

DOI:10.16136/j.joel.2023.02.0204

多尺度残差特征融合的轻量级真实图像超分辨率重建

吕 佳^{1,2*}, 许鹏程^{1,2}

(1. 重庆师范大学 计算机与信息科学学院, 重庆 401331; 2. 重庆师范大学 重庆市数字农业服务工程技术研究中心, 重庆 401331)

摘要: 基于深度学习的真实图像超分辨率 (super-resolution, SR) 重建算法目前存在参数量过大的问题, 为解决该问题, 提出了一种多尺度残差特征融合的轻量级真实图像 SR 重建算法。首先利用深度可分离卷积和复用卷积针对多尺度特征提取块进行改进, 在提取特多尺度特征的同时实现了模块的轻量化, 参数量仅为改进前的 7.5%。其次使用残差特征融合操作将 4 个多尺度深度可分离特征提取块 (multi-scale depthwise separable block, MSDSB) 聚合成一个残差特征融合块, 以减少残差路径长度。然后使用增强型注意力模块从通道和空间维度进行自适应调整以提升算法性能。最后使用自适应上采样模块获得 SR 重建图像。在消融实验中, 本文算法重建性能超过原始算法, 且参数量仅为 3.53×10^6 , 是原始算法的 34.5%。在对比实验中, 其重建性能超过了当前主流算法, 与组件分而治之 (component divide-and-conquer, CDC) 算法相比, PSNR 和 SSIM 指标分别提升了 0.01 dB 与 0.0010, 且参数量仅为组件 CDC 算法的 8.84%, 在保证重建性能的同时实现了算法的轻量化。

关键词: 真实图像; 图像超分辨率 (SR) 重建; 卷积神经网络; 深度可分离卷积; 残差特征融合
中图分类号: TP391.41 **文献标识码:** A **文章编号:** 1005-0086(2023)02-0120-12

Lightweight real-world image super-resolution reconstruction based on multi-scale residual feature aggregation

LV Jia^{1,2*}, XU Pengcheng^{1,2}(1. College of Computer and Information Sciences, Chongqing Normal University, Chongqing 401331, China;
2. Chongqing Center of Engineering Technology Research on Digital Agriculture Service, Chongqing Normal University, Chongqing 401331, China)

Abstract: At present, there exists too many parameters amount in the real-world image super-resolution (SR) reconstruction algorithm based on deep learning. To solve this problem, a lightweight real-world image SR reconstruction algorithm based on multi-scale residual feature aggregation is proposed. First, the depthwise separable convolution and multiplexing convolution are used to improve the existing multi-scale feature extraction block, which achieves the lightweight of the module while extracting the extra-multi-scale feature, with only 7.5% of the parameters amount before the improvement. Next, the residual feature aggregation is exploited to aggregate the 4 multi-scale depthwise separable blocks amount (MSDSB) into a residual feature aggregation block to reduce the length of the residual path. Then, the enhanced attention module is utilized to adaptively adjust the channel and spatial dimensions to improve the performance of the algorithm. Finally, the adaptive upsampling module is used to obtain SR reconstructed images. In ablation experiments, the reconstruction performance of the algorithm is better than that of the original algorithm, and the parameters amount is only 3.53×10^6 , which is 34.5% of the orig-

* E-mail: lvjia@cqnu.edu.cn

收稿日期: 2022-03-29 修订日期: 2022-04-29

基金项目: 国家自然科学基金重大项目 (11971084)、重庆市高校创新研究群体资助 (CXQT20015)、重庆市教委科研项目重点项目 (KJZD-K202200511) 和重庆市科技局技术预见与制度创新项目 (2022TFII-OFX0265) 资助项目

inal algorithm. In the comparative experiments, the reconstruction performance of the proposed algorithm is better than the current mainstream algorithm. Compared with the component divide-and-conquer (CDC) algorithm, the PSNR and SSIM indexes of the presented algorithm are increased by 0.01dB and 0.0010, respectively, and the parameters amount is only 8.84% of that of the CDC algorithm. The lightweight of the algorithm is realized while ensuring the reconstruction performance.

Key words: real-world image; image super-resolution (SR) reconstruction; convolutional neural network; depthwise separable convolution; residual feature aggregation

0 引言

图像超分辨率(super-resolution, SR)重建是计算机视觉中底层视觉的一个重要分支,是将低分辨率图像(low resolution, LR)恢复成高分辨率图像(high resolution, HR)的过程^[1]。SR重建主要包含基于插值、重构等传统方法以及基于机器学习和深度学习的学习方法。基于插值的方法有双三次插值(bicubic)等,虽然计算速度快,但重建的细节有限,会产生模糊等问题。基于重建的方法有迭代反投影法等,虽具有较快的计算速度和重建精度,但易产生重建图像平滑、模糊等问题^[2]。

2014年,DONG等将深度学习引入SR重建之中,提出了基于卷积神经网络的SR重建(super-resolution convolutional neural network, SRCNN)算法,采用3个卷积层即获得了比传统方法更好的性能^[3]。2016年,KIM等提出基于全局残差学习的超深卷积网络SR重建(super-resolution using very deep convolutional network, VDSR)算法,解决了超深网络中的梯度消失与梯度爆炸的问题^[3]。2017年,LEDIG等^[4]基于残差学习提出局部残差与全局残差结合的SRResNet(SR residual network)算法,在加深网络的同时可以学习更多的全局特征信息。受此启发,LIM等^[5]在SRResNet的基础上提出了增强型深度网络SR重建(enhanced deep super-resolution, EDSR)算法,在去除网络中的冗余部分的同时设置多条路径用于适配不同上采样系数。2018年,LI等^[6]基于多尺度学习提出了多尺度残差网络(multi-scale residual network, MSRN)算法,通过在每个特征提取块中设置多条不同尺寸的卷积核路径以提取多尺度特征,使其在不使用任何权重初始化和训练技巧的前提下依然具有优异的性能。2020年,LIU等^[7]提出基于残差特征融合的算法(residual feature aggregation network, RFANet),使用残差特征融合模块来更好地利用局部残差特征,同时结合改进的空间注意力(enhanced spatial attention, ESA)模块在多种图像退化方法上取得良好的重建效果。同年,LIU等^[8]基于残差特征蒸馏与ESA模块提出了超轻量图像SR算法RFDN(residual feature distillation

network),仅使用 0.6×10^6 参数量便获得了良好的性能。2021年,KONG等^[9]提出针对图像不同区域执行不同重建方法的ClassSR(Classification-SR)算法,在减少算法计算量的同时提升了性能。

