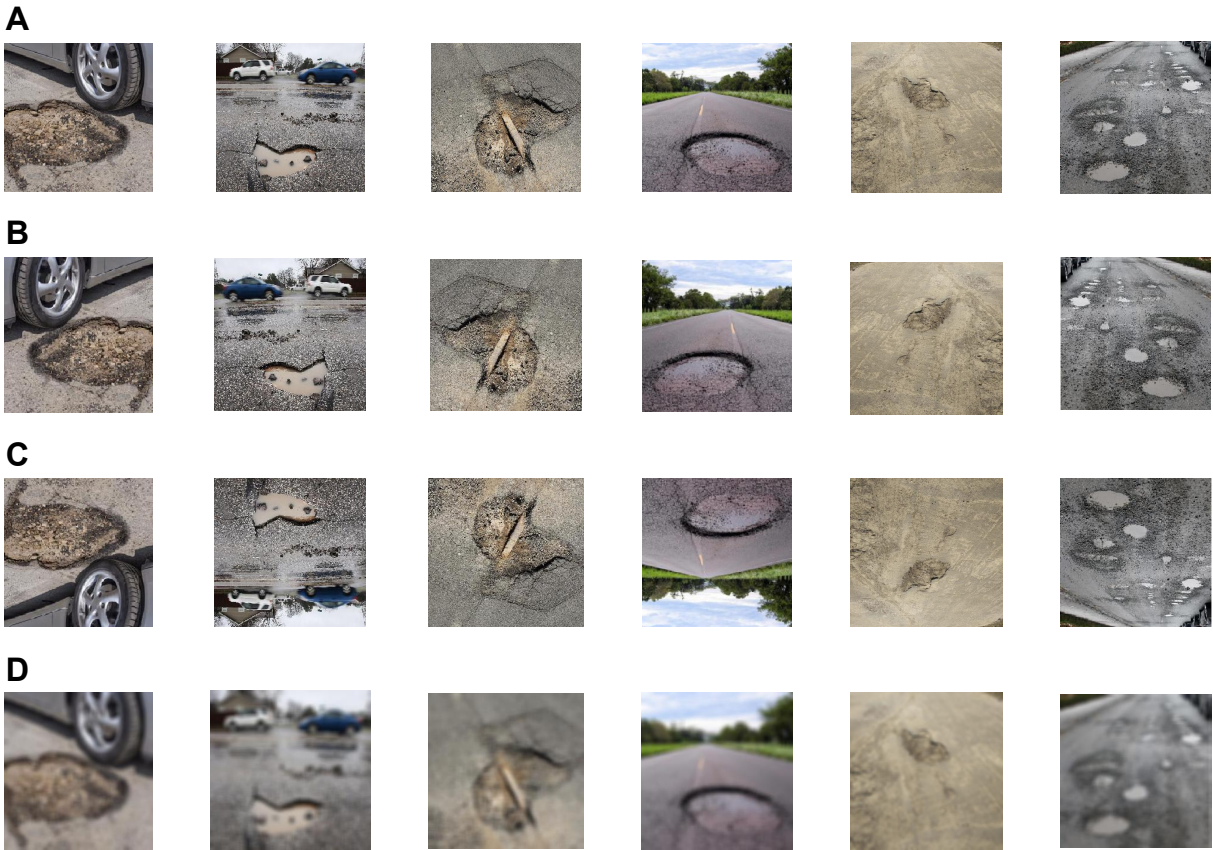


图 13 坑洼道路图像数据增广。(A) 原始图像。(B) 水平翻转。(C) 垂直翻转。(D) 高斯模糊。

对图像进行翻转可以增加数据集中的空间特征复杂性，模型在学习过程中将会见到水平和垂直方向上的不同特征，使模型对目标物体在图像中的位置变化不敏感，让模型更能适应实际场景中物体的不同姿态和方向变化，进一步帮助模型在不同的空间位置中准确地分割出坑洼图案。

考虑到坑洼道路图像获取时的成像条件不同，模型需要面临不同清晰度的坑洼道路图像，因此，本文使用高斯模糊来处理图像，增广清晰度较低的图像，以强化模型对清晰度不理想坑洼道路的分割能力。

高斯模糊通过应用高斯函数对图像进行卷积，使图像中像素值的梯度减小来使图像变模糊。其基本原理是在图像的每个像素上应用一个高斯核，将该像素及其周围的像素按照高斯分布的权重进行加权平均。高斯核的尺寸和标准差决定了模糊的程度，尺寸越大、标准差越小，模糊效果越显著。在离散形式下，高斯模糊的具体数学表达如式 (7) 所示。

$$G(i, j) = \frac{1}{2\pi\sigma^2} \exp\left(-\frac{i^2 + j^2}{2\sigma^2}\right) \quad (7)$$

设 K 是一个尺寸为 $(2k + 1) \times (2k + 1)$ 的离散高斯核, k 是一个非负整数, 那么 i 和 j 分别表示 K 的行和列索引, σ 是高斯分布的标准差。那么对原始图像 $\mathbf{x}$ 的高斯模糊卷积操作如式 (8) 所示。

$$\mathbf{x}_{\text{blurred}}(x, y) = \sum_{i=-k}^k \sum_{j=-k}^k \mathbf{x}(x + i, y + j) \cdot G(i, j) \quad (8)$$

其中 $\mathbf{x}_{\text{blurred}}$ 为高斯模糊后的图像。

标准化处理 标准化是深度学习中常用的预处理技术, 将输入数据归一到一个较小的范围, 可以减少梯度消失或梯度爆炸等问题, 有助于模型更好地收敛, 更容易找到全局最优解。

要对坑洼道路数据集进行标准化处理, 首先计算数据集整体的均值 $mean$ 和方差 std 。RGB 图像中的每个通道 c , $mean$ 的计算如式 (9) 所示。其中 H 和 W 分别是图像的宽和高, $\mathbf{x}_c(i, j)$ 是图像在位置 (i, j) 处通道 c 的像素值。

$$mean_c = \frac{1}{H \times W} \sum_{i=1}^H \sum_{j=1}^W \mathbf{x}_c(i, j) \quad (9)$$

对每个通道 c , 标准差 std 的计算如式 (10) 所示。

$$std_c = \sqrt{\frac{1}{H \times W} \sum_{i=1}^H \sum_{j=1}^W (\mathbf{x}_c(i, j) - mean_c)^2} \quad (10)$$

上述计算在整个数据集集中的图像进行, 对数据集中所有图像的所有通道进行累加求均值和标准差, 求得数据集中的均值和标准差具体数值后, 对数据集中所有数据进行如式 (11) 所示的标准化处理。

$$\mathbf{x}_{\text{Normalized}}(i, j) = \frac{\mathbf{x}(i, j) - mean}{std} \quad (11)$$

其中 $\mathbf{x}_{\text{Normalized}}$ 为标准化处理后的图像, $\mathbf{x}_{\text{Normalized}}(i, j)$ 为 $\mathbf{x}_{\text{Normalized}}$ 中位于 (i, j) 位置的像素值。

5.1.5 DeepLab v3+

建模动机 大多数用于语义分割的深度神经网络面临着两大问题。一是物体多尺度问题, 语义分割任务中数据图像内的物体通常具有不同的尺度, 而模型需要能够有效地识

别和分割这些不同尺度地物体。其次，深度神经网络的多次下采样会造成特征图分辨率变小，语义信息损失，从而导致预测精度降低，边界信息丢失。

DeepLab v3+ 通过引入空洞卷积、ASPP 模块等解决了上述问题，同时进一步优化了大尺度图像的计算量，使得语义分割任务速率高且分割精准。

网络结构 为了解决网络计算耗时过长、语义信息不够充分等问题，DeepLab v3+ 引入编码器-解码器结构。解码器进行特征提取，编码器通过上采样操作将特征图还原为低层次特征的分辨率大小。其具体网络结构如图 14 所示。

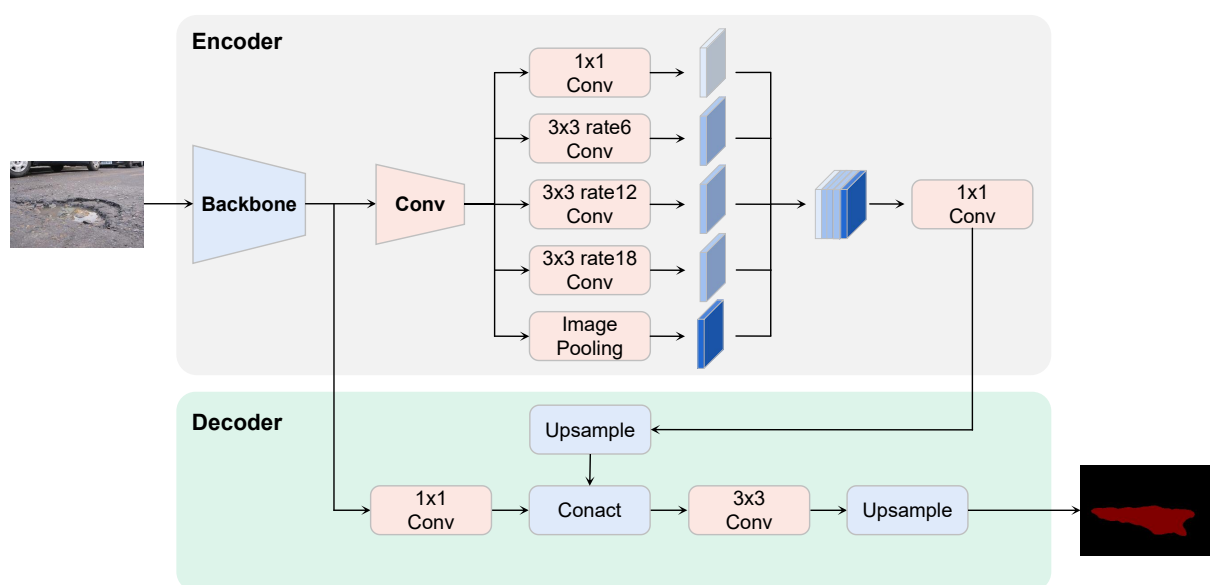


图 14 对坑洼道路执行语义分割任务的 DeepLab v3+ 总体结构。

在图 14 中，Encoder 表示提取图像特征的编码器结构，Decoder 表示还原特征图大小的解码器结构，Backbone 表示 DeepLab v3+ 中的特征提取骨干网络。

首先，将图像输入到带有空洞卷积和深度可分离卷积的骨干网络模型中提取低层次特征；然后，将图像的低层次特征输入到空间金字塔池化模块（Atrous Spatial Pyramid Pooling, ASPP）中进一步提取更深层次的多尺度特征，并使用尺寸为 1×1 的卷积层将多尺度特征融合成高层次特征；最后，对高层次特征进行双线性插值的上采样，并将其与低层次特征进行融合，通过卷积层与上采样操作输出最后的语义分割图像。

空洞卷积 空洞卷积是 DeepLab v3+ 网络中的重要组成部分。我们在进行语义分割时需要在扩大感受野的同时尽可能减少语义信息的丢失，使用空洞卷积能够在扩大卷积层中感受野的同时保持分辨率维持稳定。

在空洞卷积中， r 控制着感受野的大小，在卷积核尺寸不变的情况下， r 越大则感受野越大，具体控制效果如图 15 所示。

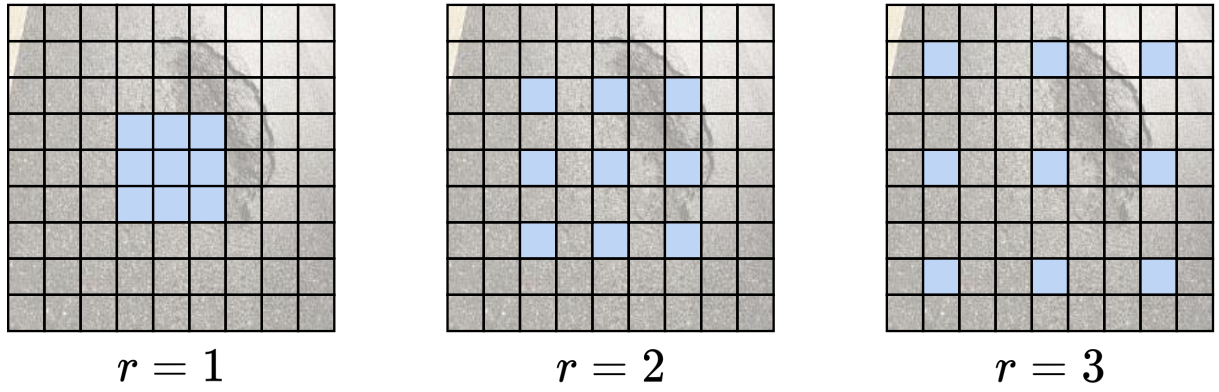


图 15 空洞卷积示意图。从左到右为 r 值逐渐增大的感受野区域。

空洞卷积对特征图位置 i 的感受野运算如式 (12) 所示。

$$\mathbf{Y}[i] = \sum_k \mathbf{x}[i + r * k] \times \mathbf{W}[k] \quad (12)$$

其中 $\mathbf{X}$ 为输入的原始图像， $\mathbf{Y}$ 为空洞卷积层卷积运算后得到的特征图， i 表示特征图上的像素位置， $\mathbf{W}$ 表示卷积核， r 为膨胀卷积中的膨胀系数。空洞卷积操作相当于将输入 X 在每个空间维度上，与两个连续的卷积核之间插入 $r - 1$ 个零值而产生的上采样卷积核进行卷积。

空间金字塔池化 当骨干模型网络提取低层次特征之后，DeepLab v3+ 使用空洞卷积构成的空间金字塔池化模块进一步提取多尺度特征。DeepLab v3+ 中的空间金字塔详细结构如图 16 所示。

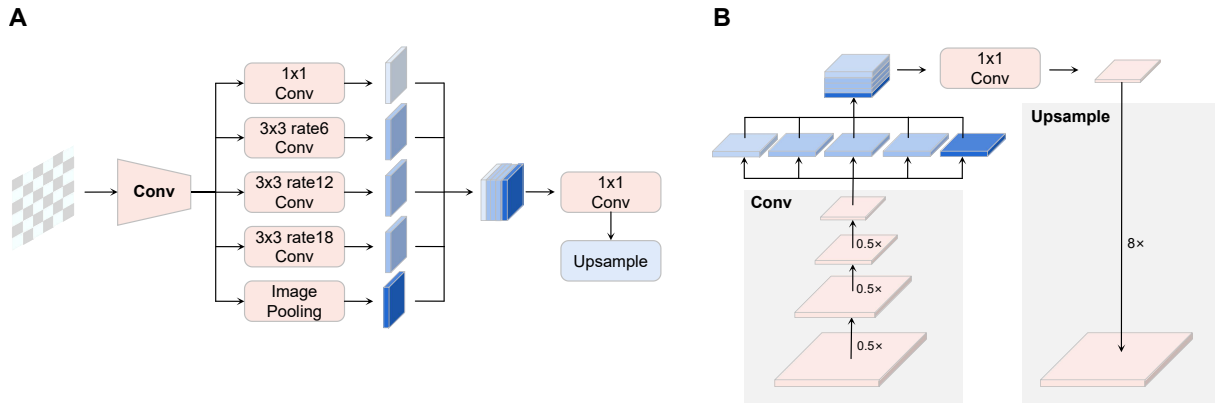


图 16 DeepLab v3+ 中的空间金字塔池化。(A) 空间金字塔网络结构。(B) 特征图在空间金字塔中的变化流程与详细网络模块。

特征图输入到空间金字塔中，首先经过三个尺寸为 3×3 且 r 不一样的空洞卷积层一个尺寸为 1×1 的卷积层，提取出特征图中的不同尺度特征，再使用一个全局平均池化层对原特征图进行处理，将处理后的特征图与之前的四张不同尺度特征图堆叠在一

起形成多尺度特征，最后将多尺度特征输入到 1×1 中的卷积层得到新的融合特征，最后对融合特征进行八倍的上采样，得到与原始特征图大小相同的特征图。

在坑洼道路的边缘检测任务中，受到拍摄角度、图像尺寸大小等因素的影响，图像中的坑洼图案也呈现出不一样的大小和尺寸，因此，在语义分割网络中加入空间金字塔模块有助于挖掘不同图像当中具有不同尺度特征的坑洼边缘特征。

解码器结构 解码器结构如图 14 下方所示。首先，上采样层对编码器输出的多尺度融合特征进行四倍的双线性插值上采样。双线性插值是一种用于在二维空间中找到非整数坐标上的像素插值方法，通常用于图像放大，以及改善图像质量使图像在变换过程中变得更加平滑。

双线性插值方法具体过程如下：考虑二维空间上的四个相邻点 Q_{11} , Q_{12} , Q_{21} , Q_{22} ，和目标点 (x, y) ，首先计算在 x 轴方向上的两个插值，分别位于 Q_{11} 和 Q_{21} 之间以及 Q_{12} 和 Q_{22} 之间，插值可通过如式 (13) 所示的线性插公式得到。

$$\begin{cases} R_1 = Q_{11} + (x_2 - x) \cdot \frac{Q_{21} - Q_{11}}{x_2 - x_1} \\ R_2 = Q_{12} + (x_2 - x) \cdot \frac{Q_{22} - Q_{12}}{x_2 - x_1} \end{cases} \quad (13)$$

再使用水平插值得到的 R_1 和 R_2 进行垂直插值，如式 (14) 所示。

$$P = R_1 + (y_2 - y) \cdot \frac{R_2 - R_1}{y_2 - y_1} \quad (14)$$

P 即为目标点 (x, y) 处的插值结果。

多尺度融合特征经过插值放大后，与来自骨干特征提取网络的低层次特征进行连接，低层次特征进行融合之前还会经过一个卷积核尺寸为 1×1 的卷积层以减少通道数量。因为低层次特征包含大量的特征通道，容易使训练变得非常困难。特征融合后，使用一个 3×3 的卷积来对融合后的特征进一步细化，并再次进行四倍的双线性插值上采样，使语义分割后的图像大小与原始输入图像大小一致。

特征提取骨干网络 本文选取特征提取任务中的经典网络残差神经网络 ResNet 作为 DeepLab v3+ 中的特征提取骨干网络模型。同时，为了挖掘图像中坑洼部分的深层特征，使用 ResNet 中层数较深的 ResNet-101 对坑洼图像执行特征提取任务。

残差神经网络 (ResNet) [2] 创造性地引入了残差模块与跳跃连接，极大程度上改善了深层网络的梯度消失与梯度爆炸问题，让神经网络的层数能够被设计地越来越深，从而提取更加复杂的语义特征。残差块与跳跃连接的结构如图 17 所示。

残差块 ResNet 通过引入残差快结构，将网络层之间的输入信号与输出信号之间的差异作为网络的学习目标，而不是直接学习完整的输出信号，学习方式如式 (15) 所示。

$$F(x) = H(x) - x \quad (15)$$

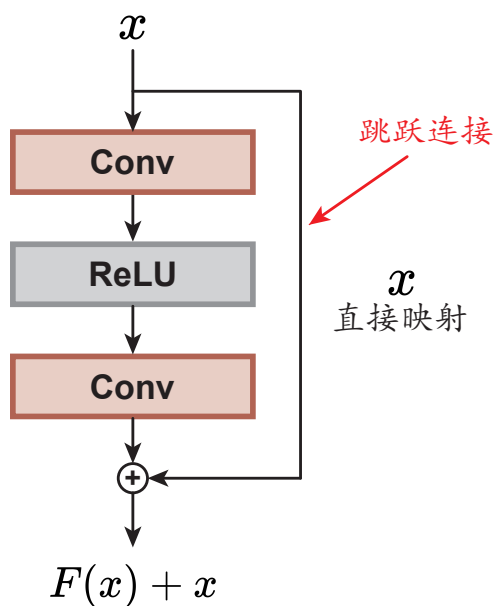


图 17 残差块与跳跃连接结构图

其中 x 为原始输入， $H(x)$ 为网络的输出映射， $F(x)$ 即为输入与输出映射之间的残差，是网络真正学习的内容。这使得网络更加容易训练，因为原始的输入信息被保留了下来，网络只需要学习输入到输出的变化，而不是学习从输入到输出的完整映射，很好地解决了梯度消失问题。

跳跃连接 ResNet 中的残差块还引入了跳跃连接，允许信号在模型的不同层之间传递，底层的信息可以直接传递到输出层，网络可以更轻松地学习恒等映射，而不会受到复杂的层次结构的阻碍。这有助于梯度的有效传播，降低了训练深度网络的难度。

