

文章编号 1004-924X(2025)14-2262-16

全局感知与多尺度特征融合的城市道路语义分割

邬开俊¹, 张治瑞^{1*}, 汪 滢¹, 安立伟²

(1. 兰州交通大学 电子与信息工程学院, 甘肃 兰州 730070;

2. 内蒙古民族大学 草业学院, 内蒙古 通辽 028000)

摘要: 语义分割在自动驾驶与智能交通工程应用中发挥着不可替代的作用。针对语义分割现存分割边界模糊、物体间相互遮挡及物体多尺度差异造成的分割精度不足问题, 提出全局感知与多尺度特征融合的城市道路语义分割网络。为改善分割边界模糊的问题, 设计全局感知模块, 通过联合空间和通道信息增强特征之间的交互以感知全局信息; 物体间相互遮挡情况下模型往往需要提升被遮挡区域的敏感度, 为此提出多尺度特征融合模块以兼顾大小物体的分割精度; 采用综合性的多约束特征平滑损失评估模型, 进一步平滑特征, 优化目标以求最优解。经实验验证, 本文方法于 Cityscapes 数据集上在不同分辨率情况下 mIoU 值分别提升 0.5%, 0.9%, 1.7%, 在 ADE20K 数据集上 mIoU 值提升 2.1%。相比于现有语义分割网络模型, 本文方法分割效果有进一步提升。

关键词: 深度学习; 图像处理; 语义分割; 特征融合; 损失函数

中图分类号: TP391.4 **文献标识码:** A

doi: 10.37188/OPE.20253314.2262 **CSTR:** 32169.14.OPE.20253314.2262

Urban road semantic segmentation with global awareness and multi-scale feature fusion

WU Kaijun¹, ZHANG Zhirui^{1*}, WANG Ying¹, AN Liwei²

(1. School of Electronic and Information Engineering, Lanzhou Jiaotong University,
Lanzhou 730070, China;

2. Pratacultural College, Inner Mongolia Minzu University, Tongliao 028000, China)

* Corresponding author, E-mail: zxr_ljd@163.com

Abstract: Semantic segmentation plays an irreplaceable role in autonomous driving and intelligent transportation systems. However, current segmentation networks often suffer from challenges such as blurred object boundaries, mutual occlusions between objects, and significant variations in object scales, which hinder segmentation accuracy. To address these issues, this paper proposed a city road scene semantic segmentation network that integrates global context awareness and multi-scale feature fusion. To mitigate the problem of blurred segmentation boundaries, a Global Awareness Module (GAM) was designed to enhance interaction between spatial and channel-wise information, enabling the network to capture comprehensive global context. For handling object occlusion and improving recognition of partially obscured re-

收稿日期: 2025-03-17; 修订日期: 2025-05-14.

基金项目: 甘肃省自然科学基金项目(No. 23JRR A913); 内蒙古自治区重点研发与成果转化计划项目(No. 2023YFDZ0043, No. 2023YFDZ0054, No. 2023YFSH0043); 兰州交通大学重点研发项目资助(No. ZDYF2304)

gions, a Multi-Scale Feature Fusion Module (MSFF) was introduced, which effectively integrated contextual cues from different receptive fields to ensure segmentation accuracy for objects of varying sizes. Furthermore, a comprehensive Multi-constraint Feature Smoothing Loss was employed to enforce spatial coherence and semantic consistency, guiding the model toward a more optimal solution by refining feature distributions around object boundaries and within complex scenes. Extensive experiments are conducted on two benchmark datasets. On the Cityscapes dataset, the proposed method achieves mIoU improvements by 0.5%, 0.9%, and 1.7% under different input resolutions. On the ADE20K dataset, an mIoU gain of 2.1% is observed. Compared with existing semantic segmentation models, the proposed approach demonstrates superior performance in urban road scene understanding, particularly in terms of boundary delineation, and occlusion robustness.

Key words: deep learning; image processing; semantic segmentation; feature fusion; loss function

1 引言

语义分割涉及对图像的像素级别分类以实现多场景的精准理解,广泛应用于自动驾驶、医学图像分析、遥感影像分析等工程任务中^[1]。从开山之作全卷积神经网络 FCN^[2]后, Ronneberger 等提出 U-Net^[3]引入编解码结构和跳跃连接以改善上采样的空间信息保留, Google 团队提出 DeepLab^[4-7]系列模型, PSPNet^[8]聚合全局上下文信息, Mask R-CNN^[9]在 Faster R-CNN^[10]的基础上加入并行分支以预测分割掩码,推进语义分割发展。然而,这些方法仍存在局限性,如细节信息丢失、复杂场景处理能力不足以及感受野限制带来的全局信息缺失等问题^[11],尤其是在自动驾驶城市道路场景中,诸如交通标志、行人轮廓、车辆边缘等小目标及边界区域,以及在多目标重叠情况下,常因高层语义特征的分辨率下降或感受野不足而出现模糊、错分的问题。针对这些问题,本文进行以下分析。

其一,空间和通道注意力通过关注空间和通道信息来强调重要特征,如 Attention U-Net^[12]。Squeeze-and-Excitation^[13] (SE) 只注重通道信息,而 Convolutional Block Attention Module^[14] (CBAM) 通过通道和空间的双重关注,有效提高关键特征表达能力。如任凤雷^[15]等在语义分支引入通道和空间注意力机制以增强特征表达能力,因此本文从空间和通道角度出发进行全局感知,有效改善分割边界模糊的问题。其二,多尺度特征融合策略确保网络既能捕捉到图像细节信息又能理解整体语义。例如,侯志强^[16]等在解

码器阶段设计多尺度融合模块,促进深层与浅层特征的融合; 闫河^[17]等通过多尺度融合,提升语义理解能力和分割准确性。此外,大核卷积能够覆盖更大的感受野,更好地捕捉到图像中的全局特征,如 ConvNext^[18], RepLKNet^[19], SlaK^[20]等,提供改善模型性能新思路。受上述启发,本文利用自注意力与大卷积核融合的多尺度特征融合策略以提升分割精度。其三,语义分割中的主要损失交叉熵损失^[21]易受类别不平衡影响,降低空间上的一致性和平滑性; Dice Loss^[22]适用于处理类别不平衡的任务,但对背景噪声较为敏感; Focal Loss^[23]的提出同样是为了解决类别不平衡的问题,但对边界区域的关注不够。由于不同任务对损失函数的需求具有高度特异性,为更适应城市道路场景语义分割任务,本文设计更具针对性的损失函数。

据上述分析,本文提出全局感知与多尺度特征融合的城市道路语义分割网络。首先,利用全局感知模块 (Global Awareness Module, GAM), 结合空间聚合 (Space Aggregation, SA) 与通道聚合 (Channel Aggregation, CA) 感知全局信息。空间聚合采集邻域空间像素信息,通道聚合捕捉不同特征通道之间的相关性,细化分割边界。其次,设计多尺度特征融合模块 (Multi-Scale Feature Fusion Module, MSFF), 浅层特征提取细节及边缘信息,深层特征提取物体轮廓及场景结构信息,以减少错分及漏分导致分割精度低的问题并兼顾大小物体的分割精度。最后,为降低背景干扰并增强特征平滑性,引入多约束特征平滑损失 (Multi-Constraint Feature Smoothing Loss,

MCFS Loss),结合交叉熵损失、特征引导损失以及基于拉普拉斯算子的空间一致性损失,以减少噪声干扰,优化模型性能。本文在 Cityscapes^[24]数据集上验证自动驾驶与智能交通工程应用下方法的有效性,在 ADE20K^[25]数据集上验证方法的泛化性。

2 网络模型

本文的整体网络模型如图 1 所示。本文模型为 4 阶段网络,阶段 1 为 1×1 卷积与 3×3 卷积的堆叠, 1×1 卷积用于降维,而 3×3 卷积则进

一步提取空间特征,实现高效的特征提取和表示。阶段 2 是 GAM,通过聚合空间和通道特征感知全局信息,细化分割边界。阶段 3 是 3×3 的卷积与 MSFF 模块的结合,改善网络在复杂背景及遮挡情况下漏检、错检导致的分割精度低的问题。阶段 4 是 MSFF 的叠加,进一步提升分割效果。同时融合阶段 2 和阶段 4 的特征,获取更丰富的语义信息。网络解码器包含解码头深度聚合金字塔池模块^[26](Deep Aggregation Pyramid Pooling Module, DAPPM),分割头由 3×3 的卷积、ReLU 激活函数、BN 以及 1×1 卷积的分类器组成。