然而,以上算法均通过已知模糊核等固定退化方法在基准数据集上生成的HR-LR图像对中训练得到,导致SR算法的泛化性较差,无法应用在实际场景中。因此,基于真实图像的SR重建成为目前的研究热点。真实图像的SR重建的HR和LR图像均通过成像设备收集,其退化函数未知,因此更具挑战性和实际应用价值。2019年,CHEN等^[10]通过分析成像系统中图像分辨率和视场之间的关系,提出了分辨率-视野域退化模型。2020年,WEI等^[11]基于角点检测提出了组件分而治之(component divide-and-conquer, CDC)算法,分别针对图像的平坦区域、边缘和角点执行重建,并组合成为SR图像。此外,WEI等提出了大规模真实图像SR数据集DRealSR,成为真实图像SR重建的基准数据集之一。然而,以上基于真实图像的SR算法参数量均在 10×10^6 以上,具有较大的参数量。虽然庞大的参数量可以保证算法的重建性能,但不利于算法的部署与应用。因此,针对真实图像SR算法的轻量化改进是当前需解决的问题之一。

为此,本文提出一种基于深度可分离卷积的轻量级多尺度残差特征融合的真实图像SR重建算法,主要改进如下:首先提出一种结合深度可分离卷积和复用卷积的轻量化多尺度深度特征提取块,在提取图像的多尺度特征的同时大幅度减少参数量;其次引入残差特征融合块将多个特征提取块组合,减少了残差路径的长度,同时实现了残差特征的非局部使用;然后提出一种改进的注意力模块,从空间和通道两个维度进行自适应调整以提升算法的重建性能;最后使用自适应上采样模块对图像执行上采样操作以获得最终的SR图像。本文算法的训练采用DRealSR数据集,从而满足真实图像的重建需求。实验结果证明,本文算法在重建性能上超过CDC等主流算法;在参数量上仅为CDC算法的8.84%;在泛化性上使用复杂退化函数时具有良好的重建性能。

1 基本原理

1.1 深度可分离卷积

深度可分离卷积是轻量级网络中常用的操作^[12]。深度可分离卷积分为两个步骤:逐通道卷积和逐点卷积。与传统卷积相比,深度可分离卷积可在轻微损失性能的前提下大幅度减少算法的参数数量。

在普通卷积中,假设输入一个通道数为 C_{in} 的特征图,卷积核大小为 k ,输出通道数为 C_{out} ,则该卷积层的参数量 N_{Conv} 如式(1)所示:

$$N_{Conv} = C_{in} \times C_{out} \times k \times k. \quad (1)$$

在逐通道卷积中,首先对每个通道执行单通道卷积,然后将输出的单通道特征图重新堆叠。该步骤仅调整输入特征图的尺寸,通道数不发生变化,因此其参数量为 $C_{in} \times k \times k$,执行过程如图1所示。

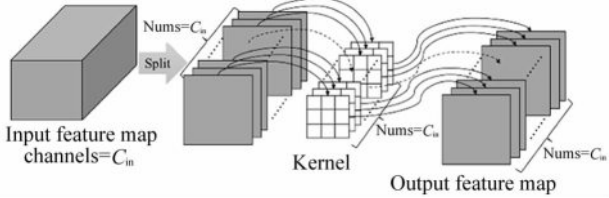


图1 逐通道卷积示意图

Fig.1 Depthwise convolution sketch

逐点卷积即 1×1 卷积,对逐通道卷积生成的特征图沿通道维度执行加权组合,从而生成新的特征图。逐点卷积的参数量为 $C_{in} \times C_{out} \times 1 \times 1$,其执行过程如图2所示。

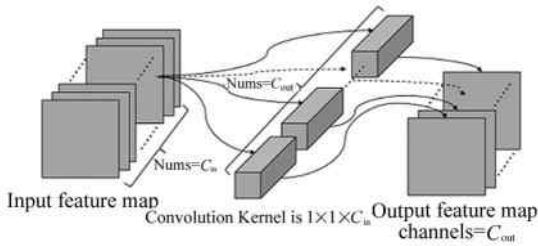


图2 逐点卷积示意图

Fig.2 Pointwise convolution sketch

因此,深度可分离卷积的参数量如式(2)所示:

$$N_{DS} = C_{in} \times (C_{out} + k \times k). \quad (2)$$

由式(2)可知,深度可分离卷积的参数量为传统卷积的 $1/k^2 + 1/C_{out}$,具有轻量化优势。

1.2 注意力机制

注意力机制根据特征的重要性来分配权重,使得算法在更关注与输出相关特征的同时弱化无关特征^[13]。注意力机制可分为通道注意力机制^[14]、空间注意力机制^[13]、自注意力机制^[15]等。

然而,以上注意力机制模块仅针对通道域或空间域分配权重,对性能的提升有限。因此,WOO等^[16]提出一种基于卷积块的注意力模块(convolutional block attention module, CBAM)。CBAM 可以从空间维度与通道维度对特征图进行自适应调整,相对于只针对通道域或空间域的注意力算法能够提升更多的性能。