5.1.6 迁移学习

本文为了更加准确地提取不规则坑洼边缘的深度特征，选取 ResNet 网络系列中较深的网络结构 ResNet-101 作为特征提取骨干网络。然而，训练一个深层次的网络模型需要大量的标签数据才能使其具有良好的泛化能力，否则网络极易陷入过拟合。在本文中，可用于训练的标签数据量较少，如果对 DeepLab v3+ 的所有网络参数进行随机初始化，然后使用较少的坑洼道路标签数据进行训练，那么模型容易对训练样本中的坑洼道路陷入过拟合，其语义分割能力无法泛化到特征分布更加复杂的测试样本上。

迁移学习 (Transfer Learning) 是一种机器学习技术，将在一个任务上学习的知识转移到不同但相关任务的模型上。先前被学习的数据成为源域，将要执行当任务的数据被称为目标域。迁移学习的目标是在源域上学习一个模型，这个模型被称为预训练模型，然后将此预训练模型应用在目标域上。

为了缓解训练数据过少带来的过拟合等问题，本文使用迁移学习将 ResNet-101 在

ImageNet 上的预训练权重加载到 DeepLab v3+ 内的特征提取骨干网络中，让网络参数从一开始就具备强大的特征提取能力。然后再使用坑洼道路的坑洼边缘标签数据微调 DeepLab v3+ 的参数，让网络具备对坑洼道路的语义分割能力。迁移学习过程如图 18 所示。

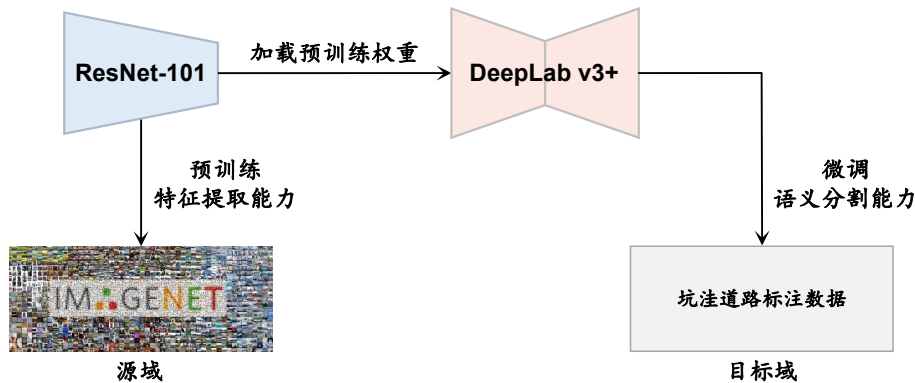


图 18 坑洼道路语义分割中的迁移学习

ImageNet 是一个大规模的图像数据库，是计算机视觉领域中的一个重要资源。它包含来自世界各地的数百万张图像，这些图像按照 WordNet 层次结构进行了标注，共涵盖数千个对象类别，是目前人工智能领域规模最大的数据集，许多经典的人工神经网络结构都有在 ImageNet 上的训练权重。

由于本文中目标域数据量较少，迁移学习可以利用源任务的大量数据，使得模型更容易泛化到目标任务。同时，预训练模型在大规模数据上进行了训练，具有更好的泛化能力。通过迁移学习，可以将这种泛化能力转移到目标任务上。通过使用坑洼道路语义标记数据进行微调，使得迁移学习后的 DeepLab v3+ 模型能更好地适应坑洼道路语义分割数据。

5.1.7 优化器设置

深度学习模型使用参数优化器来更新神经网络中的参数，使用学习率优化器来更新每一轮训练时的学习率大小。本文使用随机梯度下降（SGD）作为参数优化器，使用 PolyScheduler 作为学习率优化器。

随机梯度下降 梯度下降是优化模型参数最常用的方法，其参数更新如式 (16) 所示。其中 θ_t 表示 t 时刻的参数， $\mathcal{L}$ 表示损失函数值， η 表示学习率。

$$\theta_t = \theta_{t-1} - \eta \frac{\partial \mathcal{L}}{\partial \theta_{t-1}} \tag{16}$$

在传统的梯度下降方法中，每一次迭代都要使用整个训练数据集来计算梯度并更新模型参数。而 SGD（Stochastic Gradient Descent）每次迭代只使用训练数据集中的一个随机子集来计算梯度和更新模型，使得算法的计算开销大幅度减小。

由于随机梯度下降每次只使用一小批样本，它的计算效率相对较高，并且在处理大规模数据集时具有一定的优势。然而，由于使用的样本是随机选择的，收敛过程可能会表现得更加嘈杂，但通常能够在相对短的时间内找到局部最优解。

PolyScheduler 传统固定调整学习率的优化算法通常会在训练过程的每个固定的时间步长调整学习率。这种静态学习率优化算法有一些潜在缺陷，特别是在处理复杂的、非均匀分布的数据集时。使用固定的学习率虽然在训练初期能使模型加快收敛速度，但在接近最优值时，学习率调整时的步长太大很可能导致跳过最优值，步长太小又可能使收敛速度变慢。

本文使用动态学习率优化算法 PolyScheduler 来对模型进行学习率更新，如式 (17) 所示。

$$\eta_{epoch} = \eta_{base} \times \left(\frac{epoch}{N_{epoch}} \right)^{power} \quad (17)$$

其中, η_{epoch} 为当前训练轮次 $epoch$ 下的学习率大小, η_{base} 为初始学习率大小, $epoch$ 为当前训练轮次, N_{epoch} 为模型最大训练轮数, $power$ 为超参数, 可用于控制学习率衰减程度。

动态更新学习率可以帮助算法更快地收敛到局部最优解。在训练初期, 使用较大的学习率可以加速参数的更新, 而在接近最优值时, 逐渐减小学习率可以使模型更精确地调整参数, 防止错过最优解。同时, 动态学习率算法可以自适应地调整学习率, 根据训练过程中的模型性能和数据分布的变化来灵活地调整学习率。这种适应性使得算法更具鲁棒性, 可以更好地适应不同的数据和问题。

5.1.8 语义分割网络参数优化过程

DeepLab v3+ 等卷积神经网络使用误差反向传播算法来更新网络参数。在训练过程中, 样本在网络中的输出值将与真实标签一起输入到损失函数中计算误差值, 然后将误差值向前逐层传导, 对每一层的权重求偏导计算梯度, 让损失函数最小化。

语义分割本质为逐像素点的分类问题, 坑洼道路语义分割可看作为逐像素点的坑洼道路与正常道路二分类问题, 对此, 我们采用如式 (18) 所示的二分类交叉熵函数作为 DeepLab v3+ 的损失函数。即对图像中每一个像素点计算分类损失函数。

$$\mathcal{L} = \frac{1}{N} \cdot \frac{1}{H \times W} \sum_i^{H \times W} \sum_j^N - \left[y_i^j \cdot \log(p_i^j) + (1 - y_i^j) \cdot \log(1 - p_i^j) \right] \quad (18)$$

其中 N 表示图像数量, H 与 W 分别为图像宽高, 两者乘积即为像素点个数。 y_i^j 表示样本 y^j 中像素点 i 的标签, p_i^j 表示像素点被预测为正常道路的概率。故 DeepLab v3+ 的训练流程如算法 2 所示。

Algorithm 2: 针对坑洼道路语义分割的 DeepLab v3 + 训练

Data: 训练样本与标注数据 $\mathcal{D}(X, X_{label})$; 参数优化器 Optimizer; 学习率优化器 Scheduler

Result: 训练后的模型参数 f'_ϕ

```
1 初始化网络模型参数  $f_\phi$ , 设置  $\eta, N_{epoch}, N_{batch}$ ;  
2 for  $n = 1$  to  $N_{epoch}$  do  
3   随机打乱训练数据  $\mathcal{D}$ ;  
4   for 一个批次训练样本  $(\mathbf{x}, \mathbf{x}_{label})$  to  $\mathcal{D}$  do  
5     梯度清零: nograd();  
6     前向传播:  $f_\phi(\mathbf{x}) = \hat{\mathbf{x}}$ ;  
7     计算损失函数:  $\mathcal{L} = \text{损失函数}(\hat{\mathbf{x}}, \mathbf{x}_{label})$ ;  
8     反向传播: 计算损失函数关于模型参数的梯度;  
9     // 更新每一层模型参数;  
10     $\mathbf{W}^{(l)} \leftarrow \mathbf{W}^{(l)} - \eta \frac{\partial \mathcal{L}}{\partial \mathbf{W}^{(l)}}$ ;  
11     $b^{(l)} \leftarrow b^{(l)} - \eta \frac{\partial \mathcal{L}}{\partial b^{(l)}}$ ;  
12   更新学习率: Scheduler.step();  
13 return 训练后的模型参数  $f'_\phi$ 
```

5.1.9 深度学习语义分割实验效果

实验设置 本文将增广后的 1600 张图像随机划分为训练集 $\mathcal{D}_{train}$ 与测试集 $\mathcal{D}_{test}$, 为保证实验结果公正可靠, $\mathcal{D}_{train} \cap \mathcal{D}_{test} = \emptyset$ 其中 $\mathcal{D}_{train}$ 包含 1400 张图像, $\mathcal{D}_{test}$ 包含 200 张图像。对于参数优化器 SGD, 设置基学习率为 0.07, 用于加速 SGD 沿相关方向的收敛并减弱震荡的动量设置为 0.9, 权重衰减的正则化项设置为 0.0005, 并启用 Nesterov 动量, 帮助模型更快收敛。所有实验在单张 Tesla A100 显卡上运行。

评价指标 本文选取像素准确率 (Pixel Accuracy, PA), 类别像素准确率 (Class Pixel Accuracy, CPA), 平均交并比 (Mean Intersection over Union, MIoU), 频权交并比 (Frequency Weighted Intersection over Union, FWIoU) 作为度量语义分割性能地指标 [4]。

坑洼道路语义分割问题即为像素水平的二分类问题, 将图像中的每一个像素点分为坑洼与正常两个类别。在二分类问题中有四个常用指标: 真正例 TP (True Postitive)、假正例 FP (False Positives)、假反例 FN (False Negatives)、真反例 TN (True Negatives), 指标含义如下所示。

- **TP:** 表示模型将正类别样本正确地分类为正类别的数量。
- **FP:** 表示模型将负类别样本错误地分类为正类别的数量。

- **FN**: 表示模型将正类别样本错误地分类为负类别的数量。
- **TN**: 表示模型将正类别样本错误地分类为负类别的数量。

对于分割任务总共涉及到的 $k + 1$ 类 (含 k 个前景类和 1 个背景类), 令 p_{ii} 表示正确分类的像素数; p_{ij} 为第 i 类被预测为第 j 类的像素数, 像素准确率 PA 为分割类别正确的像素数占总像素数的比例, 如式 (19) 所示。**像素准确率 PA** 常用于总体分割效果的估计, 越接近 1 表明模型性能越好。

$$PA = \frac{TP + TN}{TP + TN + FN + TN} = \frac{\sum_{i=0}^k p_{ii}}{\sum_{i=0}^k \sum_{j=0}^k p_{ij}} p_{ij} \quad (19)$$

像素精度所包含的信息有限, 容易掩盖某一类别分割效果差的现象, 对此可以使用**类别像素准确率 CPA** 来进行更加精确的评估。对第 i 类的分割结果来说, CPA 计算公式如式 (20) 所示。

$$CPA = \frac{TP}{TP + FP} = \frac{p_{ii}}{\sum_{j=0}^k p_{ij}} \quad (20)$$

交并比 IoU 是当前语义分割研究中最常用的指标, 表示分割掩码与标签像素的交叠率, 如式 (21) 所示。

$$IoU = \frac{|A \cap B|}{|A \cup B|} = \frac{TP}{TP + FP + FN} = \frac{\sum_{i=0}^k p_{ii}}{\sum_{i=0}^k \left(\sum_{j=0}^k p_{ji} + \sum_{j=0}^k (p_{ij} - p_{ii}) \right)} \quad (21)$$

平均交并比 MIoU 是计算每个类别中 IoU 值的算数平均值, 用于总体数据集的像素重叠情况, 计算方式如式 (22) 所示。

$$\begin{aligned} MIoU &= \frac{1}{2} \left[\frac{TP}{TP + FP + FN} + \frac{TN}{TN + FN + FP} \right] \\ &= \frac{1}{k+1} \sum_{i=0}^k \frac{p_{ii}}{\sum_{j=0}^k p_{ij} + \sum_{j=0}^k (p_{ji} - p_{ii})} \end{aligned} \quad (22)$$

频率加权交并比 FWIoU 是对 MIoU 改进后的新评价标准, 需对每个像素的类别按照其出现的频率设置权重, 如式 (23) 所示。

$$\begin{aligned} FWIoU &= \frac{TP + FN}{TP + FP + TN + FN} \times \frac{TP}{TP + FP + FN} \\ &= \frac{1}{\sum_{i=0}^k \sum_{j=0}^k p_{ij}} \sum_{i=0}^k \left(\frac{\sum_{j=0}^k p_{ij} p_{ii}}{\sum_{j=0}^k p_{ij} + \sum_{j=0}^k (p_{ji} - p_{ii})} \right) \end{aligned} \quad (23)$$

特征提取骨干网络对比 为了验证本文特征提取骨干网络的性能, 本文分别选取在 ImageNet 上预训练过的 DRN, MobileNet, ResNet-101 作为 DeepLab-v3 的特征提取网络, SGD 作为参数优化器, PolyScheduler 作为学习率优化器, 在相同的 $\mathcal{D}_{train}$ 上进行 50 轮次的训练, 每轮训练后在标注数据中独立划分出的测试集上分别测试 AP、CAP、MIoU、FWIoU, 如图 19 所示。

从图 19 中可以看出, 在三个特征提取网络中, 由于 MobileNet 的参数规模较小, 性能明显低于另外两个网络。而 ResNet-101 的性能在 AP、CAP、MIoU、FWIoU 四个指标中都高于 DRN。模型在自主标注的测试数据 $\mathcal{D}_{test}$ 中性能数值如表 1 所示。

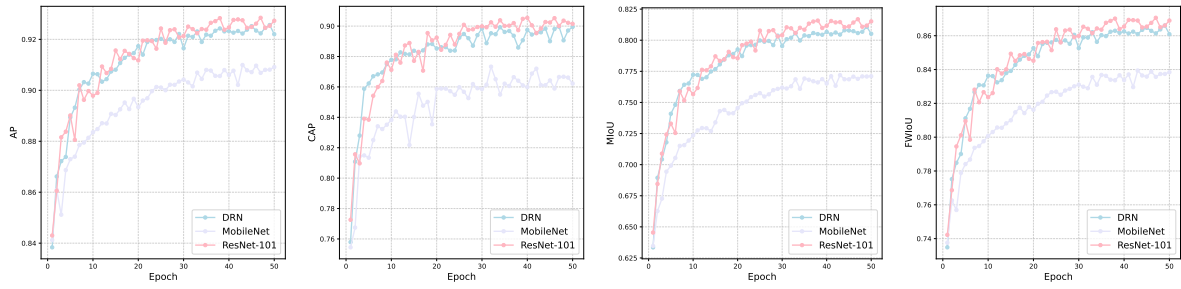


图 19 DeepLab v3+ 中不同特征提取骨干网络的性能

表 1 DeepLab v3+ 中不同特征提取骨干网络在自主标注测试集中的性能

特征提取骨干网络	AP	CAP	MIoU	FWIoU
DeepLab v3+ (DRN)	0.9220	0.8994	0.8051	0.8608
DeepLab v3+ (MobileNet)	0.9091	0.8622	0.7710	0.8383
DeepLab v3+ (ResNet-101)	0.9272	0.9014	0.8151	0.8689

语义分割模型对比 为了验证本文所选用语义分割模型 DeepLab v3+ 的性能，本文分别选取 U-Net、SegNet、DeepLab v3+ 两种网络，SGD 作为参数优化器，PolyScheduler 作为学习率优化器，在相同的 $\mathcal{D}_{train}$ 上进行 50 轮次的训练，每轮训练后在标注数据中独立划分出的测试集 $\mathcal{D}_{test}$ 上分别测试 AP、CAP、MIoU、FWIoU，如图 20 所示。

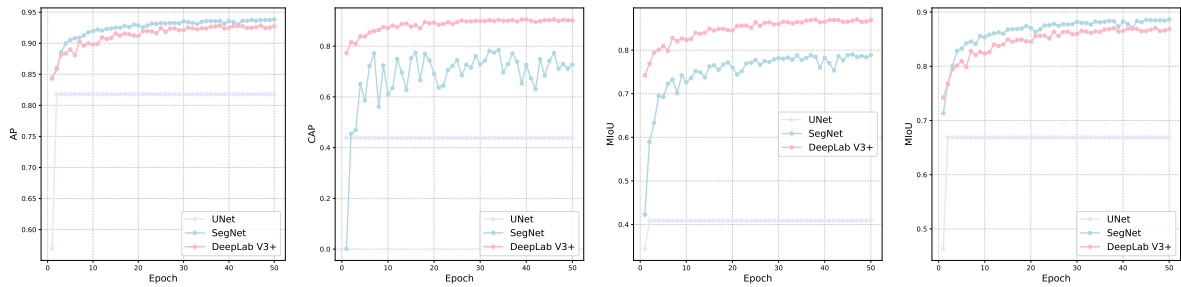


图 20 语义分割模型性能对比

从图 20 中可以在坑洼道路图像语义分割任务中，看出 UNet 比 DeepLab v3+ 模型性能差距较远，且模型仅经过一轮训练后进入收敛状态。UNet 在训练过程中，CAP 持续为零，说明模型已经对坑洼语义信息失去了识别能力，将图像中所有像素均认为是正常道路图案。因为 UNet 模型的层次结构较为简单，不具备提取类似于坑洼图案的复杂高维语义特征，模型在使用损失函数进行训练时，为保证损失的最小化而直接放弃了对坑洼图案的识别。

而对于 SegNet，虽然在 AP 与 FWIoU 上比 DeepLab v3+ 略高，但观察到 CAP 与 MIoU 中仅有 0.7274 与 0.7888，远低于 0.9014 与 0.8151，说明其对坑洼像素点的分割能

力依然远低于 ResNet-101，其他两项指标高是因为坑洼图案占图像总体相对较少。实验的具体数值如表 2 所示。

表 2 不同语义分割模型在自主标注测试集中的性能

语义分割模型	AP	CAP	MIoU	FWIoU
UNet	0.8177	0.0000	0.4088	0.6687
SegNet	0.9384	0.7274	0.7888	0.8865
DeepLab v3+	0.9272	0.9014	0.8151	0.8689

模型最终训练 基于以上对比实验结果，本文选取 DeepLab v3+ 作为语义分割模型，ResNet-101 作为特征提取骨干网络，在 $\mathcal{D}_{train}$ 上进行 150 轮次的训练，每一轮训练后在 $\mathcal{D}_{test}$ 上进行测试，采用 AP、CAP、MIoU、FWIoU 四个指标进行性能验证，结果如图 21 所示。

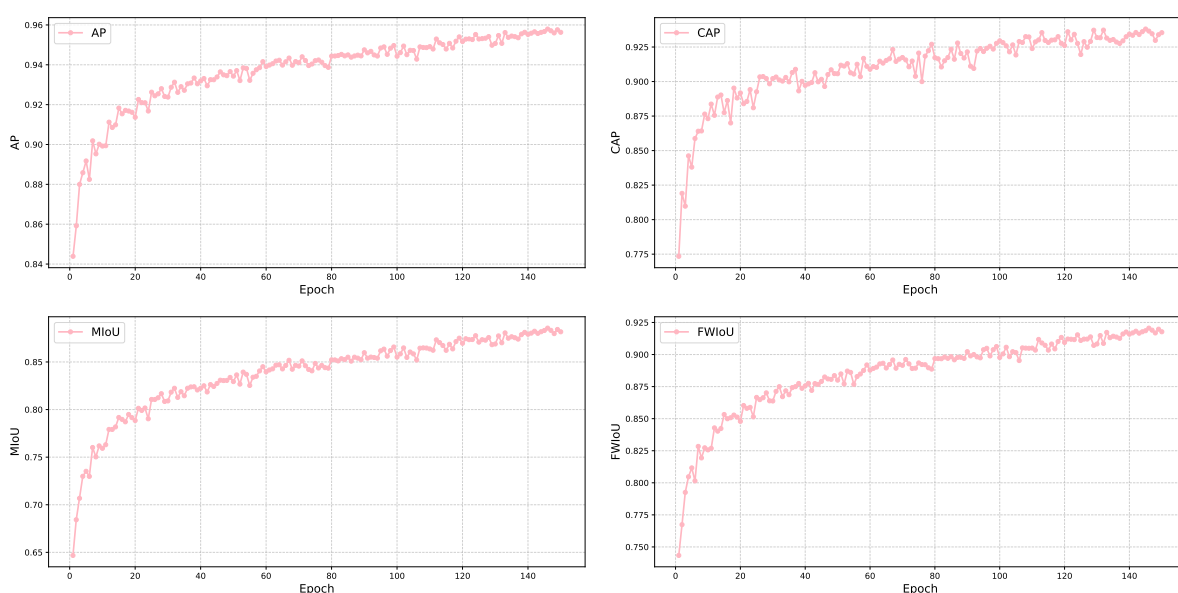


图 21 DeepLab v3+ 最终实验性能

DeepLab v3+ 在 $\mathcal{D}_{test}$ 上的 AP、CAP、MIoU、FWIoU 分别达到 0.9563、0.9353、0.8816、0.9177。