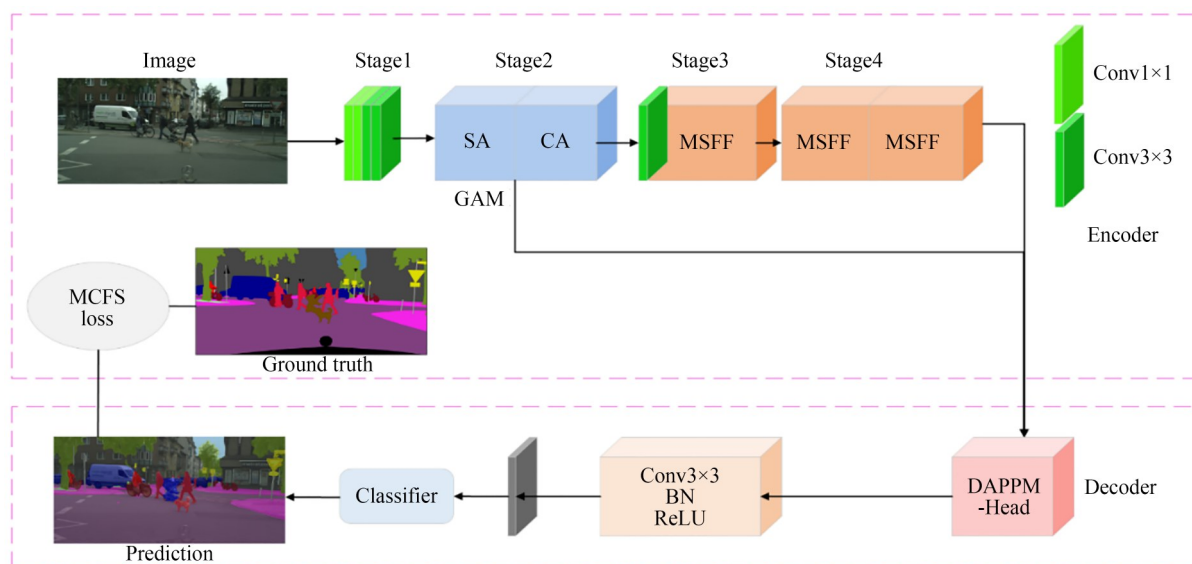


图 1 网络模型整体框架

Fig. 1 Overall framework of the network

2.1 全局感知模块

在语义分割任务中,网络模型需要同时捕捉图像中各像素的局部特征和全局上下文信息。传统的卷积操作往往仅关注空间维度,忽略通道之间的关系。对于一些复杂场景,如自动驾驶或室内场景分割,图像中的物体类别多样且密集,背景信息丰富。单纯地从空间或通道维度提取特征,会导致细节信息的遗漏。全局感知模块 GAM 通过联合空间和通道信息,在保持空间信息的同时,有效利用通道信息,增强全局感知能力,帮助模型更好地区分背景和前景,改善分割边界模糊的问题。GAM 分为两个模块,SA 和 CA,如图 2 所示。

设阶段 1 网络输出为 X ,模型在阶段 2 网络先通过空间聚合模块 SA,如图 2 所示。SA 模块保持通道不变,通过空间加权体现不同空间位置特征的重要性。SA 先经过 Layer Norm 层和 1×1 的卷积层进行前置处理,以稳定数据分布并保持通道数,其输出为 Y_1 。此后,通过差分增强项强调图像的高频信息,即边缘信息。全局平均池化操作会将特征图在空间维度($H \times W$)上进行平均,产生一个每个通道上只有一个值的全局向量,表示的是该通道的全局响应强度,其本质是对整个特征图的平滑处理,即只保留了各通道的平均能量,这正是典型的低频成分。当将原始特征图 Y_1 与其全局平均结果进行特征差分增强

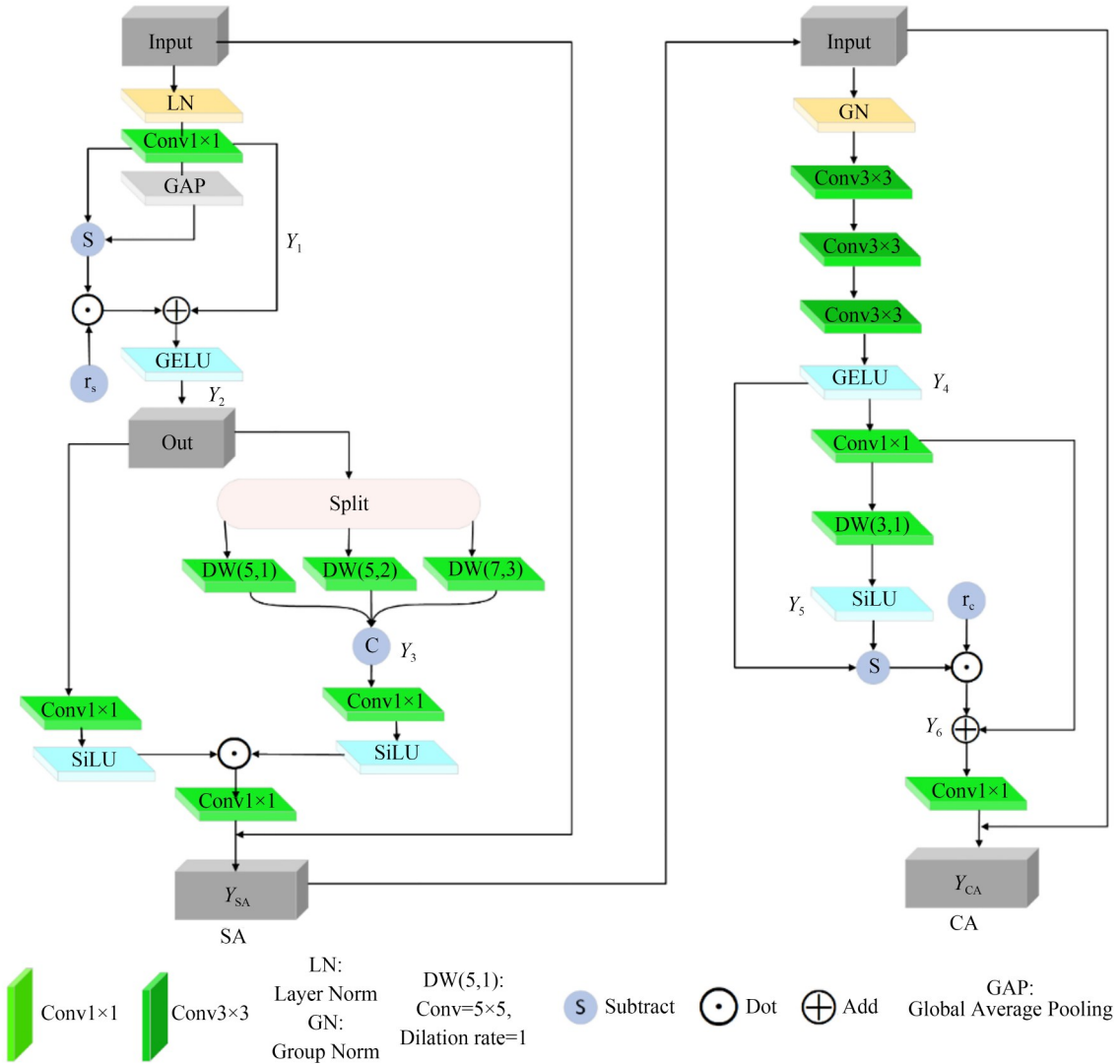


图 2 全局感知模块

Fig. 2 Global awareness module

(即: $Y_1 - \text{GAP}(Y_1)$) 时, 实际是在从原始特征中去除掉空间上恒定的低频基准值而保留高频信息, 即边缘、纹理等特征值与 GAP 差异较大的区域将被显著保留甚至放大。这种差异性在分割任务中恰好对应于物体边界等关键区域。给予差分增强项一个空间比例因子 r_s (初始化为 0), 通过可调节的标量对上述差分增强项进行调节, 从而赋予网络在训练中调整边界敏感程度的能力。而后 $r_s \odot (Y_1 - \text{GAP}(Y_1))$ 的输出与 Y_1 融合, 通过 GELU 函数激活进一步增强特征表示的非线性性, 如式(1)所示:

$$Y_2 = \text{GELU}(Y_1 + r_s \odot (Y_1 - \text{GAP}(Y_1))). \quad (1)$$

Y_2 之后并行三个深度可分离卷积, 通过不同

卷积核以及膨胀率获取不同阶特征, 第一路卷积采用核尺寸为 5×5 , 膨胀率为 1, 有效感受野为 5×5 , 用于提取精细的低阶特征; 第二路卷积同样为 5×5 卷积核, 但膨胀率提升至 2, 有效感受野增至 9×9 , 用于提取中尺度的结构信息; 第三路卷积采用 7×7 卷积核并设置膨胀率为 3, 有效感受野达到 19×19 , 更适合捕捉全局的高阶语义特征。设输入通道为 C , 低阶输入通道为 C_L , 中阶输入通道为 C_m , 高阶输入通道为 C_h , 则他们之间的关系如式(2)所示。将此三个维度的特征拼接起来, 增强特征多样性, 形成多级上下文信息, 如式(3)所示:

$$C = C_L + C_m + C_h, \quad (2)$$

$$Y_3 = \text{Concat}(Y_L, Y_m, Y_h), \quad (3)$$

其中: Y_L, Y_m, Y_h 分别表示三个维度的输出, Concat 表示拼接操作。

SA 模块的最后本文采用 SiLU 激活函数, 自适应地缩放输入, 并且保留负值信息, 避免在负值区域丢失重要信息, 适合于语义分割这种需要保留细微信息的任务。SA 模块的最终输出如式(4)所示:

$$Y_{SA} = \text{Conv}_{1 \times 1}(\text{SiLU}(\text{Conv}_{1 \times 1}(Y_2))) \odot \text{SiLU}(\text{Conv}_{1 \times 1}(Y_3)). \quad (4)$$

Y_{SA} 作为 CA 模块的输入, Group Norm 层将通道分成若干个组, 在通道维度进行标准化以捕捉到通道特征的分布, 以缓解在大批量及数据不平衡场景下的归一化难题。后通过 3×3 卷积的堆叠以及 GELU 激活函数得到 Y_4 。GELU 激活函数在负值区域仍然允许一定的信息传递, 在处理复杂特征时能够保留更多的上下文信息, 适合复杂场景的理解。为进一步挖掘通道之间的协同关系, 本文设计通道重分配项, 通过计算 Y_4 和通道投影 $Y_5 W$ 之间的差异即 $(Y_4 - Y_5 W)$, 强调与通道投影结果差异较大的部分(边缘、物体内部纹理), 并经 SiLU 激活, 反映输入特征经过非线性变换后所发生的变化, 捕捉重要特征。同样, 给予通道重分配项一个可调比例因子 r_c (初始化为 0) 以重新分配通道信息, 调整通道重分配的强度, 特别是边缘特征。如式(5)、式(6)所示:

$$Y_5 = \text{SiLU}(DW_{(3,1)}(\text{Conv}_{1 \times 1}(Y_4))), \quad (5)$$

$$Y_6 = Y_4 + r_c \odot (Y_4 - Y_5 W). \quad (6)$$

其中: $DW_{(3,1)}$ 表示卷积核为 3×3 , 膨胀系数为 1 的深度可分离卷积。 W 是一个用于降低通道数的权重矩阵, 通过 $Y_5 W$ 得到一个一维特征图, 表示每个位置的通道信息汇总。将空间信息与通道信息融合, 获取丰富的全局特征, 并通过 1×1 的卷积提高模型的非线性表达能力, 如式(7)所示:

$$Y_{CA} = Y_{SA} + \text{Conv}_{1 \times 1}(Y_6). \quad (7)$$

2.2 多尺度特征融合模块

本文多尺度特征融合模块 MSFF 通过捕捉多尺度特征提升模型对不同大小物体的分割能力, 同时减少因错检、漏检导致的分割精度低的问题。此外, MSFF 模块通过引入跨层信息传递机制, 确保模型在整合不同尺度特征时保持特征表达的一致性和完整性。

MSFF 包括上下层分支, 每层又可细分为两个分支, 共四个分支, 如图 3 所示。在上层分支, 模块的输入经过深度可分离卷积以及层归一化之后, 将此部分输出 Y 给予一个特征分割比例因子 r 。设通道数为 C , 则有 $C_A = rC$ 。 C_A 代表自注意力通道数, 与剩余块部分互不影响, 则自注意力分支 a 与剩余块分支 b 的表示如式(8)所示:

$$Y_{A,R} = \text{Split}(Y, [C_A, C - C_A]), \quad (8)$$

其中: Split 表示分割操作, Y_A 为自注意力部分, Y_R 为剩余块部分。

自注意力的表示如式(9), 则自注意力输出可表示为式(10), 上层输出可表示为式(11), 其中 σ 表示 SoftMax, Att 表示 Attention, Q 是查询, K 是键, V 是值, W^Q, W^K, W^V 为投影权值, d_{qk} 是查询和键的维度, Concat 是拼接操作。 $S_A(a$ 分支)通过矩阵乘法操作, QK^T 计算每对像素之间的相似性分数, Softmax 保证对所有类别像素赋予不同权重, 整合语义相似或高度相关位置信息, 在全局范围内捕捉长距离依赖关系。通过此方式, 模型在多尺度目标存在遮挡的情况下, 仍能保持良好的分割一致性。目标被遮挡时, 由于邻域像素特征丢失, 模型难以准确判断目标边界导致错检、漏检。针对遮挡导致的邻域特征缺失问题, 自注意力机制可利用图像上下文语义信息, 增强对被遮挡区域的判别能力, 有效弥补信息丢失, 降低漏检与错检风险。后接前馈神经网络 FFN^[27], 对各位置特征进行非线性变换与重组, 强调关键特征并抑制冗余信息。 $Y_R(b$ 分支)直接保留原始的空间特征表达, 未经过全局建模干扰。 Y_R 作为原始通道的一部分, 天然保留了与局部感受野相关的卷积特征, 尤其对于小目标, 能够保持像素级别的信息对齐, 确保细节不被丢失。相比之下, Y_A 的自注意力机制虽然能够增强全局语义建模能力, 但会对局部精度产生一定稀释效应。 Y_R 不经过全局注意力扰动, 因此作为 b 分支输出, 起到保留空间局部信息的作用, 是对自注意力分支的结构性补偿。

$$\text{Att}(Q, K, V) = \sigma \left(\frac{QK^T}{\sqrt{d_{qk}}} \right) V, \quad (9)$$

$$S_A = \text{Att}(Y_A W^Q, Y_A W^K, Y_A W^V), \quad (10)$$

$$Y_U = \text{Concat}(S_A, Y_R). \quad (11)$$

下层中 c 分支为大卷积核 (23×23) 的输出, d 分支则是将该大卷积核进行分解重组后的输出。 23×23 的卷积核计算开销较大,而通过对其进行分解,可以降低模型的计算复杂度,其分解公式如式(12)所示:

$$F_r = F_{r-1} + d_i(K_i - 1), \quad (12)$$

其中: F_r 代表大卷积核的感受野, F_{r-1} 代表分解的第一个卷积的感受野, d_i 表示膨胀系数, K_i 表示分解的第二个卷积的卷积核大小。 23×23 卷积核 ($K=23, d=1$) 可以被分解为两个较小的卷积核 ($K_1=5, d_1=1; K_2=7, d_2=3$), 在保持上下文建模

能力的同时,减少计算量。

大卷积核有助于在更广的感受野内提取丰富的大目标整体轮廓信息,增强模型对复杂空间结构的建模能力。但与此同时,其对小目标及细粒度局部特征的敏感性较弱。因此,本文融合上层具备全局建模能力的自注意力分支 a 、保持细节信息的剩余块分支 b 以及下层富含整体轮廓信息的大核卷积分支 c 、拥有局部特征的分解重组分支 d , 使网络兼具全局信息建模与局部特征提取能力,实现对多尺度目标更全面、精准的感知与理解。