2 本文算法

本文基于深度可分离卷积对现有的多尺度通道注意力特征提取块(multi-scale with channel attention block, MSCAB)^[17]进行轻量化改进,利用改进后的多尺度深度可分离特征提取块(multi-scale depthwise separate block, MSDSB)提取 LR 图像的多尺度特征。为了防止过长的残差路径导致网络退化问题,引入了残差特征融合^[7]模块,将4个 MSDSB 融合成一个多尺度残差特征融合块(multi-scale with residual feature aggregation block, MSRFAB),可以有效减少残差路径的长度,提升算法性能。为了平衡轻量化改进与算法性能之间的矛盾,提出了增强型注意力模块,从空间和通道两个维度对特征图进行自适应调整。最后使用自适应上采样模块^[17]将特征图放大至目标放大系数得到 SR。算法整体结构如图3所示。

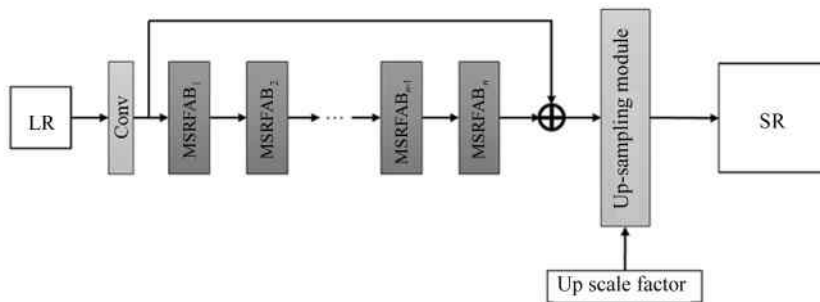


图3 本文算法网络结构图

Fig.3 The structure diagram of proposed algorithm network

2.1 MSDSB

MSDSB 以 MSCAB 为基础,在特征提取块中构建了多条不同的分支,每条分支中包含不同数量的深度可分离卷积层,以此来提取图像的多尺度特征。同时,为了进一步减少参数量,合并了不同支路中相同位置的卷积层,并对其输出的特征图进行级联。MSDSB 可在减少卷积层使用的情况下提取多尺度特征,其结构如图 4 所示。

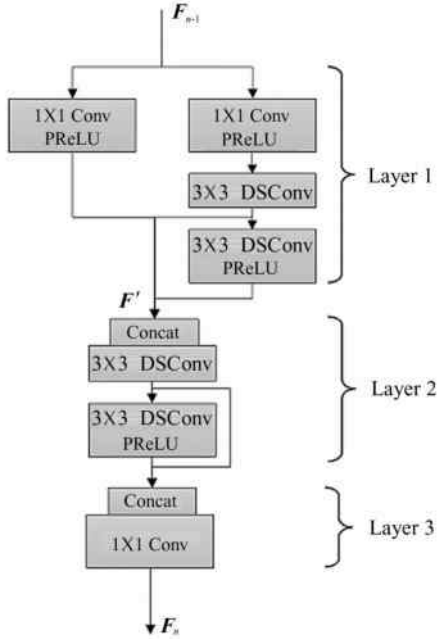


图 4 MSDSB 结构图

Fig. 4 The structure diagram of MSDSB

MSDSB 首先使用 1×1 卷积对特征图进行升降维操作,然后使用两个深度可分离卷积提取特征。特征提取完毕后依次级联每一个深度可分离卷积的输出特征图。由于 2 个 3×3 卷积的感受野与 1 个 5×5 卷积一致,因此第一层中第二个深度可分离卷积输出的特征图感受野为 5×5 ,而第一个深度可分离卷积输出特征图的感受野为 3×3 ,级联特征图后即可获得 3×3 感受野与 5×5 感受野的特征图,以此实现多尺度特征提取。之后将级联后的特征图送入第二层中执行相同的步骤,进一步提取多尺度特征。最后通过第三层的 1×1 卷积对输出特征图执行降维操作,使其输出通道数与输入通道数相同,获得最终的输出特征图。模块按照下式执行:

$$\sigma(x) = \max(0, x) + \min(ax, 0), \quad (3)$$

$$S_{1,1} = \sigma(W_{1 \times 1}^{1,1} \otimes F_{n-1} + b^{1,1}), \quad (4)$$

$$S_{1,2} = \sigma(W_{1 \times 1}^{1,2} \otimes F_{n-1} + b^{1,2}), \quad (5)$$

$$S_{1,3} = W_{3 \times 3}^{1,3} \otimes S_{1,2} + b^{1,3}, \quad (6)$$

$$S_{1,4} = \sigma(W_{3 \times 3}^{1,4} \otimes S_{1,3} + b^{1,4}), \quad (7)$$

$$F' = [S_{1,1}, S_{1,3}, S_{1,4}], \quad (8)$$

$$S_{2,1} = W_{3 \times 3}^{2,1} \otimes F' + b^{2,1}, \quad (9)$$

$$S_{2,2} = \sigma(W_{3 \times 3}^{2,2} \otimes S_{2,1} + b^{2,2}), \quad (10)$$

$$F_n = W_{1 \times 1}^{3,1} \otimes [S_{2,1}, S_{2,2}] + b^{3,1}, \quad (11)$$

式(3)表示带参数的线性整流函数(parametric rectified linear unit, PReLU)^[18],函数中的 a 由模型学习得来。式(4)一式(11)中的 W 和 b 表示卷积层中的卷积核与偏置参数,其上标的第 1 个参数表示其所在的层数,第 2 个参数表示其来自于该层的第几个卷积层; W 的下标表示其卷积核的大小; $\otimes$ 表示卷积操作。式(8)和式(11)中的 $[X_1, X_2, \dots, X_n]$ 表示级联操作,沿通道维度对特征图执行级联。式中的 S 表示卷积层输出的特征图,其下标表示其来自于第几层的第几个卷积。MSDSB 的参数量为 0.09×10^6 ,而 MSCAB 的参数量为 1.2×10^6 ,MSDSB 的参数量是后者的 7.5%,是本文算法轻量化的前提。