5.1.10 坑洼道路难样本处理效果

图 5 中展示了目标检测算法对坑洼道路图像中模糊图像、多坑图像、坑洼纹理复杂图像等难样本的检测效果，可以看出目标检测算法输出的回归框定位对于难样本的处理

是非常困难且不准确的。因为对于目标检测算法，标注信息没有精确到像素级，而图像中的坑洼特征非常复杂多样，目标检测算法无法在只有粗糙回归框标注的训练情况下对各种复杂的坑洼特征进行精准识别。

经过坑洼道路图像训练后的深度学习语义分割模型 DeepLab v3+，由于进行了像素级的语义信息标注，且通过垂直水平翻转、高斯模糊等数据增广的手段扩充了训练数据，因此具有较高的鲁棒性与坑洼特征识别能力，能够较为准确地识别各类难样本中的坑洼特征。图 22 展示了 DeepLab v3+ 对坑洼道路难样本的分割结果。

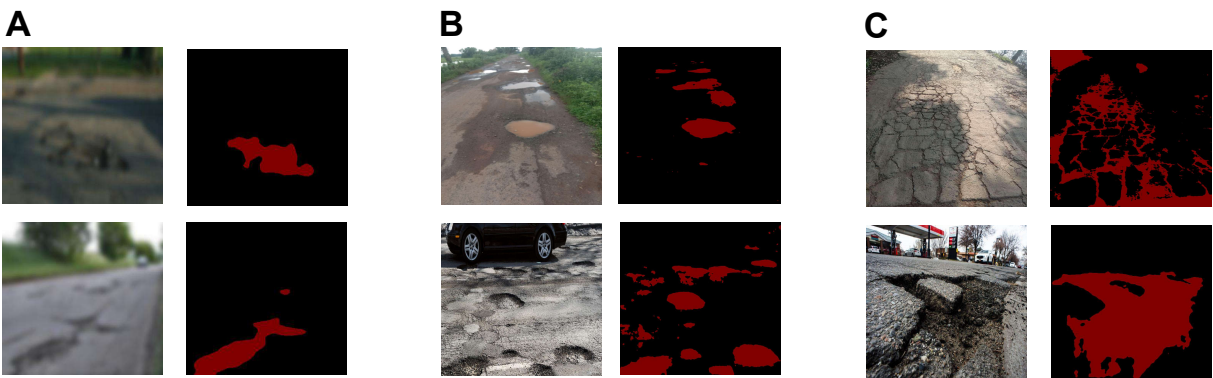


图 22 语义分割算法在坑洼道路图像难样本中的检测结果。(A) 模糊图像。(B) 多坑洼图像。(C) 坑洼纹理复杂图像。

5.2 问题二：计算坑洼面积占比

在本文中，首先使用问题一建立的分类模型从题目所给数据中分离出坑洼路面图像，然后使用训练好的语义分割模型对坑洼道路图像进行语义分割并计算坑洼面积，将结果写入到结果文件中。本问建模流程如图 23 所示。

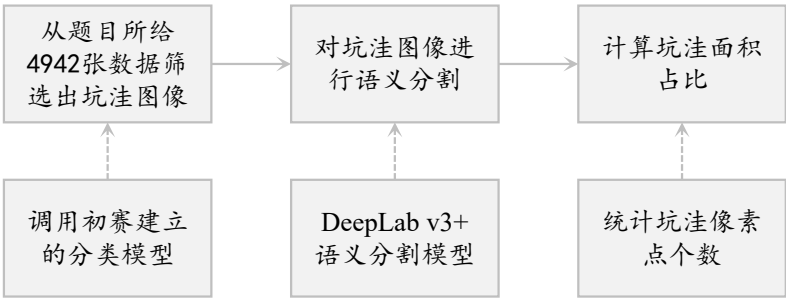


图 23 问题二建模流程

5.2.1 坑洼道路图像分类

对测试数据中所有 4942 张图像进行坑洼边缘特征提取并计算坑洼面积占比计算量较大，为简化计算，本文使用初赛问题建立的基于伪样本生成增广与通道注意力增强的

坑洼道路图像分类模型对 D_{test} 内所有图像进行分类，仅对分类为“坑洼道路”的图像进行坑洼道路语义分割。

经过分类模型进行分类后， D_{test} 4942 张图像内共有 3725 张图像被识别为“坑洼道路图像”，本文对上述图像提取坑洼边缘特征并计算坑洼面积占比。

5.2.2 面积占比计算

使用问题一建立的深度学习语义分割模型对分类出的坑洼道路图像进行语义分割，分割后的语义图像如图所示。部分具体语义分割实验效果见附录。

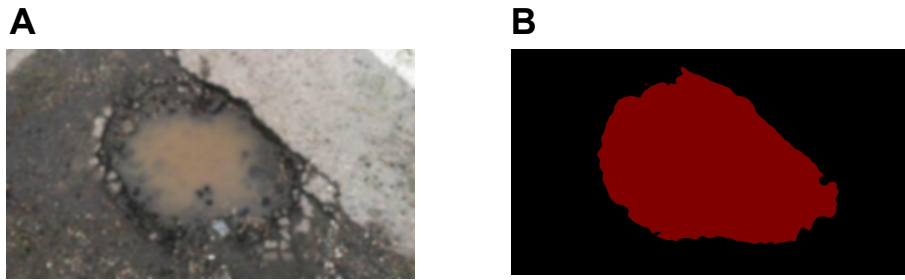


图 24 DeepLab v3+ 图像分割后的语义图

从图 24 中可以看出，在语义分割模型分割出的语义图中，被识别为坑洼图案的像素点为红色，被识别为背景图案的像素点为黑色。红色像素点在 RGB 通道的具体像素值为 (128,0,0)，黑色像素点为 (0,0,0)。因此，对于分割出的语义图 $\mathbf{x}_{seg}$ ，只需计算非零像素点占图像所有像素点的比值即可得出坑洼在原图中的面积占比。具体计算公式如式 (24) 所示。

$$S_{pothole} = \frac{N_{pothole}}{H_{\mathbf{x}_{seg}} \times W_{\mathbf{x}_{seg}}} \quad (24)$$

其中 $N_{pothole}$ 表示 $\mathbf{x}_{seg}$ 中非零像素点的个数， $H_{\mathbf{x}_{seg}}$ 和 $W_{\mathbf{x}_{seg}}$ 分别表示 $\mathbf{x}_{seg}$ 的高度与宽度，两者的乘积即为 $\mathbf{x}_{seg}$ 的像素点个数。

计算出每张图像的 $S_{pothole}$ ，将计算结果填入到 test_result2.csv 文件中。该问题完整处理流程如算法 3 所示。

六、模型的评价与推广

6.1 模型优点

- 基于深度学习的语义分割模型将图像的语义信息精确到像素级，使图像能够精准地提特征复杂地不规则坑洼边缘，并且对模糊图像、多坑洼图像、纹理特征复杂图像等坑洼道路图像中的难样本也有较好的提取效果。

Algorithm 3: 处理问题二流程

Data: 题目所给数据 $\mathcal{D}_{test}$; 初赛问题中训练好的坑洼道路图像分类模型 f_C ; 在坑洼道路图像数据集中训练好的语义分割模型 f'_ϕ ;

```
1 for  $i = 1$  to  $\mathcal{D}_{test}$  do
2   //  $\mathbf{x}_i$  为  $\mathcal{D}_{test}$  中第  $i$  个图像;
3    $f_C$  对图像进行分类;
4   if  $\mathbf{x}_i$  为正常图像 then
5     根据文件名在 test_result2.csv 中填入 0;
6   else
7      $\mathbf{x}_{seg} = f'_\phi(\mathbf{x}_i)$ ;
8     统计  $\mathbf{x}_{seg}$  中非零像素数量  $N_{p_{pothole}}$ ;
9     计算  $\mathbf{x}_{seg}$  高度  $H_{\mathbf{x}_{seg}}$  与宽度  $W_{\mathbf{x}_{seg}}$ ;
10     $S_{pothole} = \frac{N_{p_{pothole}}}{H_{\mathbf{x}_{seg}} \times W_{\mathbf{x}_{seg}}}$ ;
11    根据文件名在 test_result10.csv 中填入  $S_{pothole}$ ;
12 return 填写好面积占比结果的 test_result2.csv 文件
```

- 本文对公开数据集的数据自行进行标注后，使用垂直水平翻转，高斯模糊等图像增广手段，增加数据集中的空间特征复杂性，使模型更能适应场景中物体的不同姿态和方向变化。
- DeepLab v3+ 模型采用编码器-解码器结构来执行语义分割任务，能够将分隔后图像完全复原成原始输入图像的尺寸大小，抑制了语义信息的损失。同时，DeepLab v3+ 模型引入空洞卷积、空间特征金字塔池化，使得模型能更好地提取坑洼图像中的多尺度特征，对不同尺寸的坑洼道路图像均具备分割能力。

6.2 模型缺点

- 模型没有对学习率等超参数设置进行深入调优。
- 缺乏足够的训练数据。

6.3 模型推广

本文基于深度学习构建了用于坑洼道路语义分割的 DeepLab v3+ 语义分割模型。该模型不仅适用于坑洼道路语义分割，对于空间特征较为复杂、环境多样、成像条件参差不齐的图像语义分割任务均适用。同时，本文模型进行了充分的数据增广，对于样本量较小的语义分割任务同样能取得优越的性能。

参考文献

- [1] 冯树杰. 基于深度学习的路面缺陷检测算法研究 [D]. 陕西科技大学, 2023. DOI: 10.27290/d.cnki.gxbqc. 2023. 000082.
- [2] He K, Zhang X, Ren S, et al. Deep residual learning for image recognition[C]//Proceedings of the IEEE conference on computer vision and pattern recognition. 2016: 770-778.
- [3] Long J, Shelhamer E, Darrell T. Fully convolutional networks for semantic segmentation[C]//Proceedings of the IEEE conference on computer vision and pattern recognition. 2015: 3431-3440.
- [4] 于营, 王春平, 付强, 等. 语义分割评价指标和评价方法综述 [J]. Journal of Computer Engineering & Applications, 2023, 59(6).

附录 A 坑洼道路图像语义分割效果展示

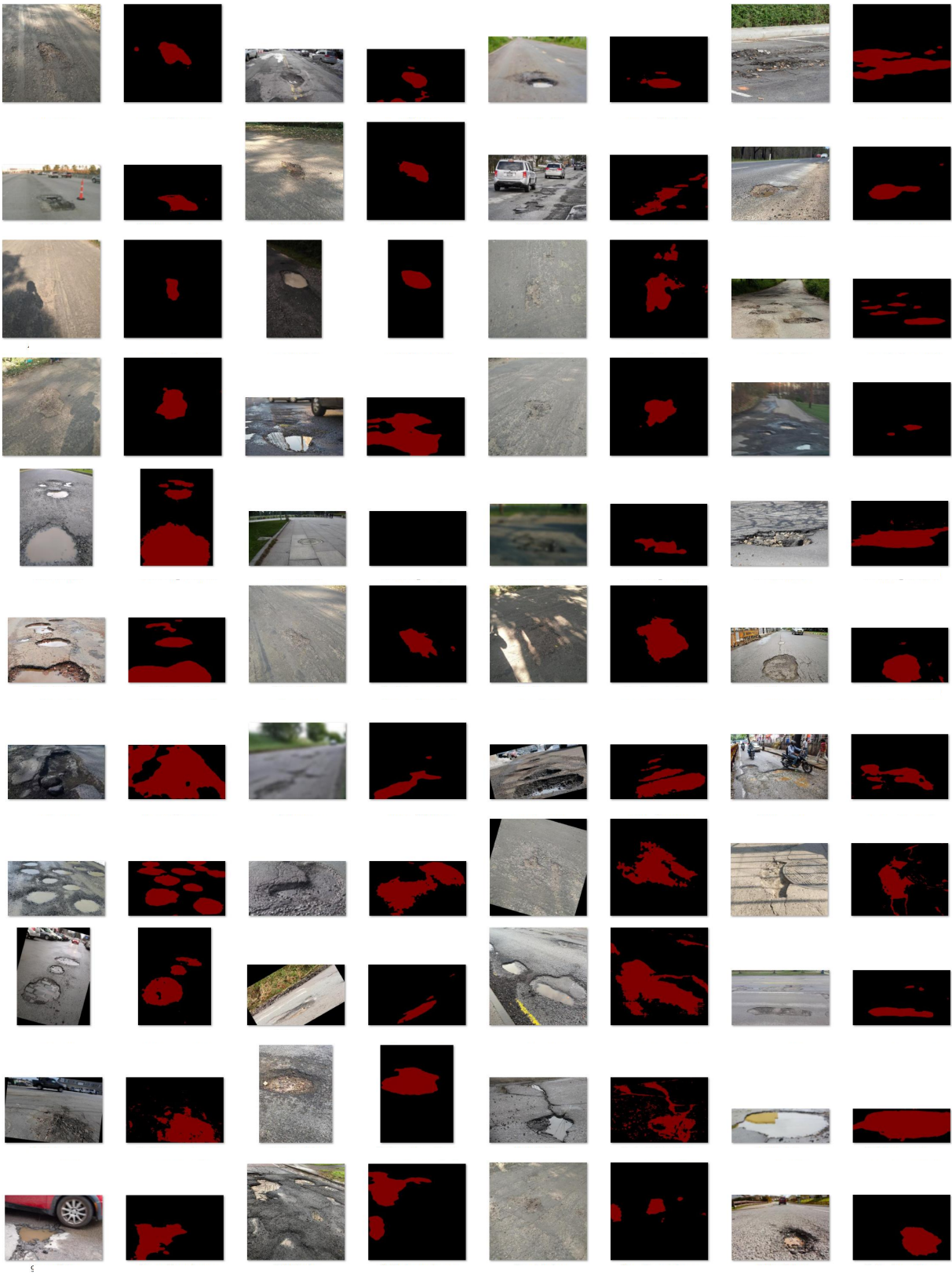


图 25 坑洼道路图像语义分割效果展示 (1)

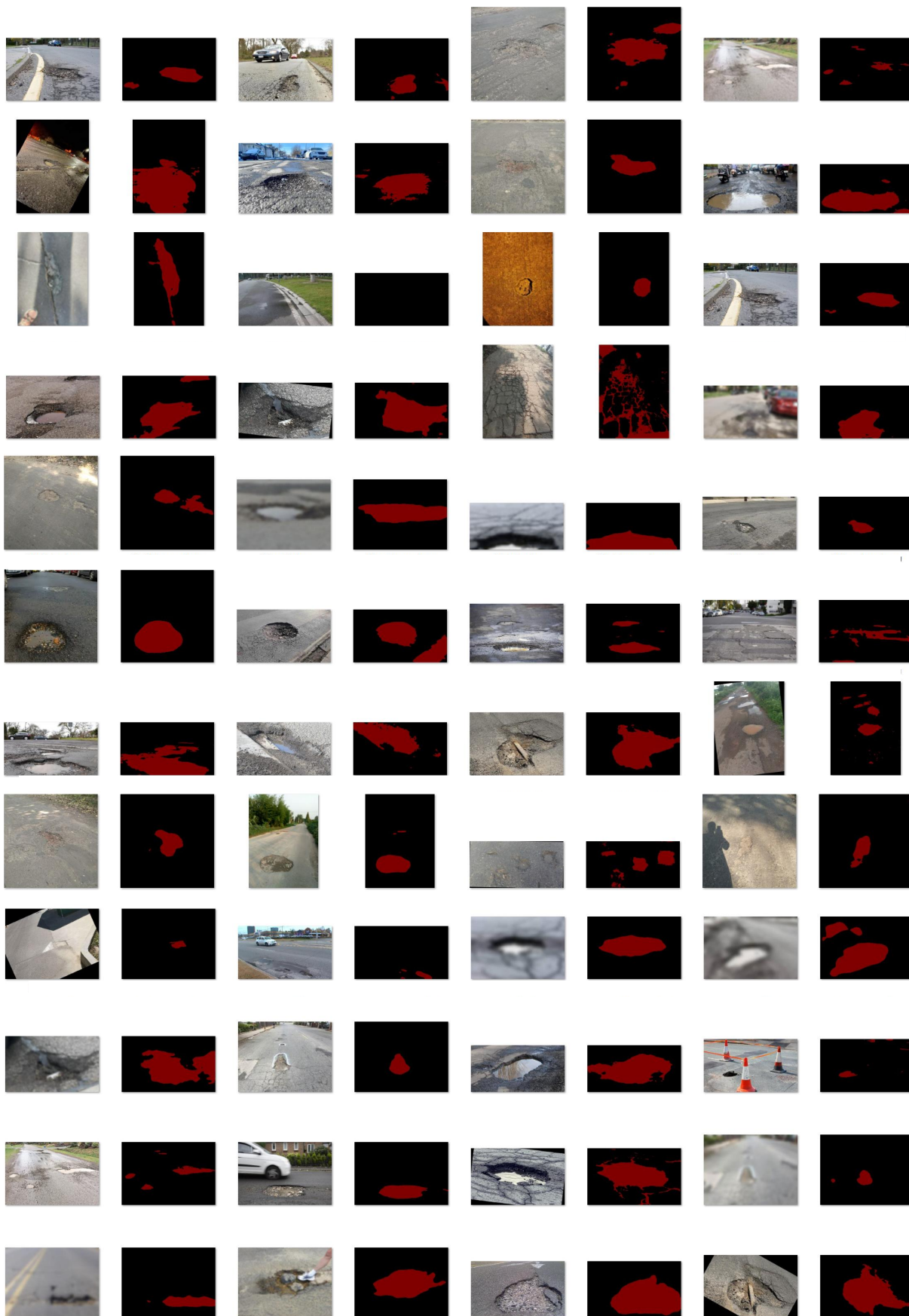


图 26 坑洼道路图像语义分割效果展示 (2)

附录 B 源代码

2.1 问题一源代码

2.1.1 训练与测试源代码

train.py

```
1  import argparse
2  import csv
3  import os
4  import numpy as np
5  from tqdm import tqdm
6
7  from mypath import Path
8  from dataloaders import make_data_loader
9  from modeling.sync_batchnorm.replicate import patch_replication_callback
10 from modeling.deeplab import *
11 from utils.loss import SegmentationLosses
12 from utils.calculate_weights import calculate_weights_labels
13 from utils.lr_scheduler import LR_Scheduler
14 from utils.saver import Saver
15 from utils.summaries import TensorboardSummary
16 from utils.metrics import Evaluator
17
18
19 class Trainer(object):
20     def __init__(self, args):
21         self.args = args
22         self.epoch_data = []
23         self.csv_filename = '../../output.csv'
24
25         # Define Saver
26         self.saver = Saver(args)
27         self.saver.save_experiment_config()
28         # Define Tensorboard Summary
29         self.summary = TensorboardSummary(self.saver.experiment_dir)
30         self.writer = self.summary.create_summary()
31
32         # Define Dataloader
33         kwargs = {'num_workers': args.workers, 'pin_memory': True}
34         self.train_loader, self.val_loader, self.test_loader, self.nclass =
            make_data_loader(args, **kwargs)
35
36         # Define network
37         model = DeepLab(num_classes=self.nclass,
38                         backbone=args.backbone,
```

```

39         output_stride=args.out_stride,
40         sync_bn=args.sync_bn,
41         freeze_bn=args.freeze_bn)
42
43     train_params = [{'params': model.get_1x_lr_params(), 'lr': args.lr},
44                     {'params': model.get_10x_lr_params(), 'lr': args.lr * 10}]
45
46     # Define Optimizer
47     optimizer = torch.optim.SGD(train_params, momentum=args.momentum,
48                                 weight_decay=args.weight_decay, nesterov=args.nesterov)
49
50     # Define Criterion
51     # whether to use class balanced weights
52     if args.use_balanced_weights:
53         classes_weights_path = os.path.join(Path.db_root_dir(args.dataset),
54                                             args.dataset+'_classes_weights.npy')
55         if os.path.isfile(classes_weights_path):
56             weight = np.load(classes_weights_path)
57         else:
58             weight = calculate_weights_labels(args.dataset, self.train_loader, self.nclass)
59             weight = torch.from_numpy(weight.astype(np.float32))
60     else:
61         weight = None
62     self.criterion = SegmentationLosses(weight=weight,
63                                         cuda=args.cuda).build_loss(mode=args.loss_type)
64     self.model, self.optimizer = model, optimizer
65
66     # Define Evaluator
67     self.evaluator = Evaluator(self.nclass)
68
69     # Define lr scheduler
70     self.scheduler = LR_Scheduler(args.lr_scheduler, args.lr,
71                                   args.epochs, len(self.train_loader))
72
73     # Using cuda
74     if args.cuda:
75         self.model = torch.nn.DataParallel(self.model, device_ids=self.args.gpu_ids)
76         patch_replication_callback(self.model)
77         self.model = self.model.cuda()
78
79     # Resuming checkpoint
80     self.best_pred = 0.0
81     if args.resume is not None:
82         if not os.path.isfile(args.resume):
83             raise RuntimeError("=> no checkpoint found at '{}'".format(args.resume))
84         checkpoint = torch.load(args.resume)
85         args.start_epoch = checkpoint['epoch']
86         if args.cuda:

```

```

84         self.model.module.load_state_dict(checkpoint['state_dict'])
85     else:
86         self.model.load_state_dict(checkpoint['state_dict'])
87     if not args.ft:
88         self.optimizer.load_state_dict(checkpoint['optimizer'])
89     self.best_pred = checkpoint['best_pred']
90     print("=> loaded checkpoint '{}' (epoch {})"
91           .format(args.resume, checkpoint['epoch']))
92
93     # Clear start epoch if fine-tuning
94     if args.ft:
95         args.start_epoch = 0
96
97 def training(self, epoch):
98     train_loss = 0.0
99     self.model.train()
100    tbar = tqdm(self.train_loader)
101    num_img_tr = len(self.train_loader)
102    for i, sample in enumerate(tbar):
103        image, target = sample['image'], sample['label']
104        if self.args.cuda:
105            image, target = image.cuda(), target.cuda()
106        self.scheduler(self.optimizer, i, epoch, self.best_pred)
107        self.optimizer.zero_grad()
108        output = self.model(image)
109        loss = self.criterion(output, target)
110        loss.backward()
111        self.optimizer.step()
112        train_loss += loss.item()
113        tbar.set_description('Train loss: %.3f' % (train_loss / (i + 1)))
114        self.writer.add_scalar('train/total_loss_iter', loss.item(), i + num_img_tr * epoch)
115
116    # Show 10 * 3 inference results each epoch
117    if i % (num_img_tr // 10) == 0:
118        global_step = i + num_img_tr * epoch
119        self.summary.visualize_image(self.writer, self.args.dataset, image, target,
120                                    output, global_step)
121
122    self.writer.add_scalar('train/total_loss_epoch', train_loss, epoch)
123    print('[Epoch: %d, numImages: %5d]' % (epoch, i * self.args.batch_size +
124        image.data.shape[0]))
125    print('Loss: %.3f' % train_loss)
126
127    if self.args.no_val:
128        # save checkpoint every epoch
129        is_best = False
130        self.saver.save_checkpoint({

```

```

129         'epoch': epoch + 1,
130         'state_dict': self.model.module.state_dict(),
131         'optimizer': self.optimizer.state_dict(),
132         'best_pred': self.best_pred,
133     }, is_best)
134
135
136 def validation(self, epoch):
137
138     self.model.eval()
139     self.evaluator.reset()
140     tbar = tqdm(self.val_loader, desc='\r')
141     test_loss = 0.0
142     for i, sample in enumerate(tbar):
143         image, target = sample['image'], sample['label']
144         if self.args.cuda:
145             image, target = image.cuda(), target.cuda()
146         with torch.no_grad():
147             output = self.model(image)
148             loss = self.criterion(output, target)
149             test_loss += loss.item()
150             tbar.set_description('Test loss: %.3f' % (test_loss / (i + 1)))
151             pred = output.data.cpu().numpy()
152             target = target.cpu().numpy()
153             pred = np.argmax(pred, axis=1)
154             # Add batch sample into evaluator
155             self.evaluator.add_batch(target, pred)
156
157     # Fast test during the training
158     Acc = self.evaluator.Pixel_Accuracy()
159     Acc_class = self.evaluator.Pixel_Accuracy_Class()
160     mIoU = self.evaluator.Mean_Intersection_over_Union()
161     FWIoU = self.evaluator.Frequency_Weighted_Intersection_over_Union()
162     self.writer.add_scalar('val/total_loss_epoch', test_loss, epoch)
163     self.writer.add_scalar('val/mIoU', mIoU, epoch)
164     self.writer.add_scalar('val/Acc', Acc, epoch)
165     self.writer.add_scalar('val/Acc_class', Acc_class, epoch)
166     self.writer.add_scalar('val/fwIoU', FWIoU, epoch)
167     print('Validation:')
168     print('[Epoch: %d, numImages: %5d]' % (epoch, i * self.args.batch_size +
169         image.data.shape[0]))
170     print("Acc:{}, Acc_class:{}, mIoU:{}, fwIoU: {}".format(Acc, Acc_class, mIoU, FWIoU))
171     print('Loss: %.3f' % test_loss)
172     self.epoch_data.append((epoch, Acc, Acc_class, mIoU, FWIoU, test_loss))
173
174     new_pred = mIoU
175     if new_pred > self.best_pred:

```

```

175         is_best = True
176         self.best_pred = new_pred
177         self.saver.save_checkpoint({
178             'epoch': epoch + 1,
179             'state_dict': self.model.module.state_dict(),
180             'optimizer': self.optimizer.state_dict(),
181             'best_pred': self.best_pred,
182         }, is_best)
183
184     def record(self):
185         with open(self.csv_filename, 'w', newline='') as csvfile:
186             # 创建CSV写入对象
187             csv_writer = csv.writer(csvfile)
188
189             # 写入表头
190             csv_writer.writerow(['Epoch', 'Acc', 'Acc_class', 'mIoU', 'FWIoU', 'Loss'])
191
192             # 写入数据
193             csv_writer.writerows(self.epoch_data)
194
195 def main():
196     parser = argparse.ArgumentParser(description="PyTorch DeeplabV3Plus Training")
197     parser.add_argument('--backbone', type=str, default='resnet',
198                         choices=['resnet', 'xception', 'drn', 'mobilenet'],
199                         help='backbone name (default: resnet)')
200     parser.add_argument('--out-stride', type=int, default=16,
201                         help='network output stride (default: 8)')
202     parser.add_argument('--dataset', type=str, default='mydata',
203                         choices=['pascal', 'coco', 'cityscapes', 'mydata'],
204                         help='dataset name (default: pascal)')
205     parser.add_argument('--use-sbd', action='store_true', default=False,
206                         help='whether to use SBD dataset (default: True)')
207     parser.add_argument('--workers', type=int, default=0,
208                         metavar='N', help='dataloader threads')
209     parser.add_argument('--base-size', type=int, default=513,
210                         help='base image size')
211     parser.add_argument('--crop-size', type=int, default=513,
212                         help='crop image size')
213     parser.add_argument('--sync-bn', type=bool, default=None,
214                         help='whether to use sync bn (default: auto)')
215     parser.add_argument('--freeze-bn', type=bool, default=False,
216                         help='whether to freeze bn parameters (default: False)')
217     parser.add_argument('--loss-type', type=str, default='ce',
218                         choices=['ce', 'focal'],
219                         help='loss func type (default: ce)')
220     # training hyper params
221     parser.add_argument('--epochs', type=int, default=150, metavar='N',

```

```

222         help='number of epochs to train (default: auto)')
223 parser.add_argument('--start_epoch', type=int, default=0,
224                     metavar='N', help='start epochs (default:0)')
225 parser.add_argument('--batch-size', type=int, default=2,
226                     metavar='N', help='input batch size for \
227                             training (default: auto)')
228 parser.add_argument('--test-batch-size', type=int, default=2,
229                     metavar='N', help='input batch size for \
230                             testing (default: auto)')
231 parser.add_argument('--use-balanced-weights', action='store_true', default=False,
232                     help='whether to use balanced weights (default: False)')
233 # optimizer params
234 parser.add_argument('--lr', type=float, default=0.007, metavar='LR',
235                     help='learning rate (default: auto)')
236 parser.add_argument('--lr-scheduler', type=str, default='poly',
237                     choices=['poly', 'step', 'cos'],
238                     help='lr scheduler mode: (default: poly)')
239 parser.add_argument('--momentum', type=float, default=0.9,
240                     metavar='M', help='momentum (default: 0.9)')
241 parser.add_argument('--weight-decay', type=float, default=5e-4,
242                     metavar='M', help='w-decay (default: 5e-4)')
243 parser.add_argument('--nesterov', action='store_true', default=False,
244                     help='whether use nesterov (default: False)')
245 # cuda, seed and logging
246 parser.add_argument('--no-cuda', action='store_true', default=
247                     False, help='disables CUDA training')
248 parser.add_argument('--gpu-ids', type=str, default='0',
249                     help='use which gpu to train, must be a \
250                             comma-separated list of integers only (default=0)')
251 parser.add_argument('--seed', type=int, default=1, metavar='S',
252                     help='random seed (default: 1)')
253 # checking point
254 parser.add_argument('--resume', type=str, default=None,
255                     help='put the path to resuming file if needed')
256 parser.add_argument('--checkname', type=str, default='deeplab-resnet',
257                     help='set the checkpoint name')
258 # finetuning pre-trained models
259 parser.add_argument('--ft', action='store_true', default=True,
260                     help='finetuning on a different dataset')
261 # evaluation option
262 parser.add_argument('--eval-interval', type=int, default=1,
263                     help='evaluation interval (default: 1)')
264 parser.add_argument('--no-val', action='store_true', default=False,
265                     help='skip validation during training')
266
267 args = parser.parse_args()
268 args.cuda = not args.no_cuda and torch.cuda.is_available()

```

```

269 if args.cuda:
270     try:
271         args.gpu_ids = [int(s) for s in args.gpu_ids.split(',')]
272     except ValueError:
273         raise ValueError('Argument --gpu_ids must be a comma-separated list of integers
274                             only')
275
276 if args.sync_bn is None:
277     if args.cuda and len(args.gpu_ids) > 1:
278         args.sync_bn = True
279     else:
280         args.sync_bn = False
281
282 # default settings for epochs, batch_size and lr
283 if args.epochs is None:
284     epoches = {
285         'coco': 30,
286         'cityscapes': 200,
287         'pascal': 50,
288     }
289     args.epochs = epoches[args.dataset.lower()]
290
291 if args.batch_size is None:
292     args.batch_size = 4 * len(args.gpu_ids)
293
294 if args.test_batch_size is None:
295     args.test_batch_size = args.batch_size
296
297 if args.lr is None:
298     lrs = {
299         'coco': 0.1,
300         'cityscapes': 0.01,
301         'pascal': 0.007,
302     }
303     args.lr = lrs[args.dataset.lower()] / (4 * len(args.gpu_ids)) * args.batch_size
304
305 if args.checkname is None:
306     args.checkname = 'deeplab-'+str(args.backbone)
307 print(args)
308 torch.manual_seed(args.seed)
309 trainer = Trainer(args)
310 print('Starting Epoch:', trainer.args.start_epoch)
311 print('Total Epoches:', trainer.args.epochs)
312 for epoch in range(trainer.args.start_epoch, trainer.args.epochs):
313     trainer.training(epoch)
314     if not trainer.args.no_val and epoch % args.eval_interval == (args.eval_interval - 1):

```

```

315         trainer.validation(epoch)
316
317     trainer.record()
318     trainer.writer.close()
319
320 if __name__ == "__main__":
321     main()

```

demo.py

```

1  #
2  # demo.py
3  #
4  import argparse
5  import os
6  import numpy as np
7  import time
8  import gc
9  from modeling.deeplab import *
10 from dataloaders import custom_transforms as tr
11 from PIL import Image
12 from torchvision import transforms
13 from dataloaders.utils import *
14 from torchvision.utils import make_grid, save_image
15
16 def main():
17
18     parser = argparse.ArgumentParser(description="PyTorch DeeplabV3Plus Training")
19     parser.add_argument('--in-path', type=str, required=False, help='image to test',
20                         default='./dataset/test_data/potholes')
21     parser.add_argument('--out-path', type=str, required=False, help='mask image to save',
22                         default='./dataset/output')
23     parser.add_argument('--backbone', type=str, default='resnet',
24                         choices=['resnet', 'xception', 'drn', 'mobilenet'],
25                         help='backbone name (default: resnet)')
26     parser.add_argument('--ckpt', type=str,
27                         default='./run/mydata/deeplab-resnet/model_best.pth.tar',
28                         help='saved model')
29     parser.add_argument('--out-stride', type=int, default=16,
30                         help='network output stride (default: 8)')
31     parser.add_argument('--no-cuda', action='store_true', default=False,
32                         help='disables CUDA training')
33     parser.add_argument('--gpu-ids', type=str, default='0',
34                         help='use which gpu to train, must be a \
35                             comma-separated list of integers only (default=0)')
36     parser.add_argument('--dataset', type=str, default='mydata',
37                         choices=['pascal', 'coco', 'cityscapes', 'invoice', 'mydata'],

```

```

35         help='dataset name (default: pascal)')
36 parser.add_argument('--crop-size', type=int, default=500,
37                     help='crop image size')
38 parser.add_argument('--num-classes', type=int, default=2,
39                     help='crop image size')
40 parser.add_argument('--sync-bn', type=bool, default=None,
41                     help='whether to use sync bn (default: auto)')
42 parser.add_argument('--freeze-bn', type=bool, default=False,
43                     help='whether to freeze bn parameters (default: False)')
44
45 args = parser.parse_args()
46 args.cuda = not args.no_cuda and torch.cuda.is_available()
47 if args.cuda:
48     try:
49         args.gpu_ids = [int(s) for s in args.gpu_ids.split(',')]
50     except ValueError:
51         raise ValueError('Argument --gpu_ids must be a comma-separated list of integers
52                             only')
53
54 if args.sync_bn is None:
55     if args.cuda and len(args.gpu_ids) > 1:
56         args.sync_bn = True
57     else:
58         args.sync_bn = False
59
60 model_s_time = time.time()
61 model = DeepLab(num_classes=args.num_classes,
62                 backbone=args.backbone,
63                 output_stride=args.out_stride,
64                 sync_bn=args.sync_bn,
65                 freeze_bn=args.freeze_bn)
66
67 ckpt = torch.load(args.ckpt, map_location='cpu')
68 model.load_state_dict(ckpt['state_dict'])
69 model = model.cuda()
70 model.eval()
71 model_u_time = time.time()
72 model_load_time = model_u_time - model_s_time
73 print("model load time is {}".format(model_load_time))
74
75 composed_transforms = transforms.Compose([
76     tr.Normalize(mean=(0.485, 0.456, 0.406), std=(0.229, 0.224, 0.225)),
77     tr.ToTensor()])
78
79 for name in os.listdir(args.in_path):
80     with torch.no_grad():
81         gc.collect()
82         torch.cuda.empty_cache()
83         s_time = time.time()

```

```

81     image = Image.open(os.path.join(args.in_path, name)).convert('RGB')
82     target = Image.open(os.path.join(args.in_path, name)).convert('L')
83     sample = {'image': image, 'label': target}
84
85     tensor_in = composed_transforms(sample)['image'].unsqueeze(0)
86     tensor_in = tensor_in.cuda()
87
88
89     output = model(tensor_in)
90
91     grid_image = make_grid(decode_seg_map_sequence(torch.max(output[:3],
92         1)[1].detach().cpu().numpy()),
93         3, normalize=False, range=(0, 255))
94     save_image(grid_image, os.path.join(args.in_path, "{}_mask.png".format(name[0:-4])))
95
96     u_time = time.time()
97     img_time = u_time - s_time
98     print("image:{} time: {}".format(name, img_time))
99
100     print("image save in in_path.")
101 if __name__ == "__main__":
102     main()
103
104 # python demo.py --in-path your_file --out-path your_dst_file

```

2.1.2 构建数据集源代码

json2dataset.py

```

1  #!/usr/bin/env python
2  # -*- coding: utf-8 -*-
3  # 这段代码负责将labelme标注后的文件转换为正确的格式
4  # @Time : 2023/12/15 下午12:56
5  # @Author : SiHang Xu
6  import os
7  files=os.listdir('./')
8  files.remove('json2dataset.py') # 删除这个py文件本身
9  for i in range(len(files)):
10     os.system('labelme_json_to_dataset '+files[i])

```

split_dataset.py

```

1  #随机生成trian.txt, trainval.txt, val.txt
2
3  import os

```

```

4 import random
5
6 # 设置路径和文件名
7 path = "./JPEGImages" # 替换为你的图片所在的路径
8 train_file = "train.txt"
9 trainval_file = "trainval.txt"
10 val_file = "val.txt"
11
12 # 获取路径下的所有图片文件
13 image_files = [f for f in os.listdir(path) if f.endswith(".jpg") or f.endswith(".png")]
14
15 # 打乱文件列表
16 random.shuffle(image_files)
17
18 # 计算划分的索引
19 total_images = len(image_files)
20 train_split = 1400
21 trainval_split = 1500
22
23 # 划分文件列表
24 train_images = [os.path.splitext(f)[0] for f in image_files[:train_split]]
25 trainval_images = [os.path.splitext(f)[0] for f in image_files[train_split:trainval_split]]
26 val_images = [os.path.splitext(f)[0] for f in image_files[trainval_split:]]
27
28 # 写入文件名到train.txt
29 with open(train_file, "w") as f:
30     f.write("\n".join(train_images))
31
32 # 写入文件名到trainval.txt
33 with open(trainval_file, "w") as f:
34     f.write("\n".join(trainval_images))
35
36 # 写入文件名到val.txt
37 with open(val_file, "w") as f:
38     f.write("\n".join(val_images))

```

custom_transform.py

```

1 import torch
2 import random
3 import numpy as np
4
5 from PIL import Image, ImageOps, ImageFilter
6
7 class Normalize(object):
8     """Normalize a tensor image with mean and standard deviation.
9     Args:

```

```

10     mean (tuple): means for each channel.
11     std (tuple): standard deviations for each channel.
12     """
13     def __init__(self, mean=(0., 0., 0.), std=(1., 1., 1.)):
14         self.mean = mean
15         self.std = std
16
17     def __call__(self, sample):
18         img = sample['image']
19         mask = sample['label']
20         img = np.array(img).astype(np.float32)
21         mask = np.array(mask).astype(np.float32)
22         img /= 255.0
23         img -= self.mean
24         img /= self.std
25
26         return {'image': img,
27                 'label': mask}
28
29
30 class ToTensor(object):
31     """Convert ndarrays in sample to Tensors."""
32
33     def __call__(self, sample):
34         # swap color axis because
35         # numpy image: H x W x C
36         # torch image: C X H X W
37         img = sample['image']
38         mask = sample['label']
39         img = np.array(img).astype(np.float32).transpose((2, 0, 1))
40         mask = np.array(mask).astype(np.float32)
41
42         img = torch.from_numpy(img).float()
43         mask = torch.from_numpy(mask).float()
44
45         return {'image': img,
46                 'label': mask}
47
48
49 class RandomHorizontalFlip(object):
50     def __call__(self, sample):
51         img = sample['image']
52         mask = sample['label']
53         if random.random() < 0.5:
54             img = img.transpose(Image.FLIP_LEFT_RIGHT)
55             mask = mask.transpose(Image.FLIP_LEFT_RIGHT)
56

```

```

57         return {'image': img,
58                'label': mask}
59
60
61 class RandomRotate(object):
62     def __init__(self, degree):
63         self.degree = degree
64
65     def __call__(self, sample):
66         img = sample['image']
67         mask = sample['label']
68         rotate_degree = random.uniform(-1*self.degree, self.degree)
69         img = img.rotate(rotate_degree, Image.BILINEAR)
70         mask = mask.rotate(rotate_degree, Image.NEAREST)
71
72         return {'image': img,
73                'label': mask}
74
75
76 class RandomGaussianBlur(object):
77     def __call__(self, sample):
78         img = sample['image']
79         mask = sample['label']
80         if random.random() < 0.5:
81             img = img.filter(ImageFilter.GaussianBlur(
82                 radius=random.random()))
83
84         return {'image': img,
85                'label': mask}
86
87
88 class RandomScaleCrop(object):
89     def __init__(self, base_size, crop_size, fill=0):
90         self.base_size = base_size
91         self.crop_size = crop_size
92         self.fill = fill
93
94     def __call__(self, sample):
95         img = sample['image']
96         mask = sample['label']
97         # random scale (short edge)
98         short_size = random.randint(int(self.base_size * 0.5), int(self.base_size * 2.0))
99         w, h = img.size
100         if h > w:
101             ow = short_size
102             oh = int(1.0 * h * ow / w)
103         else:

```

```

104         oh = short_size
105         ow = int(1.0 * w * oh / h)
106         img = img.resize((ow, oh), Image.BILINEAR)
107         mask = mask.resize((ow, oh), Image.NEAREST)
108         # pad crop
109         if short_size < self.crop_size:
110             padh = self.crop_size - oh if oh < self.crop_size else 0
111             padw = self.crop_size - ow if ow < self.crop_size else 0
112             img = ImageOps.expand(img, border=(0, 0, padw, padh), fill=0)
113             mask = ImageOps.expand(mask, border=(0, 0, padw, padh), fill=self.fill)
114         # random crop crop_size
115         w, h = img.size
116         x1 = random.randint(0, w - self.crop_size)
117         y1 = random.randint(0, h - self.crop_size)
118         img = img.crop((x1, y1, x1 + self.crop_size, y1 + self.crop_size))
119         mask = mask.crop((x1, y1, x1 + self.crop_size, y1 + self.crop_size))
120
121         return {'image': img,
122                'label': mask}
123
124
125 class FixScaleCrop(object):
126     def __init__(self, crop_size):
127         self.crop_size = crop_size
128
129     def __call__(self, sample):
130         img = sample['image']
131         mask = sample['label']
132         w, h = img.size
133         if w > h:
134             oh = self.crop_size
135             ow = int(1.0 * w * oh / h)
136         else:
137             ow = self.crop_size
138             oh = int(1.0 * h * ow / w)
139         img = img.resize((ow, oh), Image.BILINEAR)
140         mask = mask.resize((ow, oh), Image.NEAREST)
141         # center crop
142         w, h = img.size
143         x1 = int(round((w - self.crop_size) / 2.))
144         y1 = int(round((h - self.crop_size) / 2.))
145         img = img.crop((x1, y1, x1 + self.crop_size, y1 + self.crop_size))
146         mask = mask.crop((x1, y1, x1 + self.crop_size, y1 + self.crop_size))
147
148         return {'image': img,
149                'label': mask}
150

```

```

151 class FixedResize(object):
152     def __init__(self, size):
153         self.size = (size, size) # size: (h, w)
154
155     def __call__(self, sample):
156         img = sample['image']
157         mask = sample['label']
158
159         assert img.size == mask.size
160
161         img = img.resize(self.size, Image.BILINEAR)
162         mask = mask.resize(self.size, Image.NEAREST)
163
164         return {'image': img,
165                 'label': mask}

```

mydata.py

```

1  from __future__ import print_function, division
2  import os
3  from PIL import Image
4  import numpy as np
5  from torch.utils.data import Dataset
6  from mypath import Path
7  from torchvision import transforms
8  from dataloaders import custom_transforms as tr
9
10 class VOCSegmentation(Dataset):
11     """
12     PascalVoc dataset
13     """
14     NUM_CLASSES = 2
15
16     def __init__(self,
17                 args,
18                 base_dir=Path.db_root_dir('mydata'),
19                 split='train',
20                 ):
21         """
22         :param base_dir: path to VOC dataset directory
23         :param split: train/val
24         :param transform: transform to apply
25         """
26         super().__init__()
27         self._base_dir = base_dir
28         self._image_dir = os.path.join(self._base_dir, 'JPEGImages')
29         self._cat_dir = os.path.join(self._base_dir, 'SegmentationClass')

```

```

30
31     if isinstance(split, str):
32         self.split = [split]
33     else:
34         split.sort()
35         self.split = split
36
37     self.args = args
38
39     _splits_dir = os.path.join(self._base_dir, 'ImageSets', 'Segmentation')
40
41     self.im_ids = []
42     self.images = []
43     self.categories = []
44
45     for splt in self.split:
46         with open(os.path.join(os.path.join(_splits_dir, splt + '.txt')), "r") as f:
47             lines = f.read().splitlines()
48
49             for ii, line in enumerate(lines):
50                 _image = os.path.join(self._image_dir, line + ".png")
51                 _cat = os.path.join(self._cat_dir, line + ".png")
52                 assert os.path.isfile(_image)
53                 assert os.path.isfile(_cat)
54                 self.im_ids.append(line)
55                 self.images.append(_image)
56                 self.categories.append(_cat)
57
58     assert (len(self.images) == len(self.categories))
59
60     # Display stats
61     print('Number of images in {}: {}'.format(split, len(self.images)))
62
63     def __len__(self):
64         return len(self.images)
65
66
67     def __getitem__(self, index):
68         _img, _target = self._make_img_gt_point_pair(index)
69         sample = {'image': _img, 'label': _target}
70
71         for split in self.split:
72             if split == "train":
73                 return self.transform_tr(sample)
74             elif split == 'val':
75                 return self.transform_val(sample)
76

```

```

77
78 def _make_img_gt_point_pair(self, index):
79     _img = Image.open(self.images[index]).convert('RGB')
80     _target = Image.open(self.categories[index])
81
82     return _img, _target
83
84 def transform_tr(self, sample):
85     composed_transforms = transforms.Compose([
86         tr.RandomHorizontalFlip(),
87         tr.RandomScaleCrop(base_size=self.args.base_size, crop_size=self.args.crop_size),
88         tr.RandomGaussianBlur(),
89         tr.Normalize(mean=(0.485, 0.456, 0.406), std=(0.229, 0.224, 0.225)),
90         tr.ToTensor()])
91
92     return composed_transforms(sample)
93
94 def transform_val(self, sample):
95
96     composed_transforms = transforms.Compose([
97         tr.FixScaleCrop(crop_size=self.args.crop_size),
98         tr.Normalize(mean=(0.485, 0.456, 0.406), std=(0.229, 0.224, 0.225)),
99         tr.ToTensor()])
100
101     return composed_transforms(sample)
102
103 def __str__(self):
104     return 'VOC2012(split=' + str(self.split) + ')'
105
106
107 if __name__ == '__main__':
108     from dataloaders.utils import decode_segmap
109     from torch.utils.data import DataLoader
110     import matplotlib.pyplot as plt
111     import argparse
112
113     parser = argparse.ArgumentParser()
114     args = parser.parse_args()
115     args.base_size = 513
116     args.crop_size = 513
117
118     voc_train = VOCSegmentation(args, split='train')
119
120     dataloader = DataLoader(voc_train, batch_size=5, shuffle=True, num_workers=0)
121
122     for ii, sample in enumerate(dataloader):
123         for jj in range(sample["image"].size()[0]):

```

```

124         img = sample['image'].numpy()
125         gt = sample['label'].numpy()
126         tmp = np.array(gt[jj]).astype(np.uint8)
127         segmap = decode_segmap(tmp, dataset='pascal')
128         img_tmp = np.transpose(img[jj], axes=[1, 2, 0])
129         img_tmp *= (0.229, 0.224, 0.225)
130         img_tmp += (0.485, 0.456, 0.406)
131         img_tmp *= 255.0
132         img_tmp = img_tmp.astype(np.uint8)
133         plt.figure()
134         plt.title('display')
135         plt.subplot(211)
136         plt.imshow(img_tmp)
137         plt.subplot(212)
138         plt.imshow(segmap)
139
140         if ii == 1:
141             break
142
143     plt.show(block=True)

```

2.1.3 数据增广源代码

augment_image.py

```

1  #!/usr/bin/env python
2  # -*- coding: utf-8 -*-
3  # 图像增广
4  # @Time : 2023/12/12 下午8:05
5  # @Author : SiHang Xu
6  from torchvision import transforms
7  from PIL import Image
8  import os
9  import random
10
11  count = 0
12
13
14  def augment_images(input_path, output_path, seed=42):
15      global count
16      # 设置随机数种子
17
18
19      random.seed(seed)
20
21      # 创建输出路径

```

```

22 os.makedirs(output_path, exist_ok=True)
23
24 # 获取两组图片的文件名列表
25 files_group1 = os.listdir(input_path + '/JPEGImages')
26 files_group2 = os.listdir(input_path + '/SegmentationClass')
27
28 # 对每一对文件进行相同的增广操作
29 for file1, file2 in zip(files_group1, files_group2):
30     image1 = Image.open(os.path.join(input_path + '/JPEGImages', file1))
31     image2 = Image.open(os.path.join(input_path + '/SegmentationClass', file2))
32
33     # 对两组图像应用相同的随机数操作
34     random.seed(seed)
35     rand_val = random.random()
36
37     augmentations = transforms.Compose([
38         transforms.RandomHorizontalFlip(p=1),
39         # 添加其他需要的增广操作
40     ])
41     # 设置相同的随机数种子，确保增广一致
42     augmented_image1 = augmentations(image1)
43     augmented_image2 = augmentations(image2)
44
45     # 保存增广后的图片
46     augmented_image1.save(os.path.join(output_path, 'JPEGImages', str(count) + '.png'))
47     augmented_image2.save(os.path.join(output_path, 'SegmentationClass', str(count) +
48                             '.png'))
49     count += 1
50
51 def aug_images(input_path, output_path, seed=42):
52     global count
53     # 设置随机数种子
54
55
56     random.seed(seed)
57
58     # 创建输出路径
59     os.makedirs(output_path, exist_ok=True)
60
61     # 获取两组图片的文件名列表
62     files_group1 = os.listdir(input_path + '/JPEGImages')
63     files_group2 = os.listdir(input_path + '/SegmentationClass')
64
65     # 对每一对文件进行相同的增广操作
66     for file1, file2 in zip(files_group1, files_group2):
67         image1 = Image.open(os.path.join(input_path + '/JPEGImages', file1))

```

```

68     image2 = Image.open(os.path.join(input_path + '/SegmentationClass', file2))
69
70     # 对两组图像应用相同的随机数操作
71     random.seed(seed)
72     rand_val = random.random()
73
74     augmentations = transforms.Compose([
75         transforms.RandomVerticalFlip(p=1),
76         # 添加其他需要的增广操作
77     ])
78     # 设置相同的随机数种子，确保增广一致
79     augmented_image1 = augmentations(image1)
80     augmented_image2 = augmentations(image2)
81
82     # 保存增广后的图片
83     augmented_image1.save(os.path.join(output_path, 'JPEGImages', str(count) + '.png'))
84     augmented_image2.save(os.path.join(output_path, 'SegmentationClass', str(count) +
85                             '.png'))
86     count += 1
87
88 def gauss_images(input_path, output_path, seed=42):
89     global count
90     # 设置随机数种子
91
92
93     random.seed(seed)
94
95     # 创建输出路径
96     os.makedirs(output_path, exist_ok=True)
97
98     # 获取两组图片的文件名列表
99     files_group1 = os.listdir(input_path + '/JPEGImages')
100    files_group2 = os.listdir(input_path + '/SegmentationClass')
101
102    # 对每一对文件进行相同的增广操作
103    for file1, file2 in zip(files_group1, files_group2):
104        image1 = Image.open(os.path.join(input_path + '/JPEGImages', file1))
105        image2 = Image.open(os.path.join(input_path + '/SegmentationClass', file2))
106
107        # 对两组图像应用相同的随机数操作
108        random.seed(seed)
109        rand_val = random.random()
110
111        augmentations = transforms.Compose([
112            transforms.GaussianBlur(kernel_size=11, sigma=5)
113            # 添加其他需要的增广操作

```

```

114     ])
115     # 设置相同的随机数种子，确保增广一致
116     augmented_image1 = augmentations(image1)
117     augmented_image2 = image2
118
119     # 保存增广后的图片
120     augmented_image1.save(os.path.join(output_path, 'JPEGImages', str(count) + '.png'))
121     augmented_image2.save(os.path.join(output_path, 'SegmentationClass', str(count) +
122                                '.png'))
123     count += 1
124
125     # 用法示例
126     input_path = './test'
127     output_path = './output'
128
129     augment_images(input_path, output_path)
130     aug_images(input_path, output_path)
131     gauss_images(input_path, output_path)

```

2.1.4 DeepLab v3+ 网络结构源代码

aspp.py

```

1  import math
2  import torch
3  import torch.nn as nn
4  import torch.nn.functional as F
5  from modeling.sync_batchnorm.batchnorm import SynchronizedBatchNorm2d
6
7  class _ASPPModule(nn.Module):
8      def __init__(self, inplanes, planes, kernel_size, padding, dilation, BatchNorm):
9          super(_ASPPModule, self).__init__()
10          self.atrous_conv = nn.Conv2d(inplanes, planes, kernel_size=kernel_size,
11                                       stride=1, padding=padding, dilation=dilation, bias=False)
12          self.bn = BatchNorm(planes)
13          self.relu = nn.ReLU()
14
15          self._init_weight()
16
17      def forward(self, x):
18          x = self.atrous_conv(x)
19          x = self.bn(x)
20
21          return self.relu(x)
22
23      def _init_weight(self):

```

```

24     for m in self.modules():
25         if isinstance(m, nn.Conv2d):
26             torch.nn.init.kaiming_normal_(m.weight)
27         elif isinstance(m, SynchronizedBatchNorm2d):
28             m.weight.data.fill_(1)
29             m.bias.data.zero_()
30         elif isinstance(m, nn.BatchNorm2d):
31             m.weight.data.fill_(1)
32             m.bias.data.zero_()
33
34     class ASPP(nn.Module):
35     def __init__(self, backbone, output_stride, BatchNorm):
36         super(ASPP, self).__init__()
37         if backbone == 'drn':
38             inplanes = 512
39         elif backbone == 'mobilenet':
40             inplanes = 320
41         else:
42             inplanes = 2048
43         if output_stride == 16:
44             dilations = [1, 6, 12, 18]
45         elif output_stride == 8:
46             dilations = [1, 12, 24, 36]
47         else:
48             raise NotImplementedError
49
50         self.aspp1 = _ASPPModule(inplanes, 256, 1, padding=0, dilation=dilations[0],
51                                   BatchNorm=BatchNorm)
52         self.aspp2 = _ASPPModule(inplanes, 256, 3, padding=dilations[1], dilation=dilations[1],
53                                   BatchNorm=BatchNorm)
54         self.aspp3 = _ASPPModule(inplanes, 256, 3, padding=dilations[2], dilation=dilations[2],
55                                   BatchNorm=BatchNorm)
56         self.aspp4 = _ASPPModule(inplanes, 256, 3, padding=dilations[3], dilation=dilations[3],
57                                   BatchNorm=BatchNorm)
58
59         self.global_avg_pool = nn.Sequential(nn.AdaptiveAvgPool2d((1, 1)),
60                                               nn.Conv2d(inplanes, 256, 1, stride=1, bias=False),
61                                               BatchNorm(256),
62                                               nn.ReLU())
63
64         self.conv1 = nn.Conv2d(1280, 256, 1, bias=False)
65         self.bn1 = BatchNorm(256)
66         self.relu = nn.ReLU()
67         self.dropout = nn.Dropout(0.5)
68         self._init_weight()
69
70     def forward(self, x):
71         x1 = self.aspp1(x)

```

```

67     x2 = self.aspp2(x)
68     x3 = self.aspp3(x)
69     x4 = self.aspp4(x)
70     x5 = self.global_avg_pool(x)
71     x5 = F.interpolate(x5, size=x4.size()[2:], mode='bilinear', align_corners=True)
72     x = torch.cat((x1, x2, x3, x4, x5), dim=1)
73
74     x = self.conv1(x)
75     x = self.bn1(x)
76     x = self.relu(x)
77
78     return self.dropout(x)
79
80 def _init_weight(self):
81     for m in self.modules():
82         if isinstance(m, nn.Conv2d):
83             torch.nn.init.kaiming_normal_(m.weight)
84         elif isinstance(m, SynchronizedBatchNorm2d):
85             m.weight.data.fill_(1)
86             m.bias.data.zero_()
87         elif isinstance(m, nn.BatchNorm2d):
88             m.weight.data.fill_(1)
89             m.bias.data.zero_()
90
91
92 def build_aspp(backbone, output_stride, BatchNorm):
93     return ASPP(backbone, output_stride, BatchNorm)

```

decoder.py

```

1  import math
2  import torch
3  import torch.nn as nn
4  import torch.nn.functional as F
5  from modeling.sync_batchnorm.batchnorm import SynchronizedBatchNorm2d
6
7  class Decoder(nn.Module):
8      def __init__(self, num_classes, backbone, BatchNorm):
9          super(Decoder, self).__init__()
10         if backbone == 'resnet' or backbone == 'drn':
11             low_level_inplanes = 256
12         elif backbone == 'xception':
13             low_level_inplanes = 128
14         elif backbone == 'mobilenet':
15             low_level_inplanes = 24
16         else:
17             raise NotImplementedError

```

```

18
19     self.conv1 = nn.Conv2d(low_level_inplanes, 48, 1, bias=False)
20     self.bn1 = BatchNorm(48)
21     self.relu = nn.ReLU()
22     self.last_conv = nn.Sequential(nn.Conv2d(304, 256, kernel_size=3, stride=1, padding=1,
23                                     bias=False),
24                                    BatchNorm(256),
25                                    nn.ReLU(),
26                                    nn.Dropout(0.5),
27                                    nn.Conv2d(256, 256, kernel_size=3, stride=1, padding=1,
28                                             bias=False),
29                                    BatchNorm(256),
30                                    nn.ReLU(),
31                                    nn.Dropout(0.1),
32                                    nn.Conv2d(256, num_classes, kernel_size=1, stride=1))
33
34     self._init_weight()
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60
61
62
63
64
65
66
67
68
69
70
71
72
73
74
75
76
77
78
79
80
81
82
83
84
85
86
87
88
89
90
91
92
93
94
95
96
97
98
99
100
101
102
103
104
105
106
107
108
109
110
111
112
113
114
115
116
117
118
119
120
121
122
123
124
125
126
127
128
129
130
131
132
133
134
135
136
137
138
139
140
141
142
143
144
145
146
147
148
149
150
151
152
153
154
155
156
157
158
159
160
161
162
163
164
165
166
167
168
169
170
171
172
173
174
175
176
177
178
179
180
181
182
183
184
185
186
187
188
189
190
191
192
193
194
195
196
197
198
199
200
201
202
203
204
205
206
207
208
209
210
211
212
213
214
215
216
217
218
219
220
221
222
223
224
225
226
227
228
229
230
231
232
233
234
235
236
237
238
239
240
241
242
243
244
245
246
247
248
249
250
251
252
253
254
255
256
257
258
259
260
261
262
263
264
265
266
267
268
269
270
271
272
273
274
275
276
277
278
279
280
281
282
283
284
285
286
287
288
289
290
291
292
293
294
295
296
297
298
299
300
301
302
303
304
305
306
307
308
309
310
311
312
313
314
315
316
317
318
319
320
321
322
323
324
325
326
327
328
329
330
331
332
333
334
335
336
337
338
339
340
341
342
343
344
345
346
347
348
349
350
351
352
353
354
355
356
357
358
359
360
361
362
363
364
365
366
367
368
369
370
371
372
373
374
375
376
377
378
379
380
381
382
383
384
385
386
387
388
389
390
391
392
393
394
395
396
397
398
399
400
401
402
403
404
405
406
407
408
409
410
411
412
413
414
415
416
417
418
419
420
421
422
423
424
425
426
427
428
429
430
431
432
433
434
435
436
437
438
439
440
441
442
443
444
445
446
447
448
449
450
451
452
453
454
455
456
457
458
459
460
461
462
463
464
465
466
467
468
469
470
471
472
473
474
475
476
477
478
479
480
481
482
483
484
485
486
487
488
489
490
491
492
493
494
495
496
497
498
499
500
501
502
503
504
505
506
507
508
509
510
511
512
513
514
515
516
517
518
519
520
521
522
523
524
525
526
527
528
529
530
531
532
533
534
535
536
537
538
539
540
541
542
543
544
545
546
547
548
549
550
551
552
553
554
555
556
557
558
559
560
561
562
563
564
565
566
567
568
569
570
571
572
573
574
575
576
577
578
579
580
581
582
583
584
585
586
587
588
589
590
591
592
593
594
595
596
597
598
599
600
601
602
603
604
605
606
607
608
609
610
611
612
613
614
615
616
617
618
619
620
621
622
623
624
625
626
627
628
629
630
631
632
633
634
635
636
637
638
639
640
641
642
643
644
645
646
647
648
649
650
651
652
653
654
655
656
657
658
659
660
661
662
663
664
665
666
667
668
669
670
671
672
673
674
675
676
677
678
679
680
681
682
683
684
685
686
687
688
689
690
691
692
693
694
695
696
697
698
699
700
701
702
703
704
705
706
707
708
709
710
711
712
713
714
715
716
717
718
719
720
721
722
723
724
725
726
727
728
729
730
731
732
733
734
735
736
737
738
739
740
741
742
743
744
745
746
747
748
749
750
751
752
753
754
755
756
757
758
759
760
761
762
763
764
765
766
767
768
769
770
771
772
773
774
775
776
777
778
779
780
781
782
783
784
785
786
787
788
789
790
791
792
793
794
795
796
797
798
799
800
801
802
803
804
805
806
807
808
809
810
811
812
813
814
815
816
817
818
819
820
821
822
823
824
825
826
827
828
829
830
831
832
833
834
835
836
837
838
839
840
841
842
843
844
845
846
847
848
849
850
851
852
853
854
855
856
857
858
859
860
861
862
863
864
865
866
867
868
869
870
871
872
873
874
875
876
877
878
879
880
881
882
883
884
885
886
887
888
889
890
891
892
893
894
895
896
897
898
899
900
901
902
903
904
905
906
907
908
909
910
911
912
913
914
915
916
917
918
919
920
921
922
923
924
925
926
927
928
929
930
931
932
933
934
935
936
937
938
939
940
941
942
943
944
945
946
947
948
949
950
951
952
953
954
955
956
957
958
959
960
961
962
963
964
965
966
967
968
969
970
971
972
973
974
975
976
977
978
979
980
981
982
983
984
985
986
987
988
989
990
991
992
993
994
995
996
997
998
999