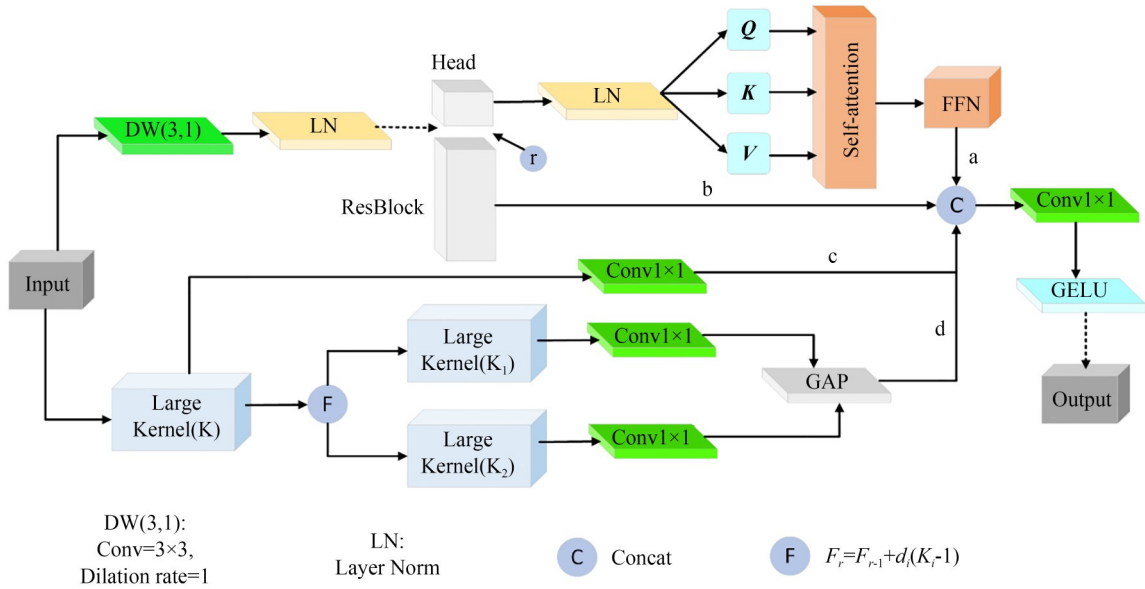


图3 多尺度特征融合模块

Fig. 3 Multi-Scale feature fusion module

2.3 多约束特征平滑损失

现有主流损失函数在处理城市道路语义分割任务中仍存在局限。CE Loss对类不平衡问题不敏感,易忽略小类别;Dice Loss易受背景干扰,背景(建筑、天空等)常占据大量像素,当前景目标(行人、交通标志)区域很小,会稀释前景目标的梯度信号,导致小目标学习困难;Focal Loss虽能处理类别不平衡,但其忽略空间结构信息,在道路边缘、车道线等连续结构中易出现预测断裂或模糊边界,影响细节表现。为适应城市道路场景,本文提出一种融合交叉熵损失(Cross-Entropy Loss, CE Loss)、特征引导损失(Feature-Guided Loss, FG Loss)以及基于拉普拉斯算子

的空间一致性损失(Spatial Consistency Loss Based on Laplacian Operator, SC Loss)的多约束特征平滑损失(MCFS Loss),以减少背景对前景预测的干扰,增强模型对目标结构细节与边界的一致性建模能力,尤其适用于小目标、道路边缘等结构,从而提升分割精度和鲁棒性。MCFS Loss如式(13)所示:

$$L = L_{CE} + \alpha L_{FG} + \beta L_{SC}, \quad (13)$$

其中, α 和 β 是可调参数。

交叉熵损失是语义分割中常用的基础损失函数。其作用是度量模型的预测结果与真实标签之间的差异,鼓励网络输出更准确的分割结果。在多约束特征平滑损失中,其充当基础的像

素级预测精度优化作用,如式(14)所示:

$$L_{\text{CE}} = - \sum_i y_i \log(y_i') + (1 - y_i) \log(1 - y_i'), \quad (14)$$

其中: y_i 表示第*i*个像素的真实标签的特征映射, y_i' 表示模型对第*i*个像素的预测结果的特征映射。

特征引导损失 FG Loss 通过对比预测分割图像的高阶特征(如车的形状、纹理、边缘或者几何信息等)与理想目标特征之间的差异来优化模型,使用欧几里得距离,即 L_2 范数来度量特征之间的差异。FG Loss 会惩罚与前景特征距离大的特征区域,自然抑制模型将背景误识为前景的行为。损失可由式(15)表示:

$$L_{\text{FG}} = \sum_i (\|F(y_i') - F(y_i)\|_2^2), \quad (15)$$

其中, $\|\cdot\|_2$ 表示 L_2 范数,度量分割结果平滑度的程度。

空间一致性约束通过惩罚分割图像中不自然的边缘或不连续的区域,使得预测结果更加光滑且自然。在平面图像中,拉普拉斯算子 ∇^2 被定义为图像的二阶导数,用来衡量图像(尤其是分割结果边缘)的局部区域变化,减少小尺度的波动和噪声,有效检测图像中突变的区域。如式(16)所示:

$$\nabla^2 I(x, y) = \frac{\partial^2 I(x, y)}{\partial x^2} + \frac{\partial^2 I(x, y)}{\partial y^2}, \quad (16)$$

其中: $I(x, y)$ 是图像的像素值, x 和 y 是图像的坐标。

在城市道路语义分割任务中,常常希望模型在边界附近保持空间一致性,即避免分割边界出现明显的不平滑或不自然的变化。本文通过对分割结果应用拉普拉斯算子来定义一个平滑损失项 SC Loss,如式(17)所示:

$$L_{\text{SC}} = \sum_i \|\nabla^2 y_i'\|_2^2, \quad (17)$$

其中: $\nabla^2 y_i'$ 表示对预测结果应用拉普拉斯算子的输出,反映该位置的局部变化程度。

3 实 验

本文实验环境为:Windows11 系统,python 版本 3.8,使用显卡为 NVIDIA GeForce RTX 3090 GPU,整个环境基于 pytorch 框架。优化器

选择 Adam 优化器;学习率为衰减式学习,初值为 4×10^{-4} ;权重衰减率为 0.012 5;共迭代 1.6×10^5 次。图像增强策略选用随机缩放,缩放比率为 (0.25, 1.5);随机裁剪,指定裁剪区域中任意单个类别的像素比例不超过 0.75;随机上下左右翻转,翻转比率为 0.5;采用光度失真对图像的颜色、亮度、对比度、饱和度等属性进行随机调整,增加图像多样性。

3.1 数据集

Cityscapes 数据集是一个广泛应用于语义分割任务的数据集,包含 5 000 张高分辨率德国多城市街道场景图像,专为城市环境下自动驾驶研究而设计,本文将其分为道路、行人、汽车等 19 类,图片输入使用三种不同分辨率 $1\,024 \times 512$ (S_1), $1\,536 \times 768$ (S_2) 以及 $2\,048 \times 1\,024$ (S_3),以获取多尺度预测结果。采用 ADE20K 数据集验证模型在不同场景下的泛化性,其包含超过 20 000 张图像,涵盖多种场景类型,包括室内外环境,提供 150 个物体类别,每张图像都经过像素级标注,适合细粒度的场景解析,图片输入大小为 512×512 。

3.2 评价指标

本文采用语义分割最重要的指标 mIoU 以及类别平均像素准确率 mPA 来评估模型。mIoU 用于衡量模型在每个类别上的预测效果,评估模型预测与真实标注之间重叠程度,能够更准确地反映模型在不同类别上的表现,尤其是在类别不平衡的情况下。其计算公式如式(18):

$$mIoU = \frac{1}{C} \sum_{i=1}^C \frac{TP_i}{TP_i + FP_i + FN_i}, \quad (18)$$

其中: TP_i 代表正确预测为类别*i*的像素数, FN_i 代表实际为类别*i*但被错误预测为其他类别的像素数, FP_i 代表错误预测为类别*i*的像素数, C 为类别数。

mPA 通过计算每个类别的像素准确率,并对所有类别的准确率进行平均,来反映模型在各个类别上的表现。其计算公式如式(19)所示:

$$mPA = \frac{1}{C} \sum_{i=1}^C \frac{P_{ii}}{\sum_{j=1}^C P_{ij}}, \quad (19)$$