2.2 残差特征融合块

残差特征融合块将多个 MSDSB 结合,采用残差连接的方式以缓解因网络深度增加而导致的退化问题。同时可以降低算法训练难度,提高学习能力。在基于残差学习的传统网络中,第一个残差块需通过较长的路径将特征图送至最后的模块,导致残差特征难以充分利用。因此采用残差特征融合块将局部特征进行更好地利用,其结构如图 5 所示。

文献[7]提出的残差特征融合块中包含了 4 个残差块,因此本文使用的残差特征融合块也包含 4 个 MSDSB。残差特征融合块将前 3 个 MSDSB 的残差特征直接传至末端,与最后一个 MSDSB 的输出级联,再通过 1×1 卷积融合以上特征并降低通道数,与文献[7]一致。模块按照如下式执行:

$$F_1 = \text{MSDSB}_1(M_{n-1}), \quad (12)$$

$$M' = F_1 + M_{n-1}, \quad (13)$$

$$F_2 = \text{MSDSB}_2(M'), \quad (14)$$

$$M' \leftarrow F_2 + M', \quad (15)$$

$$F_3 = \text{MSDSB}_3(M'), \quad (16)$$

$$M' \leftarrow F_3 + M', \quad (17)$$

$$F_4 = \text{MSDSB}_4(M'), \quad (18)$$

$$M_n' = W_{1 \times 1} \otimes [F_1, F_2, F_3, F_4] + b, \quad (19)$$

$$M_n = M_n' + M_{n-1}, \quad (20)$$

式中, M' 表示残差特征融合块中的中间特征,在模块的执行过程中不断更新, F 表示每一个 MSDSB 的输出特征图,其下标表示该特征图来自于第几个 MSDSB。

残差特征融合块可以有效减少残差路径的长度,与简单叠加多个残差块的方法相比,该模块可以实现残差特征的非局部使用。前3个MSDSB输出的有效信息可以在没有任何损失或干扰的情况下传至残差特征融合块末端,从而实现更加有效的特征表示。

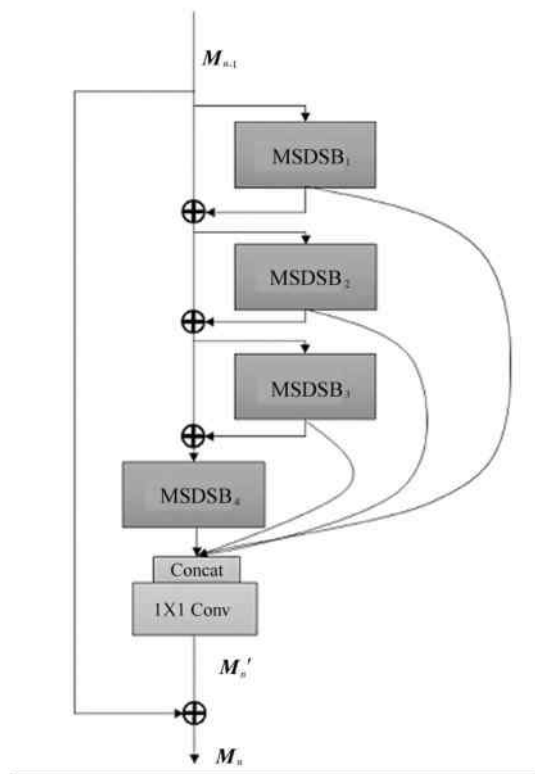


图5 残差特征融合块结构图

Fig. 5 The structure of residual feature aggregation block

2.3 增强型注意力模块

为了进一步提高多尺度残差特征融合特征提取模块的性能,在每个残差特征融合块的末端加入增强型注意力机制模块。增强型注意力机制模块类似于CBAM,其结合通道注意力和空间注意力,从空间和通道两个维度对特征图进行自适应调整以提升算法性能。增强型注意力机制模块结构如图6所示。

首先,针对残差特征融合模块输出的特征图进行一次全局平均池化(global average pooling, GAP),对特征进行向量化处理以获得特征向量,再通过两个全连接层和一个激活函数自适应地建立通道间的相互关系。然后使用Sigmoid函数将特征向量压缩至0—1之间,并加权处理原特征通道以获得经过通道维度校准的特征图。之后,将经过通道注意力模块处理后的特征图进行1×1卷积降维至1通

道,依次执行两次11×11卷积,通过大感受野获取图像的空间关系的同时,第二个11×11卷积将通道数降为1。使用Sigmoid函数将特征图压缩至0—1之间生成空间维度的掩码,最后将掩码和特征图相乘获得经模块校准后的特征图。

增强型注意力机制模块的输入和输出通道数与残差特征融合块的输出通道数相同,因此可以直接添加在MSRFB的结尾。如图7所示,每一个MSRFB中均包含增强型注意力模块。