```

deeplab.py

```

1 import torch
2 import torch.nn as nn

```

```

3 import torch.nn.functional as F
4 from modeling.sync_batchnorm.batchnorm import SynchronizedBatchNorm2d
5 from modeling.aspp import build_aspp
6 from modeling.decoder import build_decoder
7 from modeling.backbone import build_backbone
8
9 class DeepLab(nn.Module):
10     def __init__(self, backbone='resnet', output_stride=16, num_classes=21,
11                 sync_bn=True, freeze_bn=False):
12         super(DeepLab, self).__init__()
13         if backbone == 'drn':
14             output_stride = 8
15
16         if sync_bn == True:
17             BatchNorm = SynchronizedBatchNorm2d
18         else:
19             BatchNorm = nn.BatchNorm2d
20
21         self.backbone = build_backbone(backbone, output_stride, BatchNorm)
22         self.aspp = build_aspp(backbone, output_stride, BatchNorm)
23         self.decoder = build_decoder(num_classes, backbone, BatchNorm)
24
25         self.freeze_bn = freeze_bn
26
27     def forward(self, input):
28         x, low_level_feat = self.backbone(input)
29         x = self.aspp(x)
30         x = self.decoder(x, low_level_feat)
31         x = F.interpolate(x, size=input.size()[2:], mode='bilinear', align_corners=True)
32
33         return x
34
35     def freeze_bn(self):
36         for m in self.modules():
37             if isinstance(m, SynchronizedBatchNorm2d):
38                 m.eval()
39             elif isinstance(m, nn.BatchNorm2d):
40                 m.eval()
41
42     def get_1x_lr_params(self):
43         modules = [self.backbone]
44         for i in range(len(modules)):
45             for m in modules[i].named_modules():
46                 if self.freeze_bn:
47                     if isinstance(m[1], nn.Conv2d):
48                         for p in m[1].parameters():
49                             if p.requires_grad:

```

```

50         yield p
51     else:
52         if isinstance(m[1], nn.Conv2d) or isinstance(m[1], SynchronizedBatchNorm2d) \
53             or isinstance(m[1], nn.BatchNorm2d):
54             for p in m[1].parameters():
55                 if p.requires_grad:
56                     yield p
57
58     def get_10x_lr_params(self):
59         modules = [self.aspp, self.decoder]
60         for i in range(len(modules)):
61             for m in modules[i].named_modules():
62                 if self.freeze_bn:
63                     if isinstance(m[1], nn.Conv2d):
64                         for p in m[1].parameters():
65                             if p.requires_grad:
66                                 yield p
67                 else:
68                     if isinstance(m[1], nn.Conv2d) or isinstance(m[1], SynchronizedBatchNorm2d) \
69                         or isinstance(m[1], nn.BatchNorm2d):
70                         for p in m[1].parameters():
71                             if p.requires_grad:
72                                 yield p
73
74 if __name__ == "__main__":
75     model = DeepLab(backbone='mobilenet', output_stride=16)
76     model.eval()
77     input = torch.rand(1, 3, 513, 513)
78     output = model(input)
79     print(output.size())

```

2.1.5 UNet 源代码

train.py

```

1  import argparse
2  import csv
3  import logging
4  import os
5  from pathlib import Path
6
7  import numpy as np
8  import torch
9  import torch.nn as nn
10 import torch.nn.functional as F
11 import wandb

```

```

12 from torch import optim
13 from torch.utils.data import DataLoader, random_split
14 from tqdm import tqdm
15
16 from evaluate import evaluate
17 from unet import UNet
18 from utils.data_loading import BasicDataset, CarvanaDataset
19 from utils.dice_score import dice_loss
20
21 dir_img = Path('./data/subset1/folder1/')
22 dir_mask = Path('./data/subset1/folder2/')
23 dir_checkpoint = Path('./checkpoints/')
24
25
26 def calculate_iou(confusion_matrix):
27     intersection = confusion_matrix.diagonal()
28     union = confusion_matrix.sum(axis=1) + confusion_matrix.sum(axis=0) - intersection
29     iou = intersection / union
30     return iou
31
32
33 def calculate_miou(confusion_matrix):
34     return np.mean(calculate_iou(confusion_matrix))
35
36
37 def calculate_fwioi(confusion_matrix):
38     frequency = confusion_matrix.sum(axis=1) / confusion_matrix.sum()
39     fwioi = np.sum(frequency * calculate_iou(confusion_matrix))
40     return fwioi
41
42
43
44 def train_model(
45     model,
46     device,
47     epochs: int = 5,
48     batch_size: int = 1,
49     learning_rate: float = 1e-5,
50     val_percent: float = 0.1,
51     save_checkpoint: bool = True,
52     img_scale: float = 0.5,
53     amp: bool = False,
54     weight_decay: float = 1e-8,
55     momentum: float = 0.999,
56     gradient_clipping: float = 1.0,
57 ):
58     # 1. Create dataset

```

```

59     try:
60         dataset = CarvanaDataset(dir_img, dir_mask, img_scale)
61     except (AssertionError, RuntimeError, IndexError):
62         dataset = BasicDataset(dir_img, dir_mask, img_scale)
63
64     # 2. Split into train / validation partitions
65     n_val = int(len(dataset) * val_percent)
66     n_train = len(dataset) - n_val
67     train_set, val_set = random_split(dataset, [n_train, n_val],
68                                       generator=torch.Generator().manual_seed(0))
69
70     # 3. Create data loaders
71     loader_args = dict(batch_size=batch_size, num_workers=0, pin_memory=False)
72     train_loader = DataLoader(train_set, shuffle=True, **loader_args, drop_last=True)
73     val_loader = DataLoader(val_set, shuffle=False, drop_last=True, **loader_args)
74
75     # (Initialize logging)
76     experiment = wandb.init(project='U-Net', resume='allow', anonymous='must')
77     experiment.config.update(
78         dict(epochs=epochs, batch_size=batch_size, learning_rate=learning_rate,
79             val_percent=val_percent, save_checkpoint=save_checkpoint, img_scale=img_scale,
80             amp=amp)
81     )
82
83     logging.info(f'''Starting training:
84     Epochs:          {epochs}
85     Batch size:      {batch_size}
86     Learning rate:   {learning_rate}
87     Training size:   {n_train}
88     Validation size: {n_val}
89     Checkpoints:     {save_checkpoint}
90     Device:          {device.type}
91     Images scaling:  {img_scale}
92     Mixed Precision: {amp}
93 ''')
94
95     # 4. Set up the optimizer, the loss, the learning rate scheduler and the loss scaling for
96     AMP
97     optimizer = optim.RMSprop(model.parameters(),
98                               lr=learning_rate, weight_decay=weight_decay, momentum=momentum,
99                               foreach=True)
100     scheduler = optim.lr_scheduler.ReduceLROnPlateau(optimizer, 'max', patience=5) # goal:
101         maximize Dice score
102     grad_scaler = torch.cuda.amp.GradScaler(enabled=amp)
103     criterion = nn.CrossEntropyLoss() if model.n_classes > 1 else nn.BCEWithLogitsLoss()
104     global_step = 0

```

```

101 # ... (前面的代码保持不变)
102
103 # 5. Begin training
104 validate_every = 1 # Validate every epoch, you can adjust this value
105 for epoch in range(1, epochs + 1):
106     model.train()
107     epoch_loss = 0
108     confusion_matrix = np.zeros((model.n_classes, model.n_classes), dtype=int)
109
110     with tqdm(total=n_train, desc=f'Epoch {epoch}/{epochs}', unit='img') as pbar:
111         for batch in train_loader:
112             images, true_masks = batch['image'], batch['mask']
113
114             assert images.shape[1] == model.n_channels, \
115                 f'Network has been defined with {model.n_channels} input channels, ' \
116                 f'but loaded images have {images.shape[1]} channels. Please check that ' \
117                 'the images are loaded correctly.'
118
119             images = images.to(device=device, dtype=torch.float32,
120                                 memory_format=torch.channels_last)
121             true_masks = true_masks.to(device=device, dtype=torch.long)
122
123             with torch.cuda.amp.autocast(enabled=True):
124                 masks_pred = model(images)
125                 if model.n_classes == 1:
126                     loss = criterion(masks_pred.squeeze(1), true_masks.float())
127                     loss += dice_loss(F.sigmoid(masks_pred.squeeze(1)), true_masks.float(),
128                                       multiclass=False)
129                 else:
130                     loss = criterion(masks_pred, true_masks)
131                     loss += dice_loss(
132                         F.softmax(masks_pred, dim=1).float(),
133                         F.one_hot(true_masks, model.n_classes).permute(0, 3, 1, 2).float(),
134                         multiclass=True
135                     )
136
137             optimizer.zero_grad(set_to_none=True)
138             grad_scaler.scale(loss).backward()
139             torch.nn.utils.clip_grad_norm_(model.parameters(), gradient_clipping)
140             grad_scaler.step(optimizer)
141             grad_scaler.update()
142
143             pbar.update(images.shape[0])
144             global_step += 1
145             epoch_loss += loss.item()
146
147 # Evaluation round

```

```

146     if epoch % validate_every == 0:
147         model.eval()
148         confusion_matrix = np.zeros((model.n_classes, model.n_classes), dtype=int)
149
150     for batch in val_loader:
151         with torch.no_grad():
152             images, true_masks = batch['image'], batch['mask']
153             images = images.to(device=device, dtype=torch.float32,
154                                 memory_format=torch.channels_last)
155             true_masks = true_masks.to(device=device, dtype=torch.long)
156
157             masks_pred = model(images)
158             predicted_masks = masks_pred.argmax(dim=1)
159
160             confusion_matrix += np.bincount((model.n_classes * true_masks.flatten() +
161                                             predicted_masks.flatten()).cpu().numpy(),
162                                             minlength=model.n_classes**2).reshape(model.n_classes, model.n_classes)
163
164             miou = calculate_miou(confusion_matrix)
165             fwiou = calculate_fwiou(confusion_matrix)
166
167             acc = (confusion_matrix.diagonal().sum() / confusion_matrix.sum())
168             acc_class = confusion_matrix.diagonal() / confusion_matrix.sum(axis=1)
169
170             logging.info('Epoch: {}, Acc: {:.4f}, Acc_class: {}, mIoU: {:.4f}, FWIoU:
171                           {:.4f}'.format(epoch, acc,
172                                           acc_class,
173                                           miou,
174                                           fwiou))
175
176             epoch_data = (epoch, acc, acc_class, miou, fwiou)
177
178     if save_checkpoint:
179         Path(dir_checkpoint).mkdir(parents=True, exist_ok=True)
180         state_dict = model.state_dict()
181         state_dict['mask_values'] = dataset.mask_values
182         torch.save(state_dict, str(dir_checkpoint / 'checkpoint_epoch{}.pth'.format(epoch)))
183         logging.info(f'Checkpoint {epoch} saved!')
184         with open('output.csv', 'a', newline='') as csvfile:
185             # 创建CSV写入对象
186             csv_writer = csv.writer(csvfile)
187             # 写入数据
188             csv_writer.writerow(epoch_data)

```

```

189 def get_args():
190     parser = argparse.ArgumentParser(description='Train the UNet on images and target masks')
191     parser.add_argument('--epochs', '-e', metavar='E', type=int, default=50, help='Number of
        epochs')
192     parser.add_argument('--batch-size', '-b', dest='batch_size', metavar='B', type=int,
        default=8, help='Batch size')
193     parser.add_argument('--learning-rate', '-l', metavar='LR', type=float, default=1e-5,
        help='Learning rate', dest='lr')
194     parser.add_argument('--load', '-f', type=str, default=False, help='Load model from a .pth
        file')
195     parser.add_argument('--scale', '-s', type=float, default=0.5, help='Downscaling factor of
        the images')
196     parser.add_argument('--validation', '-v', dest='val', type=float, default=10.0,
        help='Percent of the data that is used as validation (0-100)')
197     parser.add_argument('--amp', action='store_true', default=False, help='Use mixed precision')
198     parser.add_argument('--bilinear', action='store_true', default=False, help='Use bilinear
        upsampling')
199     parser.add_argument('--classes', '-c', type=int, default=2, help='Number of classes')
200
201     return parser.parse_args()
202
203
204
205
206 if __name__ == '__main__':
207     args = get_args()
208
209     logging.basicConfig(level=logging.INFO, format='%(levelname)s: %(message)s')
210     device = torch.device('cuda' if torch.cuda.is_available() else 'cpu')
211     logging.info(f'Using device {device}')
212
213     # Change here to adapt to your data
214     # n_channels=3 for RGB images
215     # n_classes is the number of probabilities you want to get per pixel
216     model = UNet(n_channels=3, n_classes=args.classes, bilinear=args.bilinear)
217     model = model.to(memory_format=torch.channels_last)
218
219     logging.info(f'Network:\n'
220                 f'\t\t{model.n_channels} input channels\n'
221                 f'\t\t{model.n_classes} output channels (classes)\n'
222                 f'\t\t{"Bilinear" if model.bilinear else "Transposed conv"} upsampling')
223
224     if args.load:
225         state_dict = torch.load(args.load, map_location=device)
226         del state_dict['mask_values']
227         model.load_state_dict(state_dict)
228         logging.info(f'Model loaded from {args.load}')
229
230     model.to(device=device)

```

```

231     try:
232         train_model(
233             model=model,
234             epochs=args.epochs,
235             batch_size=args.batch_size,
236             learning_rate=args.lr,
237             device=device,
238             img_scale=args.scale,
239             val_percent=args.val / 100,
240             amp=args.amp
241         )
242     except torch.cuda.OutOfMemoryError:
243         logging.error('Detected OutOfMemoryError! '
244                       'Enabling checkpointing to reduce memory usage, but this slows down
245                           training. '
246                       'Consider enabling AMP (--amp) for fast and memory efficient training')
247         torch.cuda.empty_cache()
248         model.use_checkpointing()
249         train_model(
250             model=model,
251             epochs=args.epochs,
252             batch_size=args.batch_size,
253             learning_rate=args.lr,
254             device=device,
255             img_scale=args.scale,
256             val_percent=args.val / 100,
257             amp=args.amp
258         )

```

unet_model.py

```

1  """ Full assembly of the parts to form the complete network """
2
3  from .unet_parts import *
4
5
6  class UNet(nn.Module):
7      def __init__(self, n_channels, n_classes, bilinear=False):
8          super(UNet, self).__init__()
9          self.n_channels = n_channels
10         self.n_classes = n_classes
11         self.bilinear = bilinear
12
13         self.inc = (DoubleConv(n_channels, 64))
14         self.down1 = (Down(64, 128))
15         self.down2 = (Down(128, 256))
16         self.down3 = (Down(256, 512))

```

```

17     factor = 2 if bilinear else 1
18     self.down4 = (Down(512, 1024 // factor))
19     self.up1 = (Up(1024, 512 // factor, bilinear))
20     self.up2 = (Up(512, 256 // factor, bilinear))
21     self.up3 = (Up(256, 128 // factor, bilinear))
22     self.up4 = (Up(128, 64, bilinear))
23     self.outc = (OutConv(64, n_classes))
24
25     def forward(self, x):
26         x1 = self.inc(x)
27         x2 = self.down1(x1)
28         x3 = self.down2(x2)
29         x4 = self.down3(x3)
30         x5 = self.down4(x4)
31         x = self.up1(x5, x4)
32         x = self.up2(x, x3)
33         x = self.up3(x, x2)
34         x = self.up4(x, x1)
35         logits = self.outc(x)
36         return logits
37
38     def use_checkpointing(self):
39         self.inc = torch.utils.checkpoint(self.inc)
40         self.down1 = torch.utils.checkpoint(self.down1)
41         self.down2 = torch.utils.checkpoint(self.down2)
42         self.down3 = torch.utils.checkpoint(self.down3)
43         self.down4 = torch.utils.checkpoint(self.down4)
44         self.up1 = torch.utils.checkpoint(self.up1)
45         self.up2 = torch.utils.checkpoint(self.up2)
46         self.up3 = torch.utils.checkpoint(self.up3)
47         self.up4 = torch.utils.checkpoint(self.up4)
48         self.outc = torch.utils.checkpoint(self.outc)

```

UNET PARTS

```

1  """ Parts of the U-Net model """
2
3  import torch
4  import torch.nn as nn
5  import torch.nn.functional as F
6
7
8  class DoubleConv(nn.Module):
9      """(convolution => [BN] => ReLU) * 2"""
10
11     def __init__(self, in_channels, out_channels, mid_channels=None):
12         super().__init__()

```

```

13     if not mid_channels:
14         mid_channels = out_channels
15     self.double_conv = nn.Sequential(
16         nn.Conv2d(in_channels, mid_channels, kernel_size=3, padding=1, bias=False),
17         nn.BatchNorm2d(mid_channels),
18         nn.ReLU(inplace=True),
19         nn.Conv2d(mid_channels, out_channels, kernel_size=3, padding=1, bias=False),
20         nn.BatchNorm2d(out_channels),
21         nn.ReLU(inplace=True)
22     )
23
24     def forward(self, x):
25         return self.double_conv(x)
26
27
28 class Down(nn.Module):
29     """Downscaling with maxpool then double conv"""
30
31     def __init__(self, in_channels, out_channels):
32         super().__init__()
33         self.maxpool_conv = nn.Sequential(
34             nn.MaxPool2d(2),
35             DoubleConv(in_channels, out_channels)
36         )
37
38     def forward(self, x):
39         return self.maxpool_conv(x)
40
41
42 class Up(nn.Module):
43     """Upscaling then double conv"""
44
45     def __init__(self, in_channels, out_channels, bilinear=True):
46         super().__init__()
47
48         # if bilinear, use the normal convolutions to reduce the number of channels
49         if bilinear:
50             self.up = nn.Upsample(scale_factor=2, mode='bilinear', align_corners=True)
51             self.conv = DoubleConv(in_channels, out_channels, in_channels // 2)
52         else:
53             self.up = nn.ConvTranspose2d(in_channels, in_channels // 2, kernel_size=2, stride=2)
54             self.conv = DoubleConv(in_channels, out_channels)
55
56     def forward(self, x1, x2):
57         x1 = self.up(x1)
58         # input is CHW
59         diffY = x2.size()[2] - x1.size()[2]

```

```

60     diffX = x2.size()[3] - x1.size()[3]
61
62     x1 = F.pad(x1, [diffX // 2, diffX - diffX // 2,
63                    diffY // 2, diffY - diffY // 2])
64     # if you have padding issues, see
65     #
66     # https://github.com/HaiyongJiang/U-Net-Pytorch-Unstructured-Buggy/commit/0e854509c2cea854e247a9c615f1
67     #
68     # https://github.com/xiaopeng-liao/Pytorch-UNet/commit/8ebac70e633bac59fc22bb5195e513d5832fb3bd
69
70     x = torch.cat([x2, x1], dim=1)
71     return self.conv(x)
72
73 class OutConv(nn.Module):
74     def __init__(self, in_channels, out_channels):
75         super(OutConv, self).__init__()
76         self.conv = nn.Conv2d(in_channels, out_channels, kernel_size=1)
77
78     def forward(self, x):
79         return self.conv(x)