其中: P_{ij} 是第*i*个类别预测为第*j*个类别的总量。

FPS 反映一定时间内模型处理的总图像帧数,是衡量模型推理速度的关键指标,用以评估

模型的执行效率,但与硬件性能、模型复杂度、图像分辨率等多个因素相关。

3.3 对比实验

本文方法与语义分割经典模型 PSPNet^[8], DeeplabV3Plus^[7], SegFormer^[28], DDRNet-23^[29] 以及近几年模型 SegNext-T^[30], SeaFormer-B^[31], RTFormer-B^[32], SCTNet-B^[33] 等 8 种模型对

Params(模型参数量大小), mIoU, mPA, FPS 等指标在 cityscapes 数据集上进行不同输入分辨率图像 $1\,024\times 512(S_1)$, $1\,536\times 768(S_2)$, $2\,048\times 1\,024(S_3)$ 的对比,如表 1 所示。相较于 SCTNet-B,本文模型在 Cityscapes 数据集上 mIoU 指标有所提升,在不同分辨率情况下 mIoU 值分别达到 76.9%,79.8%,81.4%。

表 1 在 Cityscapes 数据集上的对比实验
Tab. 1 Contrast experiment on Cityscapes dataset

Models	Params/ m	S ₁			S ₂			S ₃		
		mIoU/%	mPA/%	FPS	mIoU/%	mPA/%	FPS	mIoU/%	mPA/%	FPS
PSPNet(ResNet101)	67.0	72.6	76.2	43.0	74.4	77.1	34.2	75.8	79.6	20.4
DeeplabV3Plus(Xception)	41.0	73.5	77.4	47.6	75.1	78.6	37.5	76.5	80.4	24.1
SegFormer-B0	3.8	74.9	79.0	82.1	77.1	80.2	56.7	78.4	81.8	39.3
DDRNet-23	20.1	75.8	80.5	81.6	78.1	81.3	66.8	79.2	82.7	52.1
SegNext-T	4.3	76.2	81.1	66.7	78.4	81.9	43.1	79.5	83.0	26.5
SeaFormer-B	8.6	74.2	78.8	198.9	76.0	79.7	141	77.3	81.2	100.3
RTFormer-B	16.8	75.6	80.2	78.6	78.3	81.0	58.3	79.0	82.4	46.7
SCTNet-B	17.4	76.4	81.6	131.6	78.9	82.3	95.3	79.7	83.4	58.3
Ours	24.3	76.9	82.5	117.4	79.8	83.4	82.4	81.4	85.8	53.7

表 2 是当前主流模型 DeeplabV3Plus, SegFormer 等在 Cityscapes 数据集上的 19 个类的交并对比表。显而易见,杆、交通信号灯、摩托车、自行车等小目标类别上的 IoU 远低于道路、建筑物、植被等大目标类别,反映出当前模型在复杂街道场景中细节建模能力仍存在一定不足。本文模型对小目标如杆、摩托车、自行车等的分割交并比

分别比基线模型提升 2.1%,1.9%,2.4%,并且通过扩大感受野,进一步提高大目标如道路、建筑物、植被等的交并比,提升幅度分别为 0.9%,1.6%,1.1%。上述数据说明本文方法在兼顾细节信息和全局上下文建模时的有效性。

为直观地表明本文方法在自动驾驶城市道路场景下的有效性,本文采用 6 个模型在

表 2 在 Cityscapes 数据集上 19 个类的对比实验
Tab. 2 Contrast experiment of 19 classes on Cityscapes dataset (%)

类别/模型/mIoU	DeeplabV3Plus	SegFormer	DDRNet	SegNext	SeaFormer	SCTNet	Ours
道路	96.8	97.2	97.3	97	95.8	97.6	98.5
人行道	87.2	88.1	87.9	89.2	88.5	88.2	88.9
建筑物	89.7	91.2	91.5	92.1	88.9	92.2	93.8
墙壁	56.3	59.6	62.2	61.7	56.7	62.9	67.7
栅栏	65.1	66.7	66.3	66.6	65.4	67.1	68.4
杆	51.8	53.4	55.7	56.2	49.8	55.4	57.5

续表 2 在 Cityscapes 数据集上 19 个类的对比实验

Tab. 2 Contrast experiment of 19 classes on Cityscapes dataset (%)

类别/模型/mIoU	DeeplabV3Plus	SegFormer	DDRNet	SegNext	SeaFormer	SCTNet	Ours
交通信号灯	65.6	67.8	70.5	69.2	68.3	69.7	70.8
交通标志	71.3	73.1	74.2	75.3	72.4	74.9	75.1
植被	90.6	91.3	92.1	92.7	91.1	92.6	93.7
地形	65.7	67.1	68.4	67.9	64.6	68.9	71.5
天空	95.3	96.2	96.6	95.8	94.9	96.2	96.7
行人	80.4	82.3	83.1	83.5	81.7	82.8	84.9
骑行者	58.4	61.1	62.3	62.1	59.2	63.7	65.3
小轿车	93.6	94.9	95.1	95.3	94.5	96.2	97.5
大卡车	86.3	88	86.5	87.2	88.6	87.5	90.7
公共汽车	91.5	92.7	93	92.5	92.1	94.4	95.4
火车	86.6	87.9	88.2	88.9	87.3	88.6	90.7
摩托车	59.1	65.7	66.4	67.4	64.4	67.5	69.4
自行车	62.6	65.2	68.3	70.3	65.2	67.8	70.2
平均值	76.5	78.4	79.2	79.5	77.3	79.7	81.4

Cityscapes 数据集上进行可视化效果图对比分析,如图 4 所示。A 中,DeepLabV3Plus, SegFormer-B0 等模型对左边墙壁下的道路分割不完整,相比于 SCTNet-B,本文模型对路面边界的分割边界更加平滑。B 中,大多数模型对右边的草丛分割不完整,分隔边界模糊,且没有分割出被树木部分遮挡的交通信号灯,本文模型细化草丛

的分割边界,并正确分割出被部分遮挡的交通信号灯。C 中,左边在大棚底下的路人被周围物体遮挡,多数模型并没有分割出原有结果;以及远处的道路标志,DeepLabV3Plus, SegFormer-B0, SCTNet-B 等模型均漏检一个,本文模型呈现出完整结果,有效减少错检及漏检的情况,说明本文方法的有效性。

表 3 在 ADE20K 数据集上的对比实验

Tab. 3 Contrast experiment on ADE20K dataset

Models	Params/m	mIoU/%	MPA/%	FPS
PSPNet(ResNet101)	67.0	33.1	51.3	50.7
DeeplabV3Plus(Xception)	41.0	34.5	53.9	58.4
SegFormer-B0	3.8	36.9	56.4	79.6
DDRNet-23	20.1	38.7	58.5	90.3
SegNext-T	4.3	40.6	61.4	60.0
SeaFormer-B	8.6	37.8	57.6	43.6
RTFormer-B	16.8	40.5	61.2	91.7
SCTNet-B	17.4	41.5	63.0	139.7
Ours	24.3	43.6	65.4	125.9