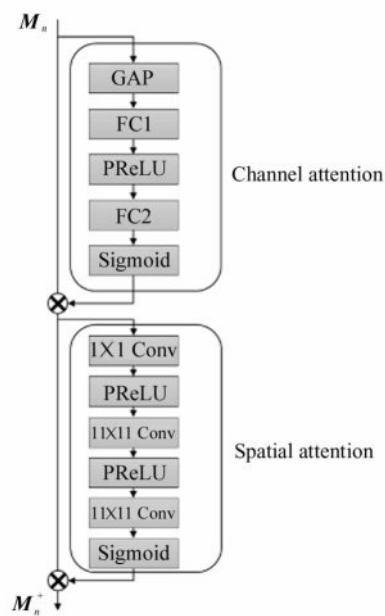


图6 增强型注意力模块结构图

Fig. 6 The structure of enhanced attention module

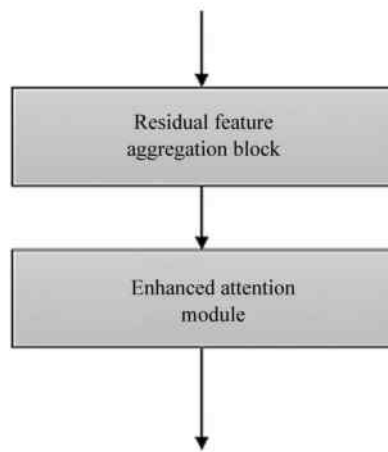


图7 MSRFB结构图

Fig. 7 The structure of MSRFB

3 实验及结果分析

3.1 实验环境

本文算法的硬件实验环境为AMDRyzenR9-

3900X CPU,64 GB 内存,Nvidia GeForce RTX2080 和 Nvidia Tesla M40 GPU;操作系统为 Windows 10 21H1 专业版 64 位;HR-LR 图像对的生成使用了 Matlab R2019b;算法的训练采用 PyTorch 1.5.1 框架,测试使用 Numpy 1.19.5 科学计算库。

3.2 评价指标

为了验证本文算法的有效性,采用峰值信噪比 (peak signal to noise ratio,PSNR)和结构相似性指标(structural similarity,SSIM)作为客观评价标准。考虑到比较的公平性,消融实验中的指标在计算时先将图像色彩空间从 RGB 转换为 YCbCr,然后在其 Y 通道进行,同时从图像的每个边缘裁剪 s 个像素^[6], s 为放大系数;在对比实验中,使用文献[6]的策略,在 YCbCr 色彩空间的 Y 通道计算 PSNR 指标,在 RGB 色彩空间计算 SSIM 指标。

3.3 训练细节及参数设置

本文的消融实验采用 SR 使用最广泛的 DIV2K 作为训练集,Set5、Set14 和 Urban100 作为测试集^[17]。首先对训练集中的 HR 图像进行 90°、180°和 270°旋转,得到 3 600 张 HR 图像,然后使用 Matlab 中的 imresize 函数对 HR 图像执行双三次插值退化获得 HR-LR 图像对。分别采用 2、3、4 倍放大系数训练,每个放大系数包含 1 000 个轮次。训练时将 LR 图像随机裁剪成 48×48 的图像块,对应的 HR 图像大小为 48 s ×48 s , s 为放大系数。每次迭代送入 8 个图像块,每个轮次包含 450 次迭代。设置初始学习率为 0.000 1,每 250 轮次后下降至原先的 0.5 倍。算法采用 8 个 MSRFAB、Charbonnier 损失函数和 Adam 优化器进行训练。训练开始时对每个 MSRFAB 中的增强型注意力模块执行 Xavier 初始化,在每一个放大系数训练完成后,使用当前放大系数的模型作为下一个放大系数的预训练模型继续训练。

本文的对比实验采用真实图像数据集 DRealSR 的训练集进行训练,使用 DRealSR 和 RealSR^[19]的测试集对算法性能进行测试,在训练过程中不采用任

何数据集扩充策略。使用 2、3、4 倍放大系数进行训练,每个放大系数包含 500 个轮次,每轮次平均包含 3 821 次反向传播迭代。设置初始学习率为 0.000 1,每 125 轮次后下降至原先的 0.5 倍。其余训练策略与消融实验保持一致。

3.4 实验结果分析

3.4.1 消融实验

为了验证本文算法中提出的不同模块的有效性,在不使用 Xavier 初始化的前提下,使用 2、3 倍放大系数在 Set5、Set14 测试集上对算法进行消融实验。其中,原始算法的实验结果来自于文献[17],打勾表示使用了该模块。实验结果如表 1 所示。

从表 1 的结果可以得出,使用深度可分离卷积和复用卷积后,模型的参数量急剧下降,但是重建性能也随之下降。将多尺度特征提取块数量提升至 32 后算法性能有了一定的提升,达到原始算法的性能水平。将 4 个多尺度特征提取块聚合成一个残差特征融合块,使得算法的性能获得轻微的提升。相较于不使用残差特征融合,算法的 PSNR 指标平均提升了 0.01 dB。增强型注意力机制模块从通道和空间维度对特征图进行自适应调整,在增强与重建相关特征的同时屏蔽与重建不相关的特征。增强型注意力机制模块使得算法的 PSNR 指标提升了 0.03 dB,SSIM 指标提升了 0.000 3。本文算法参数量为 3.53×10^6 ,仅为原始算法的 34.5%的基础上 PSNR 指标提升了 0.05 dB,SSIM 指标提升了 0.000 6,在保证重建性能的同时实现了算法的轻量化。

为了验证 Xavier 初始化对本文算法性能的影响,针对算法的不同模块使用 Xavier 初始化。在 2、3 倍放大系数下在 Set5、Set14 测试集上对算法进行消融实验。其中,该部分消融实验全部使用深度可分离卷积,“全部初始化”表示针对算法的所有模块使用 Xavier 初始化,“注意力初始化”表示仅针对增强型注意力模块执行 Xavier 初始化。实验结果如表 2 所示。

表 1 本文算法不同模块有效性 PSNR/dB|SSIM 结果

Tab. 1 The effectiveness PSNR/dB|SSIM results of different modules of the proposed algorithm