```

2.1.6 SegNet 源代码

train.py

```

1  #!/usr/bin/env python
2  # -*- coding: utf-8 -*-
3  # @Time : 2023/12/15 下午2:11
4  # @Author : SiHang Xu
5  import csv
6  import logging
7  import time
8
9  import cv2
10 import numpy as np
11 from torch.utils import data as Data
12 from MyDataset import MyDataset
13 from SegNet import *
14
15
16 def calculate_iou(confusion_matrix):
17     intersection = confusion_matrix.diagonal()
18     union = confusion_matrix.sum(axis=1) + confusion_matrix.sum(axis=0) - intersection
19     iou = intersection / union
20     return iou
21

```

```

22
23 def calculate_miou(confusion_matrix):
24     return np.mean(calculate_iou(confusion_matrix))
25
26
27 def calculate_fwioi(confusion_matrix):
28     frequency = confusion_matrix.sum(axis=1) / confusion_matrix.sum()
29     fwioi = np.sum(frequency * calculate_iou(confusion_matrix))
30     return fwioi
31
32 epoch_data = []
33 def train_and_validate(SegNet, train_data, val_data, BATCH_SIZE, LR, MOMENTUM, EPOCH,
34     CATE_WEIGHT, WEIGHTS, PRE_TRAINING):
35     SegNet = SegNet(3, 2).cuda()
36     SegNet.load_weights(PRE_TRAINING)
37
38     train_loader = torch.utils.data.DataLoader(train_data, batch_size=BATCH_SIZE, shuffle=True,
39         drop_last=True, num_workers=12)
40     val_loader = torch.utils.data.DataLoader(val_data, batch_size=BATCH_SIZE, shuffle=False,
41         drop_last=True, num_workers=12)
42
43     optimizer = torch.optim.SGD(SegNet.parameters(), lr=LR, momentum=MOMENTUM)
44
45     loss_func =
46         nn.CrossEntropyLoss(weight=torch.from_numpy(np.array(CATE_WEIGHT)).float()).cuda()
47
48     for epoch in range(EPOCH):
49         # Training
50         SegNet.train()
51         for step, (b_x, b_y) in enumerate(train_loader):
52             b_x = b_x.cuda()
53             b_y = b_y.cuda()
54             b_y = b_y.view(BATCH_SIZE, 224, 224)
55             output = SegNet(b_x)
56             loss = loss_func(output, b_y.long())
57             loss = loss.cuda()
58             optimizer.zero_grad()
59             loss.backward()
60             optimizer.step()
61
62         print(f"Epoch: {epoch} || Training Loss: {loss:.4f}")
63
64         # Validation
65         SegNet.eval()
66         confusion_matrix = np.zeros((2, 2), dtype=int)
67         with torch.no_grad():
68             val_loss = 0.0

```

```

65
66     for step, (val_x, val_y) in enumerate(val_loader):
67         val_x = val_x.cuda()
68         val_y = val_y.cuda()
69         val_y = val_y.view(BATCH_SIZE, 224, 224)
70         val_output = SegNet(val_x)
71         val_loss += loss_func(val_output, val_y.long()).item()
72
73         val_output = torch.squeeze(val_output)
74
75         predict = val_output.argmax(dim=1)
76         target = val_y
77         confusion_matrix += np.bincount(
78             (2 * target.flatten() + predict.flatten()).cpu().numpy(),
79             minlength=2 ** 2).reshape(2, 2)
80
81         miou = calculate_miou(confusion_matrix)
82         fwiou = calculate_fwiou(confusion_matrix)
83
84         acc = (confusion_matrix.diagonal().sum() / confusion_matrix.sum())
85         acc_class = confusion_matrix.diagonal() / confusion_matrix.sum(axis=1)
86         print('Epoch: {}, Acc: {:.4f}, Acc_class: {}, mIoU: {:.4f}, FWIoU:
87                                     acc_class,
88                                     miou,
89                                     fwiou))
89                                     {:.4f}'.format(epoch, acc,
90
91                                     epoch_data.append((epoch, acc, acc_class, miou, fwiou))
92
93     torch.save(SegNet.state_dict(), WEIGHTS + "SegNet_weights" + str(time.time()) + ".pth")
94     with open('output.csv', 'w', newline='') as csvfile:
95         # 创建CSV写入对象
96         csv_writer = csv.writer(csvfile)
97
98         # 写入表头
99         csv_writer.writerow(['Epoch', 'Acc', 'Acc_class', 'mIoU', 'FWIoU'])
100
101         # 写入数据
102         csv_writer.writerows(epoch_data)
103
104 # 示例调用
105 BATCH_SIZE = 8
106 LR = 0.001
107 MOMENTUM = 0.9
108 EPOCH = 50
109 CATE_WEIGHT = [1, 1] # 根据实际情况调整
110 WEIGHTS = "weights/"
111 PRE_TRAINING = "vgg16_bn-6c64b313.pth"

```

```

111 train_data = MyDataset(txt_path='train.txt')
112 val_data = MyDataset(txt_path='trainval.txt')
113
114 train_and_validate(SegNet, train_data, val_data, BATCH_SIZE, LR, MOMENTUM, EPOCH, CATE_WEIGHT,
    WEIGHTS, PRE_TRAINING)

```

MyDataset.py

```

1  #!/usr/bin/env python
2  # -*- coding: utf-8 -*-
3  # @Time : 2023/12/15 下午2:09
4  # @Author : SiHang Xu
5  import numpy as np
6  import torch
7  from PIL import Image
8  from torch.utils.data import Dataset
9  import cv2
10
11 class MyDataset(Dataset):
12     def __init__(self, txt_path):
13         super(MyDataset, self).__init__()
14
15         with open(txt_path, "r") as file:
16             image_label = [line.strip() for line in file]
17
18         self.image_label = image_label
19
20     def __getitem__(self, item):
21         root = self.image_label[item]
22         image_path = f'./400dataset/JPEGImages/{root}.png'
23         label_path = f'./400dataset/SegmentationClass/{root}.png'
24
25         image = cv2.imread(image_path)
26         image = cv2.resize(image, (224, 224))
27         image = cv2.cvtColor(image, cv2.COLOR_BGR2RGB) # Convert to RGB
28         image = image / 255.0 # Normalize input
29         image = torch.tensor(image, dtype=torch.float32)
30         image = image.permute(2, 0, 1)
31
32         label = Image.open(label_path)
33         label = label.resize((224, 224), Image.NEAREST) # 使用最近邻插值
34         label = np.array(label)
35         label = torch.tensor(label, dtype=torch.long)
36         # label = cv2.imread(label_path, cv2.IMREAD_GRAYSCALE)
37         # label = cv2.resize(label, (224, 224))
38         # label = torch.tensor(label, dtype=torch.long) # Assuming segmentation labels are
            integers

```

```
39
40     return image, label
41
42     def __len__(self):
43         return len(self.image_label)
```

SegNet.pt

```
1  #!/usr/bin/env python
2  # -*- coding: utf-8 -*-
3  # @Time : 2023/12/15 下午2:06
4  # @Author : SiHang Xu
5  import torch
6  from torch import nn
7  import torch.nn.functional as F
8
9
10 class Encoder(nn.Module):
11     def __init__(self, input_channels):
12         super(Encoder, self).__init__()
13
14         self.enco1 = nn.Sequential(
15             nn.Conv2d(input_channels, 64, kernel_size=3, stride=1, padding=1),
16             nn.BatchNorm2d(64),
17             nn.ReLU(),
18             nn.Conv2d(64, 64, kernel_size=3, stride=1, padding=1),
19             nn.BatchNorm2d(64),
20             nn.ReLU()
21         )
22         self.enco2 = nn.Sequential(
23             nn.Conv2d(64, 128, kernel_size=3, stride=1, padding=1),
24             nn.BatchNorm2d(128),
25             nn.ReLU(),
26             nn.Conv2d(128, 128, kernel_size=3, stride=1, padding=1),
27             nn.BatchNorm2d(128),
28             nn.ReLU()
29         )
30         self.enco3 = nn.Sequential(
31             nn.Conv2d(128, 256, kernel_size=3, stride=1, padding=1),
32             nn.BatchNorm2d(256),
33             nn.ReLU(),
34             nn.Conv2d(256, 256, kernel_size=3, stride=1, padding=1),
35             nn.BatchNorm2d(256),
36             nn.ReLU(),
37             nn.Conv2d(256, 256, kernel_size=3, stride=1, padding=1),
38             nn.BatchNorm2d(256),
39             nn.ReLU()
```

```

40     )
41     self.enco4 = nn.Sequential(
42         nn.Conv2d(256, 512, kernel_size=3, stride=1, padding=1),
43         nn.BatchNorm2d(512),
44         nn.ReLU(),
45         nn.Conv2d(512, 512, kernel_size=3, stride=1, padding=1),
46         nn.BatchNorm2d(512),
47         nn.ReLU(),
48         nn.Conv2d(512, 512, kernel_size=3, stride=1, padding=1),
49         nn.BatchNorm2d(512),
50         nn.ReLU()
51     )
52     self.enco5 = nn.Sequential(
53         nn.Conv2d(512, 512, kernel_size=3, stride=1, padding=1),
54         nn.BatchNorm2d(512),
55         nn.ReLU(),
56         nn.Conv2d(512, 512, kernel_size=3, stride=1, padding=1),
57         nn.BatchNorm2d(512),
58         nn.ReLU(),
59         nn.Conv2d(512, 512, kernel_size=3, stride=1, padding=1),
60         nn.BatchNorm2d(512),
61         nn.ReLU()
62     )
63
64     def forward(self, x):
65         id = []
66
67         x = self.enco1(x)
68         x, id1 = F.max_pool2d(x, kernel_size=2, stride=2, return_indices=True) # 保留最大值的位置
69         id.append(id1)
70         x = self.enco2(x)
71         x, id2 = F.max_pool2d(x, kernel_size=2, stride=2, return_indices=True)
72         id.append(id2)
73         x = self.enco3(x)
74         x, id3 = F.max_pool2d(x, kernel_size=2, stride=2, return_indices=True)
75         id.append(id3)
76         x = self.enco4(x)
77         x, id4 = F.max_pool2d(x, kernel_size=2, stride=2, return_indices=True)
78         id.append(id4)
79         x = self.enco5(x)
80         x, id5 = F.max_pool2d(x, kernel_size=2, stride=2, return_indices=True)
81         id.append(id5)
82
83         return x, id
84
85
86     class SegNet(nn.Module):

```

```

87 def __init__(self, input_channels, output_channels):
88     super(SegNet, self).__init__()
89
90     self.weights_new = self.state_dict()
91     self.encoder = Encoder(input_channels)
92
93     self.deco1 = nn.Sequential(
94         nn.Conv2d(512, 512, kernel_size=3, stride=1, padding=1),
95         nn.BatchNorm2d(512),
96         nn.ReLU(),
97         nn.Conv2d(512, 512, kernel_size=3, stride=1, padding=1),
98         nn.BatchNorm2d(512),
99         nn.ReLU(),
100        nn.Conv2d(512, 512, kernel_size=3, stride=1, padding=1),
101        nn.BatchNorm2d(512),
102        nn.ReLU()
103    )
104    self.deco2 = nn.Sequential(
105        nn.Conv2d(512, 512, kernel_size=3, stride=1, padding=1),
106        nn.BatchNorm2d(512),
107        nn.ReLU(),
108        nn.Conv2d(512, 512, kernel_size=3, stride=1, padding=1),
109        nn.BatchNorm2d(512),
110        nn.ReLU(),
111        nn.Conv2d(512, 256, kernel_size=3, stride=1, padding=1),
112        nn.BatchNorm2d(256),
113        nn.ReLU()
114    )
115    self.deco3 = nn.Sequential(
116        nn.Conv2d(256, 256, kernel_size=3, stride=1, padding=1),
117        nn.BatchNorm2d(256),
118        nn.ReLU(),
119        nn.Conv2d(256, 256, kernel_size=3, stride=1, padding=1),
120        nn.BatchNorm2d(256),
121        nn.ReLU(),
122        nn.Conv2d(256, 128, kernel_size=3, stride=1, padding=1),
123        nn.BatchNorm2d(128),
124        nn.ReLU()
125    )
126    self.deco4 = nn.Sequential(
127        nn.Conv2d(128, 128, kernel_size=3, stride=1, padding=1),
128        nn.BatchNorm2d(128),
129        nn.ReLU(),
130        nn.Conv2d(128, 64, kernel_size=3, stride=1, padding=1),
131        nn.BatchNorm2d(64),
132        nn.ReLU()
133    )

```

```

134     self.deco5 = nn.Sequential(
135         nn.Conv2d(64, 64, kernel_size=3, stride=1, padding=1),
136         nn.BatchNorm2d(64),
137         nn.ReLU(),
138         nn.Conv2d(64, output_channels, kernel_size=3, stride=1, padding=1),
139     )
140
141     def forward(self, x):
142         x, id = self.encoder(x)
143
144         x = F.max_unpool2d(x, id[4], kernel_size=2, stride=2)
145         x = self.deco1(x)
146         x = F.max_unpool2d(x, id[3], kernel_size=2, stride=2)
147         x = self.deco2(x)
148         x = F.max_unpool2d(x, id[2], kernel_size=2, stride=2)
149         x = self.deco3(x)
150         x = F.max_unpool2d(x, id[1], kernel_size=2, stride=2)
151         x = self.deco4(x)
152         x = F.max_unpool2d(x, id[0], kernel_size=2, stride=2)
153         x = self.deco5(x)
154
155         return x
156
157     def load_weights(self, weights_path):
158         weights = torch.load(weights_path)
159         del weights["classifier.0.weight"]
160         del weights["classifier.0.bias"]
161         del weights["classifier.3.weight"]
162         del weights["classifier.3.bias"]
163         del weights["classifier.6.weight"]
164         del weights["classifier.6.bias"]
165
166         names = []
167         for key, value in self.encoder.state_dict().items():
168             if "num_batches_tracked" in key:
169                 continue
170             names.append(key)
171
172         for name, dict in zip(names, weights.items()):
173             self.weights_new[name] = dict[1]
174
175         self.encoder.load_state_dict(self.weights_new)

```

2.1.7 边缘检测与目标检测

edge.py

```
1 import cv2
2 import os
3 import numpy as np
4
5 def laplacian_edge_detection(image_path, alpha=3, beta=30):
6     # Read the image
7     original_image = cv2.imread("images/" + image_path, cv2.IMREAD_GRAYSCALE)
8     original_images = cv2.imread("images/" + image_path)
9
10    # Apply GaussianBlur to reduce noise and help with edge detection
11    blurred_image = cv2.GaussianBlur(original_image, (5, 5), 0)
12
13    # Apply Laplacian edge detection
14    laplacian = cv2.Laplacian(blurred_image, cv2.CV_64F)
15
16    # Convert the Laplacian result to uint8 (8-bit) for display
17    laplacian = np.uint8(np.absolute(laplacian))
18
19    # Adjust brightness and contrast
20    laplacian_adjusted = cv2.convertScaleAbs(laplacian, alpha=alpha, beta=beta)
21
22    # Display the original and adjusted Laplacian edge-detected images
23    cv2.imwrite(os.path.join('output', image_path), laplacian_adjusted)
24    cv2.waitKey(0)
25    cv2.destroyAllWindows()
26
27    # Replace 'path/to/your/image.jpg' with the actual path to your image file
28    for image in os.listdir("images"):
29        laplacian_edge_detection(image)
```

detect.py

```
1 import argparse
2 import os
3 import platform
4 import shutil
5 import time
6 from pathlib import Path
7
8 import cv2
9 import torch
10 import torch.backends.cudnn as cudnn
11 import torchvision.utils
```

```

12 from numpy import random
13
14 from models.experimental import attempt_load
15 from utils.datasets import LoadStreams, LoadImages
16 from utils.general import (
17     check_img_size, non_max_suppression, apply_classifier, scale_coords,
18     xyxy2xywh, plot_one_box, strip_optimizer, set_logging)
19 from utils.torch_utils import select_device, load_classifier, time_synchronized
20
21
22 save_csv = True
23 csv_f = open("results.csv", "w")
24
25
26 def detect(save_img=False):
27     out, source, weights, view_img, save_txt, imgsz = \
28         opt.output, opt.source, opt.weights, opt.view_img, opt.save_txt, opt.img_size
29     webcam = source.isnumeric() or source.startswith(('rtsp://', 'rtmp://', 'http://')) or
30         source.endswith('.txt')
31
32     # Initialize
33     set_logging()
34     device = select_device(opt.device)
35
36     if os.path.exists(out):
37         shutil.rmtree(out) # delete output folder
38     os.makedirs(out) # make new output folder
39     half = device.type != 'cpu' # half precision only supported on CUDA
40
41     # Load model
42     model = attempt_load(weights, map_location=device) # load FP32 model
43     imgsz = check_img_size(imgsz, s=model.stride.max()) # check img_size
44     if half:
45         model.half() # to FP16
46
47     # Second-stage classifier
48     classify = False
49     if classify:
50         modelc = load_classifier(name='resnet101', n=2) # initialize
51         modelc.load_state_dict(torch.load('weights/resnet101.pt', map_location=device)['model'])
52         # load weights
53         modelc.to(device).eval()
54
55     # Set Dataloader
56     vid_path, vid_writer = None, None
57     if webcam:
58         view_img = True

```

```

57     cudnn.benchmark = True # set True to speed up constant image size inference
58     dataset = LoadStreams(source, img_size=imsgsz)
59 else:
60     save_img = True
61     dataset = LoadImages(source, img_size=imsgsz)
62
63 # Get names and colors
64 names = model.module.names if hasattr(model, 'module') else model.names
65 colors = [[random.randint(0, 255) for _ in range(3)] for _ in range(len(names))]
66
67 # Run inference
68 t0 = time.time()
69 img = torch.zeros((1, 3, imsgsz, imsgsz), device=device) # init img
70 _ = model(img.half() if half else img) if device.type != 'cpu' else None # run once
71 for path, img, im0s, vid_cap in dataset:
72     img = torch.from_numpy(img).to(device)
73     img = img.half() if half else img.float() # uint8 to fp16/32
74     img /= 255.0 # 0 - 255 to 0.0 - 1.0
75     if img.ndimension() == 3:
76         img = img.unsqueeze(0)
77
78 # Inference
79 t1 = time_synchronized()
80 pred = model(img, augment=opt.augment)[0]
81
82 # Apply NMS
83 pred = non_max_suppression(pred, opt.conf_thres, opt.iou_thres, classes=opt.classes,
84                             agnostic=opt.agnostic_nms)
85
86 # Apply Classifier
87 if classify:
88     pred = apply_classifier(pred, modelc, img, im0s)
89
90 # Process detections
91 for i, det in enumerate(pred): # detections per image
92     if save_csv:
93         csv_f.write(os.path.basename(path)+",")
94     if webcam: # batch_size >= 1
95         p, s, im0 = path[i], '%g: ' % i, im0s[i].copy()
96     else:
97         p, s, im0 = path, '', im0s
98
99     save_path = str(Path(out) / Path(p).name)
100    txt_path = str(Path(out) / Path(p).stem) + ('_%g' % dataset.frame if dataset.mode ==
101        'video' else '')
    s += '%gx%g ' % img.shape[2:] # print string

```

```

102 gn = torch.tensor(im0.shape)[[1, 0, 1, 0]] # normalization gain whwh
103 if det is not None and len(det):
104     # Rescale boxes from img_size to im0 size
105     det[:, :4] = scale_coords(img.shape[2:], det[:, :4], im0.shape).round()
106
107     # Print results
108     for c in det[:, -1].unique():
109         n = (det[:, -1] == c).sum() # detections per class
110         s += '%g %ss, ' % (n, names[int(c)]) # add to string
111
112     # Write results
113     for *xyxy, conf, cls in reversed(det):
114         if save_txt: # Write to file
115             xywh = (xyxy2xywh(torch.tensor(xyxy).view(1, 4)) / gn).view(-1).tolist() #
116                 normalized xywh
117             with open(txt_path + '.txt', 'a') as f:
118                 f.write('%g ' * 5 + '\n' % (cls, *xywh)) # label format
119
120             if save_img or view_img: # Add bbox to image
121                 label = '%s %.2f' % (names[int(cls)], conf)
122                 plot_one_box(xyxy, im0, label=label, color=colors[int(cls)],
123                             line_thickness=3)
124
125             if save_csv:
126                 csv_f.write("{} {} {} {} {}
127                     ".format(str(int(cls.detach().cpu().numpy()+1),
128                               str(int(xyxy[0].detach().cpu().numpy()),
129                               str(int(xyxy[1].detach().cpu().numpy()),
130                               str(int(xyxy[2].detach().cpu().numpy()),
131                               str(int(xyxy[3].detach().cpu().numpy()))))
132
133             csv_f.write("\n")
134
135     # Print time (inference + NMS)
136     print('%sDone. (%.3fs)' % (s, t2 - t1))
137
138     # Stream results
139     if view_img:
140         cv2.imshow(p, im0)
141         if cv2.waitKey(1) == ord('q'): # q to quit
142             raise StopIteration
143
144     # Save results (image with detections)
145     if save_img:
146         if dataset.mode == 'images':
147             cv2.imwrite(save_path, im0)
148         else:
149             if vid_path != save_path: # new video