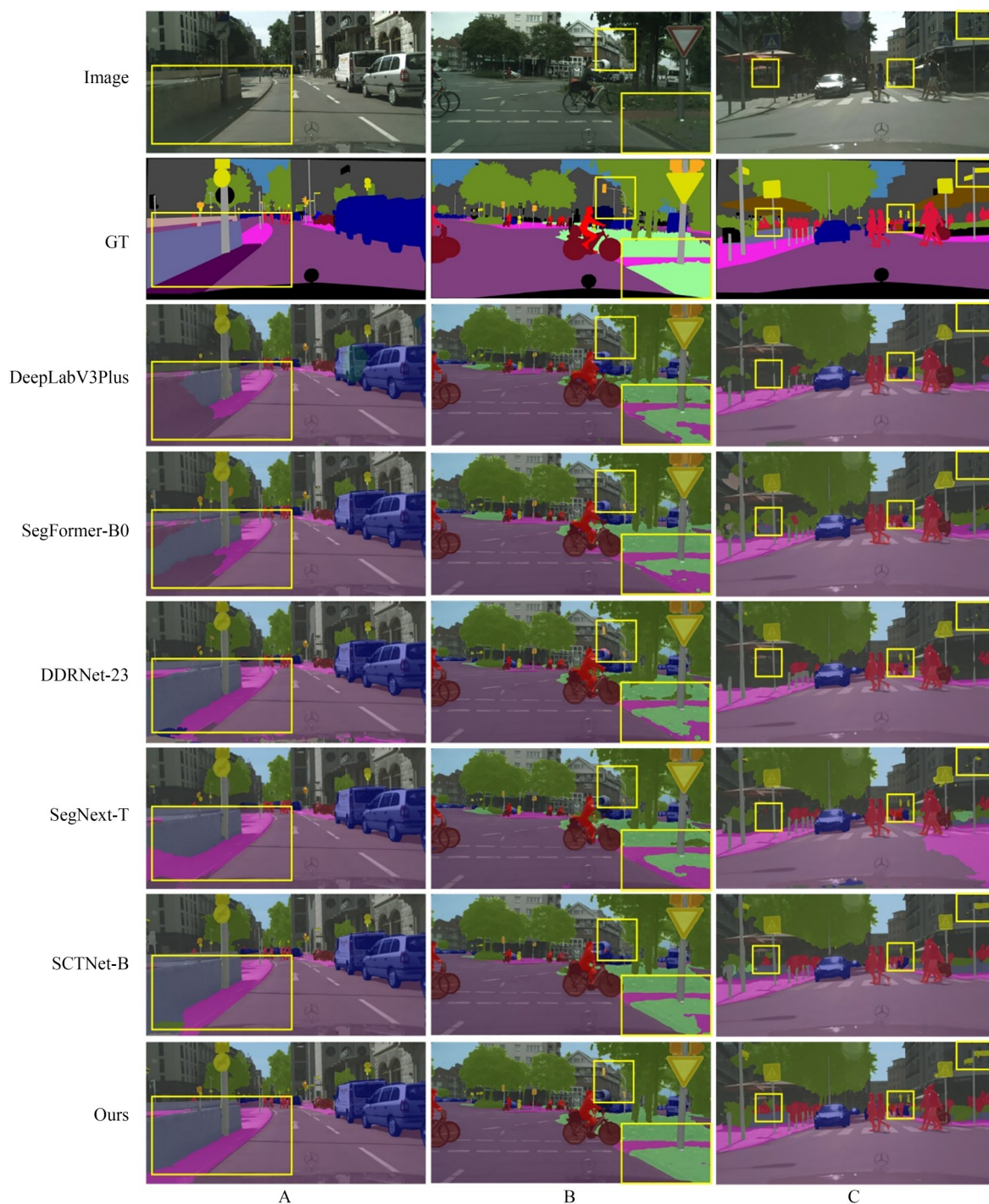


图4 Cityscapes数据集可视化效果图对比

Fig. 4 Cityscapes data set visualization comparison

为验证方法在多场景下的泛化性,本文在ADE20K数据集上进行对比,如表3所示。ADE20K数据集的mIoU达到43.6%,从数值上说明本文方法的泛化性。从FPS数值说明,本文方法比大多数分割模型推理速度快,在保持分割

精度的情况下仍然具有较好推理速度,达到了精度与效率之间的相对平衡。

在ADE20K数据集上,本文采用5个模型进行可视化效果图对比,如图5所示。D中,左边大房间和右边小房间共存,PSPNet,SegForme-B0

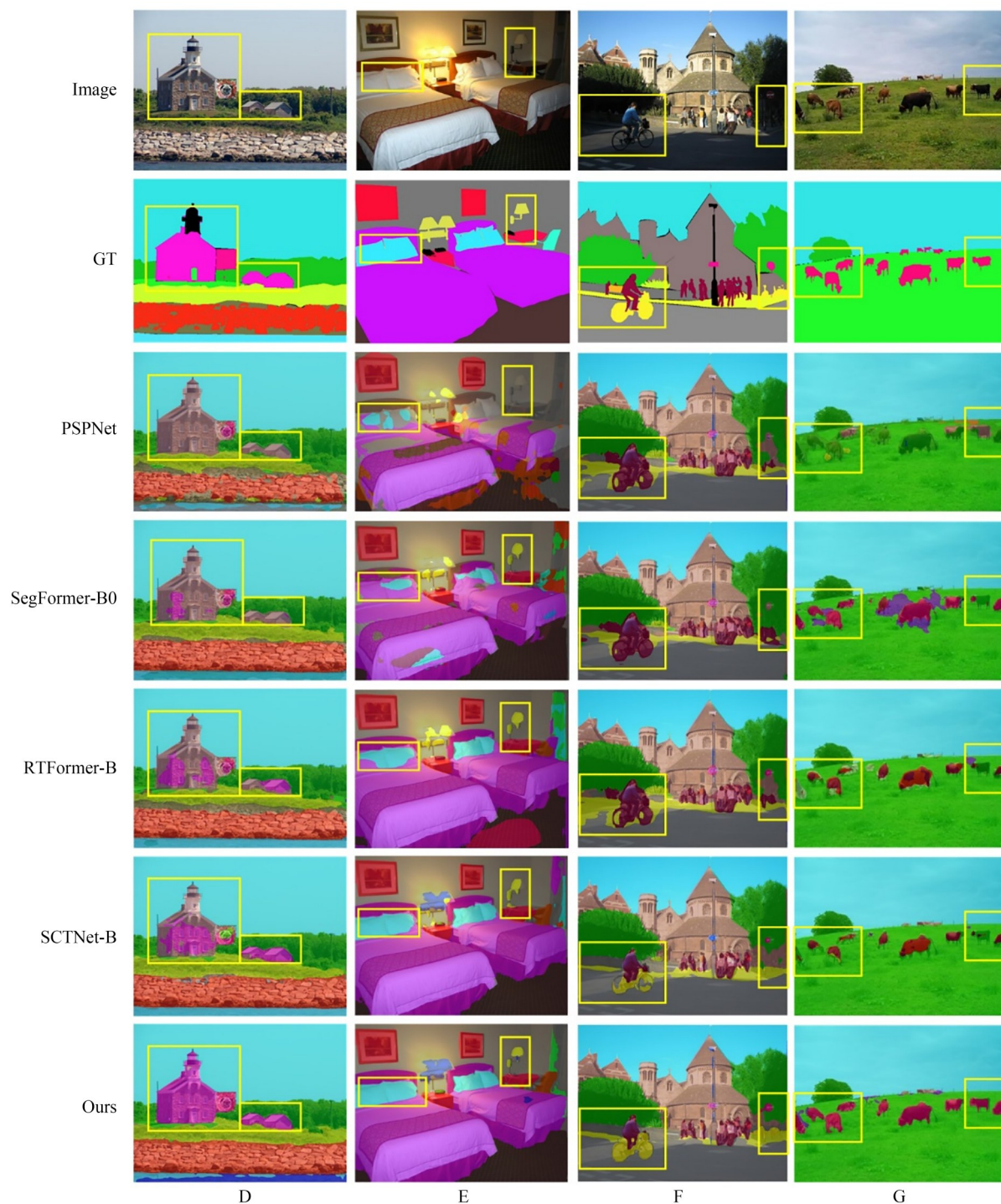


图5 ADE20K数据集可视化效果图对比

Fig. 5 ADE20K data set visualization comparison

等模型无法做到两者之间分割精度的权衡,出现大小物体均分割不完整的问题,本文方法在提升大房间的分割精度的同时,兼顾小屋子的分割精度。对E中枕头、台灯等小物体,本文方法分割结果更完整,分割边界更平滑。F中,PSPNet,

SegFormer, RTformer等将自行车与骑行者混为一类,且对右边的交通标志分割不完整,G中大多模型出现对牛个体的不同程度的漏检,本文方法有效减少错检、漏检的数量,说明本文方法在多场景下的泛化性。

3.4 消融实验

3.4.1 全局感知模块分析

全局感知模块 GAM 从空间和通道角度出发,交互空间与通道信息,增强模型的全局感知能力。本文对空间比例因子 r_s 和通道比例因子 r_c 进行探讨,如表 3 所示。实验组 1 为对 Baseline 的复现结果,通过实验组 2、实验组 3 发现 GAM 模块对 SA 的依赖比 CA 略大,因此本文从区间 $[0,1]$ 将 r_s 从 0 开始往上调,将 r_c 从 1 开始往下调,调整幅度为 0.1,以权衡 r_s 与 r_c 的最优值,最终发现 $r_s=0.6, r_c=0.4$ 时(实验组 5)指标达到最优,如实验 4、实验 5 所示。分析原因,语义分割要求模型对图像的空间布局有较强理解,空间信息帮助模型识别物体的形状、边界等信息,这对于准确