DSConv	Reuse conv	RFA	Enhance attention	Numbers of MSDSB	Parameters / $\times 10^6$	Scale $\times 2$		Scale $\times 3$	
						Set5	Set14	Set5	Set14
				8	10.23	38.12 0.9610	33.76 0.9186	34.46 0.9274	30.32 0.8412
✓				8	1.92	37.99 0.9606	33.61 0.9184	34.21 0.9256	30.21 0.8399
✓	✓			8	1.22	37.95 0.9604	33.59 0.9176	34.15 0.9251	30.17 0.8388
✓	✓			32	3.30	38.12 0.9610	33.85 0.9203	34.36 0.9270	30.36 0.8427
✓	✓	✓		32	3.43	38.13 0.9611	33.81 0.9194	34.42 0.9271	30.37 0.8430
✓	✓	✓	✓	32	3.53	38.15 0.9611	33.85 0.9198	34.47 0.9277	30.37 0.8421

表2 Xavier初始化对本文算法重建性能影响的PSNR/dB|SSIM结果

Tab. 2 PSNR/dB|SSIM results of Xavier initialization on the reconstruction performance of the proposed algorithm

Reuse conv	RFA	Enhance attention	Numbers of MSDSB	In all modules	In attention modules	Scale×2		Scale×3	
						Set5	Set14	Set5	Set14
			8			37.99 0.9606	33.61 0.9184	34.21 0.9256	30.21 0.8399
			8	✓		37.97 0.9605	33.56 0.9175	34.16 0.9255	30.18 0.8393
✓			8	✓		37.88 0.9602	33.48 0.9167	34.07 0.9245	30.13 0.8379
✓			32	✓		38.07 0.9609	33.78 0.9193	34.30 0.9263	30.27 0.8408
✓	✓		32	✓		38.10 0.9610	33.83 0.9201	34.37 0.9270	30.31 0.8417
✓	✓	✓	32			38.15 0.9611	33.85 0.9198	34.47 0.9277	30.37 0.8421
✓	✓	✓	32	✓		38.14 0.9611	33.89 0.9206	34.42 0.9273	30.36 0.8423
✓	✓	✓	32		✓	38.15 0.9612	33.94 0.9210	34.46 0.9277	30.37 0.8424

结合表1和表2的实验结果可以看出,针对算法的所有模块执行Xavier初始化后的重建性能并不理想,其原因在于文献[17]中提到多尺度特征提取模块不需要使用任何权重初始化方法,因此对其执行初始化的操作会影响其性能;仅针对增强型注意力模块执行Xavier初始化后可以使得增强型注意力模块收敛速度快于残差特征融合块,从而指导残差特征融合块更加有效地提取相关特征,因此提升了算法的性能。相对于不使用Xavier初始化,针对增强型注意力模块执行初始化的操作使得算法PSNR指标平均提升0.02 dB,SSIM指标平均提升0.0004。

为了验证MSRFAB数量对算法性能的影响,分别使用1、2、4、8、16、24个MSRFAB进行实验,并在3、4倍放大系数下使用Set5、Set14、Urban100作为测试集进行测试。同时分别验证了增强型注意力模

块是否进行Xavier初始化对算法性能的影响。实验结果如表3所示,其中最好的实验结果加粗表示,第二好的实验结果用下划线表示。

从表3的实验结果可以看出,在不使用初始化时,随着MSRFAB的增加,算法重建性能随之增加。当MSRFAB数量为16时,重建性能达到顶峰,然而其参数量为 6.54×10^6 ,超过了轻量级算法 5×10^6 参数量的分界线,因此不具备轻量化优势;而针对增强型注意力模块使用初始化时,8个MSRFAB的算法重建性能仅次于上述的16个MSRFAB的算法,参数量仅为其54%,即 3.53×10^6 ,具有轻量化优势。因此,本文采用8个MSRFAB作为特征提取模块。当MSRFAB数量超过16时,由于出现主干网络残差路径过长的问题,从而导致特征传播困难和网络退化现象的发生,因此当MSFRAB数量大于16时,

表3 不同MSRFAB对本文算法重建性能影响的PSNR/dB|SSIM结果

Tab. 3 PSNR/dB|SSIM results of the impact of different MSRFAB on the reconstruction performance of the proposed algorithm

Numbers of MSDSB	Parameters / $\times 10^6$	Attention initialized	Scale×3			Scale×4		
			Set5	Set14	Urban100	Set5	Set14	Urban100
1	0.90		33.99 0.9243	30.08 0.8372	27.49 0.8382	31.90 0.8910	28.39 0.7771	25.60 0.7702
1	0.90	✓	34.00 0.9243	30.10 0.8372	27.50 0.8380	31.93 0.8914	28.43 0.7773	25.61 0.7703
2	1.27		34.18 0.9257	30.18 0.8388	27.72 0.8429	32.04 0.8932	28.49 0.7794	25.81 0.7776
2	1.27	✓	34.19 0.9255	30.21 0.8394	27.77 0.8444	32.06 0.8934	28.51 0.7798	25.85 0.7791
4	2.03		34.35 0.9269	30.32 0.8415	28.12 0.8520	32.26 0.8959	28.63 0.7826	26.16 0.7890
4	2.03	✓	34.34 0.9267	30.30 0.8413	28.05 0.8507	32.24 0.8956	28.59 0.7822	26.13 0.7878
8	3.53		<u>34.47</u> <u>0.9277</u>	30.37 0.8421	28.30 0.8549	<u>32.29</u> <u>0.8966</u>	28.69 0.7841	26.33 0.7934
8	3.53	✓	34.46 <u>0.9277</u>	<u>30.37</u> <u>0.8424</u>	<u>28.32</u> <u>0.8563</u>	32.27 0.8964	<u>28.70</u> <u>0.7844</u>	<u>26.35</u> <u>0.7945</u>
16	6.54		34.56 0.9286	30.43 0.8441	28.52 0.8599	32.40 0.8977	28.76 0.7859	26.49 0.7992
16	6.54	✓	32.94 0.9121	29.00 0.8190	25.50 0.7803	31.15 0.8790	27.50 0.7557	24.09 0.7081
24	9.56		32.39 0.9057	29.07 0.8174	25.72 0.7862	30.22 0.8586	27.34 0.7493	24.16 0.7104
24	9.56	✓	32.27 0.9004	28.99 0.8168	25.68 0.7849	30.23 0.8586	27.19 0.7458	24.02 0.7078