```

```

142         vid_path = save_path
143         if isinstance(vid_writer, cv2.VideoWriter):
144             vid_writer.release() # release previous video writer
145
146         fourcc = 'mp4v' # output video codec
147         fps = vid_cap.get(cv2.CAP_PROP_FPS)
148         w = int(vid_cap.get(cv2.CAP_PROP_FRAME_WIDTH))
149         h = int(vid_cap.get(cv2.CAP_PROP_FRAME_HEIGHT))
150         vid_writer = cv2.VideoWriter(save_path, cv2.VideoWriter_fourcc(*fourcc),
151                                     fps, (w, h))
152         vid_writer.write(im0)
153
154     if save_txt or save_img:
155         print('Results saved to %s' % Path(out))
156         if platform.system() == 'Darwin' and not opt.update: # MacOS
157             os.system('open ' + save_path)
158
159     print('Done. (%.3fs)' % (time.time() - t0))
160
161 if __name__ == '__main__':
162     parser = argparse.ArgumentParser()
163     parser.add_argument('--weights', nargs='+', type=str,
164                         default='weights/IMSC/last_120_640_32_aug2.pt', help='model.pt path(s)')
165     parser.add_argument('--source', type=str, default='inference/images', help='source') #
166         file/folder, 0 for webcam
167     parser.add_argument('--output', type=str, default='inference/output', help='output folder')
168         # output folder
169     parser.add_argument('--img-size', type=int, default=640, help='inference size (pixels)')
170     parser.add_argument('--conf-thres', type=float, default=0.4, help='object confidence
171         threshold')
172     parser.add_argument('--iou-thres', type=float, default=0.5, help='IOU threshold for NMS')
173     parser.add_argument('--device', default='cpu', help='cuda device, i.e. 0 or 0,1,2,3 or cpu')
174     parser.add_argument('--view-img', action='store_true', help='display results')
175     parser.add_argument('--save-txt', action='store_true', help='save results to *.txt')
176     parser.add_argument('--classes', nargs='+', type=int, help='filter by class: --class 0, or
177         --class 0 2 3')
178     parser.add_argument('--agnostic-nms', action='store_true', help='class-agnostic NMS')
179     parser.add_argument('--augment', action='store_true', help='augmented inference')
180     parser.add_argument('--update', action='store_true', help='update all models')
181     opt = parser.parse_args()
182     print(opt)
183
184     with torch.no_grad():
185         if opt.update: # update all models (to fix SourceChangeWarning)
186             for opt.weights in ['yolov5s.pt', 'yolov5m.pt', 'yolov5l.pt', 'yolov5x.pt']:
187                 detect()

```

```

183         strip_optimizer(opt.weights)
184     else:
185         detect()
186 csv_f.close()

```

yolo.py

```

1  import argparse
2  import logging
3  import math
4  from copy import deepcopy
5  from pathlib import Path
6
7  import torch
8  import torch.nn as nn
9
10 from models.common import Conv, Bottleneck, SPP, DWConv, Focus, BottleneckCSP, Concat, NMS
11 from models.experimental import MixConv2d, CrossConv, C3
12 from utils.general import check_anchor_order, make_divisible, check_file, set_logging
13 from utils.torch_utils import (
14     time_synchronized, fuse_conv_and_bn, model_info, scale_img, initialize_weights,
15     select_device)
16
17 logger = logging.getLogger(__name__)
18
19 class Detect(nn.Module):
20     stride = None # strides computed during build
21     export = False # onnx export
22
23     def __init__(self, nc=80, anchors=(), ch=()): # detection layer
24         super(Detect, self).__init__()
25         self.nc = nc # number of classes
26         self.no = nc + 5 # number of outputs per anchor
27         self.nl = len(anchors) # number of detection layers
28         self.na = len(anchors[0]) // 2 # number of anchors
29         self.grid = [torch.zeros(1)] * self.nl # init grid
30         a = torch.tensor(anchors).float().view(self.nl, -1, 2)
31         self.register_buffer('anchors', a) # shape(nl,na,2)
32         self.register_buffer('anchor_grid', a.clone().view(self.nl, 1, -1, 1, 1, 2)) #
33             shape(nl,1,na,1,1,2)
34         self.m = nn.ModuleList(nn.Conv2d(x, self.no * self.na, 1) for x in ch) # output conv
35
36     def forward(self, x):
37         # x = x.copy() # for profiling
38         z = [] # inference output
39         self.training |= self.export

```

```

39     for i in range(self.nl):
40         x[i] = self.m[i](x[i]) # conv
41         bs, _, ny, nx = x[i].shape # x(bs,255,20,20) to x(bs,3,20,20,85)
42         x[i] = x[i].view(bs, self.na, self.no, ny, nx).permute(0, 1, 3, 4, 2).contiguous()
43
44         if not self.training: # inference
45             if self.grid[i].shape[2:4] != x[i].shape[2:4]:
46                 self.grid[i] = self._make_grid(nx, ny).to(x[i].device)
47
48             y = x[i].sigmoid()
49             y[..., 0:2] = (y[..., 0:2] * 2. - 0.5 + self.grid[i].to(x[i].device)) *
                self.stride[i] # xy
50             y[..., 2:4] = (y[..., 2:4] * 2) ** 2 * self.anchor_grid[i] # wh
51             z.append(y.view(bs, -1, self.no))
52
53     return x if self.training else (torch.cat(z, 1), x)
54
55     @staticmethod
56     def _make_grid(nx=20, ny=20):
57         yv, xv = torch.meshgrid([torch.arange(ny), torch.arange(nx)])
58         return torch.stack((xv, yv), 2).view((1, 1, ny, nx, 2)).float()
59
60
61 class Model(nn.Module):
62     def __init__(self, cfg='yolov5s.yaml', ch=3, nc=None): # model, input channels, number of
        classes
63         super(Model, self).__init__()
64         if isinstance(cfg, dict):
65             self.yaml = cfg # model dict
66         else: # is *.yaml
67             import yaml # for torch hub
68             self.yaml_file = Path(cfg).name
69             with open(cfg) as f:
70                 self.yaml = yaml.load(f, Loader=yaml.FullLoader) # model dict
71
72         # Define model
73         if nc and nc != self.yaml['nc']:
74             print('Overriding model.yaml nc=%g with nc=%g' % (self.yaml['nc'], nc))
75             self.yaml['nc'] = nc # override yaml value
76         self.model, self.save = parse_model(deepcopy(self.yaml), ch=[ch]) # model, savelist,
            ch_out
77         # print([x.shape for x in self.forward(torch.zeros(1, ch, 64, 64))])
78
79         # Build strides, anchors
80         m = self.model[-1] # Detect()
81         if isinstance(m, Detect):
82             s = 128 # 2x min stride

```

```

83         m.stride = torch.tensor([s / x.shape[-2] for x in self.forward(torch.zeros(1, ch, s,
84                                     s))] # forward
85         m.anchors /= m.stride.view(-1, 1, 1)
86         check_anchor_order(m)
87         self.stride = m.stride
88         self._initialize_biases() # only run once
89         # print('Strides: %s' % m.stride.tolist())
90
91     # Init weights, biases
92     initialize_weights(self)
93     self.info()
94     print('')
95
96 def forward(self, x, augment=False, profile=False):
97     if augment:
98         img_size = x.shape[-2:] # height, width
99         s = [1, 0.83, 0.67] # scales
100        f = [None, 3, None] # flips (2-ud, 3-lr)
101        y = [] # outputs
102        for si, fi in zip(s, f):
103            xi = scale_img(x.flip(fi) if fi else x, si)
104            yi = self.forward_once(xi)[0] # forward
105            # cv2.imwrite('img%g.jpg' % s, 255 * xi[0].numpy().transpose((1, 2, 0))[:, :,
106            # ::-1]) # save
107            yi[..., :4] /= si # de-scale
108            if fi == 2:
109                yi[..., 1] = img_size[0] - yi[..., 1] # de-flip ud
110            elif fi == 3:
111                yi[..., 0] = img_size[1] - yi[..., 0] # de-flip lr
112            y.append(yi)
113        return torch.cat(y, 1), None # augmented inference, train
114    else:
115        return self.forward_once(x, profile) # single-scale inference, train
116
117 def forward_once(self, x, profile=False):
118     y, dt = [], [] # outputs
119     for m in self.model:
120         if m.f != -1: # if not from previous layer
121             x = y[m.f] if isinstance(m.f, int) else [x if j == -1 else y[j] for j in m.f] #
122             from earlier layers
123
124         if profile:
125             try:
126                 import thop
127                 o = thop.profile(m, inputs=(x,), verbose=False)[0] / 1E9 * 2 # FLOPS
128             except:
129                 o = 0

```

```

127         t = time_synchronized()
128         for _ in range(10):
129             _ = m(x)
130         dt.append((time_synchronized() - t) * 100)
131         print('%10.1f%10.0f%10.1fms %-40s' % (o, m.np, dt[-1], m.type))
132
133     x = m(x) # run
134     y.append(x if m.i in self.save else None) # save output
135
136     if profile:
137         print('%10.1fms total' % sum(dt))
138     return x
139
140 def _initialize_biases(self, cf=None): # initialize biases into Detect(), cf is class
141     frequency
142     # cf = torch.bincount(torch.tensor(np.concatenate(dataset.labels, 0)[: , 0]).long(),
143     minlength=nc) + 1.
144     m = self.model[-1] # Detect() module
145     for mi, s in zip(m.m, m.stride): # from
146         b = mi.bias.view(m.na, -1) # conv.bias(255) to (3,85)
147         b[:, 4] += math.log(8 / (640 / s) ** 2) # obj (8 objects per 640 image)
148         b[:, 5:] += math.log(0.6 / (m.nc - 0.99)) if cf is None else torch.log(cf /
149             cf.sum()) # cls
150         mi.bias = torch.nn.Parameter(b.view(-1), requires_grad=True)
151
152 def _print_biases(self):
153     m = self.model[-1] # Detect() module
154     for mi in m.m: # from
155         b = mi.bias.detach().view(m.na, -1).T # conv.bias(255) to (3,85)
156         print(('6g Conv2d.bias:' + '%10.3g' * 6) % (mi.weight.shape[1],
157             *b[:5].mean(1).tolist(), b[5:].mean()))
158
159 def fuse(self): # fuse model Conv2d() + BatchNorm2d() layers
160     print('Fusing layers... ')
161     for m in self.model.modules():
162         if type(m) is Conv and hasattr(Conv, 'bn'):
163             m._non_persistent_buffers_set = set() # pytorch 1.6.0 compatability
164             m.conv = fuse_conv_and_bn(m.conv, m.bn) # update conv
165             delattr(m, 'bn') # remove batchnorm
166             m.forward = m.fuseforward # update forward
167     self.info()
168     return self
169
170 def add_nms(self): # fuse model Conv2d() + BatchNorm2d() layers
171     if type(self.model[-1]) is not NMS: # if missing NMS
172         print('Adding NMS module... ')
173         m = NMS() # module

```

```

170         m.f = -1 # from
171         m.i = self.model[-1].i + 1 # index
172         self.model.add_module(name='%s' % m.i, module=m) # add
173     return self
174
175     def info(self, verbose=False): # print model information
176         model_info(self, verbose)
177
178
179     def parse_model(d, ch): # model_dict, input_channels(3)
180         logger.info('\n%3s%18s%3s%10s %-40s%-30s' % ('', 'from', 'n', 'params', 'module',
181             'arguments'))
182         anchors, nc, gd, gw = d['anchors'], d['nc'], d['depth_multiple'], d['width_multiple']
183         na = (len(anchors[0]) // 2) if isinstance(anchors, list) else anchors # number of anchors
184         no = na * (nc + 5) # number of outputs = anchors * (classes + 5)
185
186         layers, save, c2 = [], [], ch[-1] # layers, savelist, ch out
187         for i, (f, n, m, args) in enumerate(d['backbone'] + d['head']): # from, number, module, args
188             m = eval(m) if isinstance(m, str) else m # eval strings
189             for j, a in enumerate(args):
190                 try:
191                     args[j] = eval(a) if isinstance(a, str) else a # eval strings
192                 except:
193                     pass
194
195             n = max(round(n * gd), 1) if n > 1 else n # depth gain
196             if m in [Conv, Bottleneck, SPP, DWConv, MixConv2d, Focus, CrossConv, BottleneckCSP, C3]:
197                 c1, c2 = ch[f], args[0]
198
199                 # Normal
200                 # if i > 0 and args[0] != no: # channel expansion factor
201                 #     ex = 1.75 # exponential (default 2.0)
202                 #     e = math.log(c2 / ch[1]) / math.log(2)
203                 #     c2 = int(ch[1] * ex ** e)
204                 # if m != Focus:
205
206                 c2 = make_divisible(c2 * gw, 8) if c2 != no else c2
207
208                 # Experimental
209                 # if i > 0 and args[0] != no: # channel expansion factor
210                 #     ex = 1 + gw # exponential (default 2.0)
211                 #     ch1 = 32 # ch[1]
212                 #     e = math.log(c2 / ch1) / math.log(2) # level 1-n
213                 #     c2 = int(ch1 * ex ** e)
214                 # if m != Focus:
215                 #     c2 = make_divisible(c2, 8) if c2 != no else c2

```

```

216         args = [c1, c2, *args[1:]]
217         if m in [BottleneckCSP, C3]:
218             args.insert(2, n)
219             n = 1
220         elif m is nn.BatchNorm2d:
221             args = [ch[f]]
222         elif m is Concat:
223             c2 = sum([ch[-1 if x == -1 else x + 1] for x in f])
224         elif m is Detect:
225             args.append([ch[x + 1] for x in f])
226             if isinstance(args[1], int): # number of anchors
227                 args[1] = [list(range(args[1] * 2))] * len(f)
228         else:
229             c2 = ch[f]
230
231         m_ = nn.Sequential(*[m(*args) for _ in range(n)] if n > 1 else m(*args) # module
232         t = str(m)[8:-2].replace('__main__', '') # module type
233         np = sum([x.numel() for x in m_.parameters()]) # number params
234         m_.i, m_.f, m_.type, m_.np = i, f, t, np # attach index, 'from' index, type, number
235             params
236         logger.info('%3s%18s%3s%10.0f %-40s%-30s' % (i, f, n, np, t, args)) # print
237         save.extend(x % i for x in ([f] if isinstance(f, int) else f) if x != -1) # append to
238             savelist
239         layers.append(m_)
240         ch.append(c2)
241
242     return nn.Sequential(*layers), sorted(save)
243
244 if __name__ == '__main__':
245     parser = argparse.ArgumentParser()
246     parser.add_argument('--cfg', type=str, default='yolov5s.yaml', help='model.yaml')
247     parser.add_argument('--device', default='', help='cuda device, i.e. 0 or 0,1,2,3 or cpu')
248     opt = parser.parse_args()
249     opt.cfg = check_file(opt.cfg) # check file
250     set_logging()
251     device = select_device(opt.device)
252
253     # Create model
254     model = Model(opt.cfg).to(device)
255     model.train()

```

2.2 问题二源代码

2.2.1 坑洼图像筛选源代码

classification.py

```

1  import os
2  import shutil
3
4  import torch
5  from torchvision.transforms import transforms
6  import cv2
7  from models.SEnet import *
8  import csv
9
10 # model = SE_ResNet18().cuda().eval()
11 #
12 # model.load_state_dict(torch.load('SENet.pt'))
13
14 model = torch.load('best.pt').cuda().eval()
15
16 data_root = './data_test_V2'
17
18 transform_test = transforms.Compose([
19     transforms.ToPILImage(),
20     transforms.Resize((224, 224)),
21     transforms.ToTensor(),
22 ])
23
24
25 # 初始化一个CSV文件
26 csv_file = 'test_result1.csv'
27
28 # 打开CSV文件以写入结果
29 with open(csv_file, mode='w', newline='') as file:
30     fieldnames = ['fname', 'label']
31     writer = csv.DictWriter(file, fieldnames=fieldnames)
32     writer.writeheader()
33
34 # 遍历数据根目录中的文件/子目录
35 for i in os.listdir(data_root):
36     image_root = os.path.join(data_root, i)
37     image = cv2.imread(image_root)
38     image = transform_test(image)
39     image = image.unsqueeze(0)
40     image = image.cuda()
41     output = model(image)
42     _, predicted = torch.max(output, 1)
43
44 # 设置label值
45 label = 1 if predicted.item() == 0 else 0
46

```

```

47     # 写入CSV文件
48     writer.writerow({'fname': i, 'label': label})
49     if label == 1:
50         shutil.copyfile(os.path.join(data_root, i), os.path.join('./normal/',i))
51     else:
52         shutil.copyfile(os.path.join(data_root, i), os.path.join('./potholes/',i))
53 print("CSV文件已创建并写入结果。")

```

2.2.2 填写结果文件源代码

fill_normal.py

```

1  #!/usr/bin/env python
2  # -*- coding: utf-8 -*-
3  # @Time : 2023/12/13 下午7:23
4  # @Author : SiHang Xu
5  import csv
6  import os
7
8  file_name = 'test_result2.csv'
9  # 指定结果文件的路径
10 csv_file_path = "test_result2.csv"
11
12 # 打开CSV文件, 使用 'a' 模式
13 with open(csv_file_path, 'a', newline='') as csvfile:
14     # 创建CSV写入对象
15     csv_writer = csv.writer(csvfile)
16
17
18 # 遍历目录中的文件
19 for i in os.listdir('./dataset/test_data/normal'):
20     image_path = './dataset/test_data/normal/' + i
21     result = 0
22     image_name = i
23
24     # 将结果写入CSV文件
25     if result is not None:
26         csv_writer.writerow([image_name, int(result * 100)])
27         print(image_name, result*100)
28     else:
29         print(f"无法处理 {image_name}")

```

calculate.py

```

1  #!/usr/bin/env python
2  # -*- coding: utf-8 -*-

```

```

3  # @Time : 2023/12/13 下午7:10
4  # @Author : SiHang Xu
5  import csv
6  import os
7
8  from PIL import Image
9
10 def calculate_non_black_pixel_ratio(image_path):
11     try:
12         # 打开图片
13         image = Image.open(image_path)
14
15         # 获取图片大小
16         width, height = image.size
17
18         # 初始化非黑色像素计数
19         non_black_count = 0
20
21         # 遍历图片像素
22         for y in range(height):
23             for x in range(width):
24                 # 获取像素颜色
25                 pixel = image.getpixel((x, y))
26
27                 # 检查是否为纯黑色
28                 if pixel != (0, 0, 0):
29                     non_black_count += 1
30
31         # 计算非黑色像素占比
32         total_pixels = width * height
33         non_black_ratio = non_black_count / total_pixels
34
35         return non_black_ratio
36
37     except Exception as e:
38         print(f"Error: {e}")
39         return None
40
41
42
43 # 指定结果文件的路径
44 csv_file_path = "test_result2.csv"
45
46 # 打开CSV文件，使用 'w' 模式
47 with open(csv_file_path, 'w', newline='') as csvfile:
48     # 创建CSV写入对象
49     csv_writer = csv.writer(csvfile)

```

```
50
51 # 写入表头
52 csv_writer.writerow(['fnames', 'pothole %'])
53
54 # 遍历目录中的文件
55 for i in os.listdir('./result_mask'):
56     image_path = './result_mask/' + i
57     result = calculate_non_black_pixel_ratio(image_path)
58     image_name = i.split('_')[0] + '.jpg'
59
60 # 将结果写入CSV文件
61 if result is not None:
62     csv_writer.writerow([image_name, int(result * 100)])
63     print(image_name, result*100)
64 else:
65     print(f"无法处理 {image_name}")
66
67 print(f"数据已保存到 {csv_file_path}")
```