分割更加重要。

实验 6、实验 7 是本文与典型通道及空间模块 SE 和 CBAM 的对比,如表 4 所示。SE 将整个通道的信息压缩为一个全局标量,这在需要精确空间感知的语义分割任务中表现不佳,其主要关注通道间的权重分配,忽略了特征的空间维度。CBAM 结合通道注意力和空间注意力,旨在增强模型对特征的选择性关注能力,因其结构简单、易于集成的特性成为有效的注意力机制,但其在独立建模通道和空间注意力时,捕获通道与空间联合关系的效率略低。本文 GAM 通过空间聚合和通道聚合实现对空间和通道的联合优化,并采用深度可分离卷积减少计算开销,达到更佳分割精度。

表 4 全局感知模块分析
Tab. 4 Global Awareness Module analysis

实验组	Scaling factor		Models				mIoU/%	
	r_s	r_c	Baseline	+SE	+CBAM	+GAM	Cityscapes	ADE20K
1	0.0	0.0	✓				79.69	41.52
2	1	0.0				✓	80.21	42.03
3	0.0	1.0				✓	79.94	41.60
4	0.5	0.5				✓	80.16	42.17
5	0.6	0.4				✓	80.33	42.30
6	0.0	0.0		✓			79.60	41.39
7	0.0	0.0			✓		80.12	41.86

3.4.2 多尺度特征融合模块分析

MSFF 分别从 a,b,c,d 四个分支进行特征融合以获取丰富的多尺度特征,本文对四个分支的影响进行分析,如表 5 所示。a,b 是注意力部分与剩余块结合的上层分支。为寻找最优的比例因子 r , 本文将上层分支特征分割比例系数 r 从 0 到 1 以 0.1 为步长进行遍历,以得到最优值。自注意力机制能够增强对长距离依赖的建模能力,使模型更好地捕捉到目标区域之间的语义关系,但自注意力本身容易过度关注显著区域,会弱化边界细节或干扰较小目标的表征,因此如果 r 过大(如 1 和 0.7),融合结果更偏向于注意力输出,导致目标边缘模糊或小目标遗漏,模型性能有一定波动;剩余块能够有效保留输入特征的结构性

表 5 多尺度融合模块分析
Tab. 5 Multi-scale fusion module analysis

实验组	Scaling factor	Branch				mIoU/%	
	r	a	b	c	d	Cityscapes	ADE20K
8	0.0	✓				80.36	42.32
	1.0	✓				80.44	42.41
	0.7	✓	✓			80.49	42.44
	0.3	✓	✓			80.67	42.76
9	---			✓		80.61	42.61
10	---				✓	80.55	42.57
11	0.3	✓	✓	✓		80.81	43.22
12	0.3	✓	✓	✓	✓	81.06	43.41

信息,具备一定的抗噪能力和泛化能力,当 r 较小(接近0时),模型特征几乎未被注意力机制调制,虽然稳定但缺乏语义辨别力,mIoU值偏低; $r=0.3$ 时恰好在保留原有优势的同时,引入一定比例的语义增强,有效缓解自注意力的噪声放大效应,最大程度发挥多尺度信息互补性,如实验8所示。实验9和实验10表明,大卷积核(c分支)和分解卷积核(d分支)对网络模型的性能具有一定影响;进一步,实验11和12则证明,将两者融合后对整体模型带来更大的性能提升。这种效果主要得益于大卷积核在捕获图像中大范围上下文信息方面的优势,尤其在需要识别大物体或全局上下文关系的场景中表现突出。同时,通过卷积核的分解,网络能够更加精细地提取局部和细粒度特征,增强模型的灵活性。此设计不仅提高对被遮挡部分的预测准确性,还兼顾大小物体的分割精度,实现多尺度物体分割精度的有效平衡。

3.4.3 多约束特征平滑损失分析

本文针对语义分割中的不同损失函数进行分析与对比,并对本文多约束特征平滑损失参数 α 和 β 进行讨论,如表6所示。实验13中,本文分别对参数 α,β 进行网格化超参数搜索, α 控制特征引导损失的贡献度,根据复杂场景下特征分布不均匀特性,设置其搜索范围为 $\alpha \in$

$[0.5, 2.5]$; β 控制平滑约束项的强度,考虑其对细节保留的影响,其值会略小,故设置 $\beta \in [0.0, 2.0]$ 。 α,β 步长均为0.5,通过枚举组合在验证集上评估,逐步逼近最优组合。对于 α ,当 $\alpha \leq 1.0$ 时,FG Loss在优化过程中权重略低,导致模型对语义特征的引导能力减弱,复杂区域(如街道遮挡区域、室内小物体)表现不佳,mIoU值略低;当 $\alpha \geq 2.0$ 时,FG Loss占比过高,会干扰主任务损失的收敛方向,导致局部过拟合及训练不稳定;当 $\alpha = 1.5$ 时,模型表现出较强的特征提取与类别区分能力,兼顾模型稳定性和泛化能力。对于 β ,当 $\beta \leq 0.5$ 时,平滑约束不足,分割图中会出现明显边界断裂与噪声伪影;当 $\beta \geq 1.5$ 时,平滑约束过强,尤其在ADE20K数据集中,导致小目标区域被过度平滑、细节消失;当 $\beta = 1.0$ 时,能够较好地平衡边缘平滑与结构保持。据此分析,本文选取 $\alpha = 1.5, \beta = 1.0$ 为最佳值。实验14和实验15是其他语义分割常用损失函数与本文损失函数的比较。Dice Loss的优化目标是区域级重叠(整体匹配程度),而不是逐像素的分类精度,其依赖区域特性,对像素级分割精度不佳;Focal Loss对语义分割中类别不平衡问题表现良好,但其在梯度更新效率和物体分割边界等方面略显不足。本文损失从多角度出发,综合利用不同约束条件来提升模型性

表 6 多约束特征平滑损失分析
Tab. 6 Multi-constraint feature smoothing loss analysis

实验组	MCFS Loss		mIoU/%		Loss	
	α	β	Cityscapes	ADE20K	Cityscapes	ADE20K
13	0.0	0.0	81.06	43.41	0.178 9	0.702 9
	1.0	0.0	81.36	43.45	0.169 3	0.683 3
	1.5	0.0	81.38	43.51	0.167 1	0.682 5
	2.5	0.0	81.23	43.43	0.168 0	0.690 7
	0.0	0.5	81.14	43.34	0.179 6	0.708 7
	0.0	1.0	81.22	43.39	0.170 6	0.695 1
	0.0	1.5	81.19	43.36	0.178 1	0.702 8
	2.0	0.5	81.29	43.40	0.167 7	0.686 5
	1.5	1.0	81.41	43.63	0.163 2	0.679 6
14	Dice Loss		80.97	43.11	0.181 2	0.703 4
15	Focal Loss		81.18	43.37	0.172 1	0.698 6

能,提升边缘和过渡区域的平滑性。

4 结 论

本文提出全局感知与多尺度特征融合的城市道路语义分割网络,以解决城市道路语义分割任务中分割边界模糊、物体间相互遮挡、多尺度物体共存带来的分割精度不足问题。GAM通过在空间和通道上进行加权,感知全局信息,以细化分割边界;MSFF通过在不同语义层次的特征间共享信息,有效改善物体间相互遮挡和多尺度物体共存导致的分割精度不足;多约束特征平滑

损失 MCFS Loss 确保分割结果既符合高层语义,又具备细节一致性。经实验表明,该方法在 Cityscapes 和 ADE20K 数据集上表现出良好的性能,mIoU 值分别达到 81.4% 和 43.6%。但在模型推理效率以及轻量化方面略有欠缺,未来将尝试通过模型剪枝、蒸馏等方法进一步提升。

作者贡献声明:

邬开俊:论文写作指导,论文资助获取;
张治瑞:数据管理,论文构思与初稿写作;
汪滢:可视化呈现;
安立伟:论文审核。

参考文献:

- [1] 高常鑫,徐正泽,吴东岳,等.深度学习实时语义分割综述[J].中国图象图形学报,2024,29(5): 1119-1145.
GAO C X, XU Z Z, WU D Y, *et al.* Deep learning-based real-time semantic segmentation: a survey [J]. *Journal of Image and Graphics*, 2024, 29(5): 1119-1145. (in Chinese)
- [2] SHELHAMER E, LONG J, DARRELL T. Fully convolutional networks for semantic segmentation [C]. *IEEE Transactions on Pattern Analysis and Machine Intelligence*. May 24, 2016, IEEE, 2017: 640-651.
- [3] RONNEBERGER O, FISCHER P, BROX T. U-Net: Convolutional Networks for Biomedical Image Segmentation [M]. Medical Image Computing and Computer-Assisted Intervention-MICCAI 2015. Cham: Springer International Publishing, 2015: 234-241.
- [4] CHEN L C. Semantic image segmentation with deep convolutional nets and fully connected CRFs [J]. *arXiv preprint arXiv:1412.7062*, 2014.
- [5] CHEN L C, PAPANDREOU G, KOKKINOS I, *et al.* DeepLab: semantic image segmentation with deep convolutional nets, atrous convolution, and fully connected CRFs [J]. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2018, 40 (4): 834-848.
- [6] CHEN L C. Rethinking atrous convolution for semantic image segmentation [J]. *arXiv preprint arXiv:1706.05587*, 2017.
- [7] CHEN L C, ZHU Y K, PAPANDREOU G, *et al.* Encoder-decoder with atrous separable convolution for semantic image segmentation [C]. *Computer Vision-ECCV 2018*. Cham: Springer International Publishing, 2018: 833-851.
- [8] ZHAO H S, SHI J P, QI X J, *et al.* Pyramid scene parsing network [C]. 2017 *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*. July 21-26, 2017, Honolulu, HI, USA. IEEE, 2017: 6230-6239.
- [9] HE K M, GKIOXARI G, DOLLÁR P, *et al.* Mask R-CNN [C]. 2017 *IEEE International Conference on Computer Vision (ICCV)*. October 22-29, 2017, Venice, Italy. IEEE, 2017: 2980-2988.
- [10] REN S Q, HE K M, GIRSHICK R, *et al.* Faster R-CNN: towards real-time object detection with region proposal networks [J]. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2017, 39(6): 1137-1149.
- [11] 何家峰,陈宏伟,骆德汉.深度学习实时语义分割算法研究综述[J].计算机工程与应用,2023, 59(8): 13-27.
HE J F, CHEN H W, LUO D H. Review of real-time semantic segmentation algorithms for deep learning [J]. *Computer Engineering and Applications*, 2023, 59(8): 13-27. (in Chinese)
- [12] OKTAY O, SCHLEMPER J, FOLGOC L L, *et al.* Attention u-net: Learning where to look for the pancreas [J]. *arXiv preprint arXiv: 1804.03999*, 2018.
- [13] HU J, SHEN L, SUN G. Squeeze-and-excitation Networks [C]. 2018 *IEEE/CVF Conference on Computer Vision and Pattern Recognition*. June 18-

- 23, 2018, Salt Lake City, UT, USA. IEEE, 2018: 7132-7141.
- [14] WOO S, PARK J, LEE J Y, *et al.* CBAM: Convolutional block attention module [C]. *Computer Vision-ECCV 2018. Cham: Springer International Publishing*, 2018: 3-19.
- [15] 任凤雷, 杨璐, 周海波, 等. 基于改进 BiSeNet 的实时图像语义分割[J]. *光学精密工程*, 2023, 31(8): 1217-1227.
- REN F L, YANG L, ZHOU H B, *et al.* Real-time semantic segmentation based on improved BiSeNet[J]. *Opt. Precision Eng.*, 2023, 31(8): 1217-1227. (in Chinese)
- [16] 侯志强, 程敏捷, 马素刚, 等. 基于跨层次聚合网络的实时城市街景语义分割[J]. *光学精密工程*, 2024, 32(8): 1212-1226.
- HOU Z Q, CHENG M J, MA S G, *et al.* Real-time urban street view semantic segmentation based on cross-layer aggregation network[J]. *Opt. Precision Eng.*, 2024, 32(8): 1212-1226. (in Chinese)
- [17] 闫河, 雷秋霞, 王旭. 融合注意力机制的改进型 DeepLabv3+ 语义分割[J]. *光学精密工程*, 2025, 33(1): 123-134.
- YAN H, LEI Q X, WANG X. Improved DeepLabv3+ semantic segmentation incorporating attention mechanisms [J]. *Opt. Precision Eng.*, 2025, 33(1): 123-134. (in Chinese)
- [18] LIU Z, MAO H Z, WU C Y, *et al.* A convnet for the 2020s [C]. *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). June 18-24, 2022, New Orleans, LA, USA. IEEE*, 2022: 11966-11976.
- [19] DING X H, ZHANG X Y, HAN J G, *et al.* Scaling up your kernels to 31×31 : revisiting large kernel design in CNNs [C]. *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). June 18-24, 2022. New Orleans, LA, USA. IEEE*, 2022: 11963-11975.
- [20] LIU S, CHEN T, CHEN X, *et al.* More convnets in the 2020s: Scaling up kernels beyond 51×51 using sparsity [J]. *arXiv preprint arXiv: 2207.03620*, 2022.
- [21] SERT E, ALKAN A. Image edge detection based on neutrosophic set approach combined with Chan-veese algorithm [J]. *International Journal of Pattern Recognition and Artificial Intelligence*, 2019, 33(3): 1954008.
- [22] SUDRE C H, LI W Q, VERCAUTEREN T, *et al.* Generalised dice overlap as a deep learning loss function for highly unbalanced segmentations [C]. *Deep Learning in Medical Image Analysis and Multimodal Learning for Clinical Decision Support. Cham: Springer International Publishing*, 2017: 240-248.
- [23] LIN T Y, GOYAL P, GIRSHICK R, *et al.* Focal loss for dense object detection [C]. *2017 IEEE International Conference on Computer Vision (ICCV). October 22-29, 2017, Venice, Italy. IEEE*, 2017: 2999-3007.
- [24] CORDTS M, OMRAN M, RAMOS S, *et al.* The cityscapes dataset for semantic urban scene understanding [C]. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. Los Alamitos: IEEE Computer Society Press*, 2016: 3213-3223.
- [25] ZHOU B L, ZHAO H, PUIG X, *et al.* Scene parsing through ADE20K dataset [C]. *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). July 21-26, 2017, Honolulu, HI, USA. IEEE*, 2017: 5122-5130.
- [26] HONG Y, PAN H, SUN W, *et al.* Deep dual-resolution networks for real-time and accurate semantic segmentation of road scenes [J]. *arXiv preprint arXiv:2101.06085*, 2021.
- [27] VASWANI A, SHAZEER N, PARMAR N, *et al.* Attention is all you need [J]. *Advances in neural information processing systems*, 2017, 30.
- [28] XIE E Z, WANG W H, YU Z D, *et al.* SegFormer: simple and efficient design for semantic segmentation with transformers [C]. *Neural Information Processing Systems*, 2021.
- [29] PAN H H, HONG Y D, SUN W C, *et al.* Deep dual-resolution networks for real-time and accurate semantic segmentation of traffic scenes [J]. *IEEE Transactions on Intelligent Transportation Systems*, 2023, 24(3): 3448-3460.
- [30] GUO M H, LU C Z, HOU Q, *et al.* Segnext: Rethinking convolutional attention design for semantic segmentation [J]. *Advances in Neural Information Processing Systems*, 2022, 35: 1140-1156.
- [31] WAN Q, HUANG Z, LU J, *et al.* Seaformer: squeeze-enhanced axial transformer for mobile se-

- mantic segmentation[C]. *The eleventh international conference on learning representations*, 2023.
- [32] WANG J, GOU C, WU Q, *et al.* RTFormer: Efficient design for real-time semantic segmentation with transformer[J]. *Advances in Neural Information Processing Systems*, 2022, 35: 7423-7436.
- [33] XU Z, WU D, YU C, *et al.* SCTNet: Single-Branch CNN with transformer semantic information for real-time segmentation[C]. *Proceedings of the AAAI Conference on Artificial Intelligence*. 2024, 38(6): 6378-6386.

作者简介:



邬开俊(1978—),男,山东莒南人,教授,博士生导师,现为兰州交通大学电子与信息工程学院副院长,主要研究方向为深度学习、计算机视觉。E-mail:wkj@mail.lzjtu.cn

通讯作者:



张治瑞(1998—),男,甘肃金昌人,硕士研究生,2022年于江西理工大学软件工程学院获取工学学士学位,现就读于兰州交通大学,主要研究方向为计算机视觉、语义分割。E-mail:zzr_ljd@163.com