算法的重建性能急剧下降。

此外,虽然针对增强型注意力模块执行 Xavier 初始化的操作可以提升算法重建性能,但会导致增强型注意力模块和残差特征融合块收敛速度不一致。在 MSRFAB 数量较少时可以更快、更好地针对残差特征融合块进行自适应调整,从而提升性能;但是随着 MSRFAB 数量的增加,增强型注意力模块和残差特征融合块收敛速度不一致的缺点会逐渐显现。当 MSRFAB 数量大于 8 时,由于主干网络的路径较长,其收敛速度不一致的特性会导致特征传播困难问题的出现,从而影响算法性能。因此当 MSRFAB 数量为 16 或 24 时,针对增强型注意力模块执行 Xavier 初始化的操作会导致算法的重建性能急剧下降。

3.4.2 对比实验

为了测试本文算法在真实图像中的重建效果,在 2、3、4 倍放大系数下对各测试集进行重建,并从 PSNR 和 SSIM 两个维度进行评判。使用 Bicubic、SRResNet、EDSR、ESRGAN (enhanced super resolution generative adversarial network)^[20]、RCAN (residual channel attention network)^[21]、CDC、RFDN-L^[8] 和 MSCAN^[17] 作为对比算法。其中, Bicubic、SRResNet、EDSR、ESRGAN、RCAN 和 CDC 的实验结果来自于文献[11]; RFDN-L 与 MSCAN 的实验结果分别采用文献[8]与文献[17]的训练策略,使用 DRealSR 作为训练集重新训练得到。实验结果如表 4 所示,最好的实验结果加粗表示,第二好的实验结果用下划线表示。

从表 4 的实验结果可以看出,本文算法的 PSNR 与 SSIM 指标相对于 Bicubic 分别提升了 1.37 dB 和 0.0353;对比当下性能最好的 CDC 算法,

本文算法的 PSNR 与 SSIM 指标分别提升了 0.01 dB 与 0.0010;对比最新的 MSCAN 算法,本文算法的 PSNR 与 SSIM 指标分别提升了 0.17 dB 与 0.0025;对比轻量级算法 RFDN-L,本文算法虽然参数量是其 560%,但是本文算法的 PSNR 与 SSIM 指标比其提升了 0.17 dB 与 0.0047,具有更好的重建效果。除了基于插值的 Bicubic 算法和轻量级的 RFDN-L 算法外,其余各对比算法的参数量均在 10×10^6 上,不具备轻量化优势。对比 SRResNet、EDSR、ESRGAN、RCAN、CDC 和 MSCAN 算法,本文算法的参数量分别为以上各算法的 12.81%、8.11%、21.14%、22.58%、8.84% 和 34.5%。与上述各算法相比,本文算法平衡了参数量与重建性能,适合落地与部署。

图 8 所示为 Bicubic、RFDN-L、CDC 和本文算法在 2 倍放大系数下对 DRealSR 中 DSC_0988 的重建效果对比;图 9 所示为以上各算法在 3 倍放大系数下对 DRealSR 中 panasonic_50 的重建效果对比;图 10 所示为以上各算法在 4 倍放大系数下对 DRealSR 中 Canon_10 的重建效果对比。

从图 8—图 10 可以看出,Bicubic 的重建图像较为模糊;RFDN-L 虽然具有最少的参数量,但是其对于图像细节和纹理的重建效果不好,在纹理部分会产生涂抹感;采用 CDC 的方法针对图像的平坦区域、边缘和角点分别进行重建,具有良好的重建效果。本文算法采用深度可分离卷积和复用卷积操作大幅降低参数量,并使用残差特征融合操作有效减少了残差路径的长度,同时使用增强型注意力模块从通道和空间两个维度进行自适应调整,重建效果与 CDC 不相上下。但本文算法的参数量仅为 CDC 的 8.84%,因此具有轻量化优势。

表 4 不同算法在不同放大系数下的 PSNR/dB|SSIM 性能比较

Tab. 4 Comparison of PSNR/dB|SSIM performance of different algorithms under different amplification factors

Algorithm	Proposed Parameters		Scale×2		Scale×3		Scale×4	
	year	/×10 ⁶	RealSR	DRealSR	RealSR	DRealSR	RealSR	DRealSR
Bicubic	1981	—	31.67 0.8870	32.67 0.8770	28.63 0.8090	31.50 0.8350	27.24 0.7640	30.56 0.8200
SRResNet	2017	27.55	32.65 0.9070	33.56 0.9000	28.85 0.8320	31.16 0.8590	27.63 0.7850	31.63 0.8470
EDSR	2018	43.70	32.71 0.9060	34.24 0.9080	29.50 0.8410	32.93 0.8760	27.77 0.7920	32.03 0.8550
ESRGAN	2018	16.70	32.25 0.9000	33.89 0.9060	29.57 0.8410	32.39 0.8730	27.82 0.7940	31.92 0.8570
RCAN	2018	15.63	32.88 0.9080	34.34 0.9080	<u>29.68</u> 0.8410	33.03 0.8760	27.93 0.7950	31.85 0.8570
RFDN-L	2020	0.63	32.88 0.9113	34.26 0.9060	29.58 0.8395	32.81 0.8678	28.00 0.7954	31.94 0.8556
CDC	2020	39.92	<u>32.81</u> <u>0.9100</u>	34.45 <u>0.9100</u>	29.57 0.8410	<u>33.06</u> <u>0.8760</u>	<u>28.11</u> 0.8000	32.42 <u>0.8610</u>
MSCAN	2021	10.23	32.62 0.9083	34.33 0.9091	29.50 0.8408	32.98 0.8741	27.96 0.7969	32.08 0.8598
Ours	2022	3.53	32.77 0.9121	<u>34.44</u> 0.9121	29.69 0.8424	33.08 0.8768	28.13 <u>0.7992</u>	<u>32.38</u> 0.8612

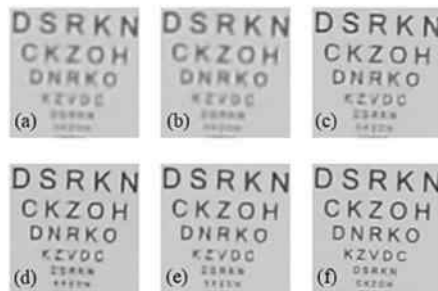


图 8 2 倍放大系数下 DSC_0988 重建效果对比:(a) LR; (b) 双三次插值; (c) RFDN-L; (d) CDC; (e) 本文算法; (f) HR

Fig. 8 Image “DSC_0988” at $2\times$ magnification reconstruction effect comparison:

(a) LR; (b) Bicubic; (c) RFDN-L; (d) CDC; (e) Ours; (f) HR

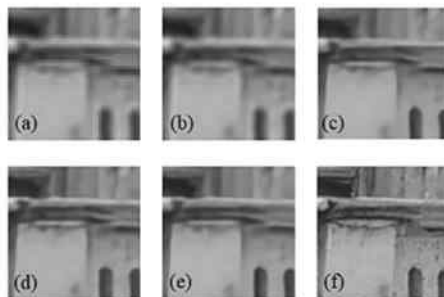


图 9 3 倍放大系数下 panasonic_50 重建效果对比:(a) LR; (b) 双三次插值; (c) RFDN-L; (d) CDC; (e) 本文算法; (f) HR

Fig. 9 Image “panasonic_50” at $3\times$ magnification reconstruction effect comparison:

(a) LR; (b) Bicubic; (c) RFDN-L; (d) CDC; (e) Ours; (f) HR

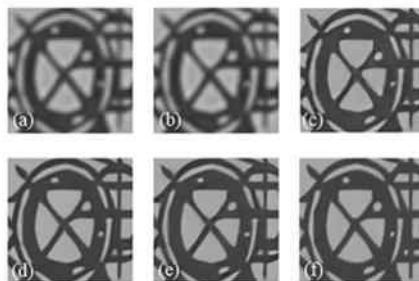
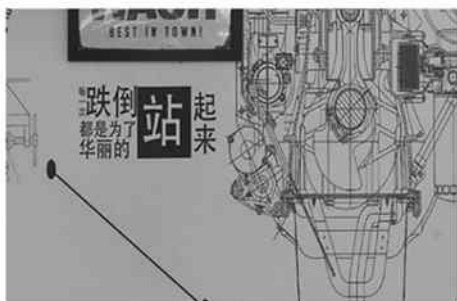


图 10 4 倍放大系数下 Canon_10 重建效果对比:(a) LR; (b) 双三次插值; (c) RFDN-L; (d) CDC; (e) 本文算法; (f) HR

Fig. 10 Image “Canon_10” at $4\times$ magnification reconstruction effect comparison:

(a) LR; (b) Bicubic; (c) RFDN-L; (d) CDC; (e) Ours; (f) HR

3.4.3 算法泛化性实验

面向真实图像的 SR 算法要求在不同设备拍摄的图像上均具有良好的重建性能,因此对本文算法进行泛化性测试十分必要。为了有效测试算法的泛化性能,本文使用基准数据集在不同退化函数下生成 HR-LR 图像对,然后针对两张个人拍摄的不同图像进行可视化分析。使用 Bicubic 和对比实验中训练的 RFDN-L 作为对比算法,测试本文算法在不同情形下的泛化性能。

首先,使用来自文献[22]中的高斯模糊(Gaussian blur, GB)、双三次插值退化(bicubic interpola-

tion, BI) 和 JPEG 压缩 3 种退化方法作为本实验的退化方法。使用 Urban100 作为测试集, SSIM 作为评价指标,在 3 倍放大系数下进行实验,实验结果如表 5 所示。其中,高斯模糊采用 5×5 的模糊核, JPEG 压缩的图像质量设置为 50,打勾表示使用了该退化方法。

从表 5 的实验结果可以看出,在仅使用双三次插值退化时,本文算法与 RFDN-L 的实验结果不佳,其原因在于本文算法与 RFDN-L 均采用真实图像数据集训练得到,在针对较为简单的退化方法时重建性能会比较差。而随着高斯模糊和 JPEG 压缩退化

方法的加入,重建难度逐渐增加,RFDN-L 和本文算法相较于 Bicubic 的重建性能也随之提升。在使用以上 3 种退化方法时,本文算法的 SSIM 结果相较于 RFDN-L 算法平均提升了 0.005 4,具有良好的泛化性能。

表 5 本文算法在不同退化方法下的 SSIM 结果
Tab. 5 SSIM results of proposed algorithm under different degradation methods

BI	GB	JPEG	Bicubic	RFDN-L	Ours
✓			0.720 8	0.679 9	0.672 3
✓	✓		0.663 9	0.675 2	0.690 4
✓	✓	✓	0.610 0	0.613 8	0.622 3

为了进一步测试本文算法的泛化性,在 4 倍放大系数下使用以上 3 种算法针对 iPhone 11 Pro 手机和 DJI Spark 无人机两种设备拍摄的图片进行 SR 重建,并对其进行可视化分析,如图 11 和图 12 所示。由于以上各图像无对应的 HR 图像,因此主要针对其重建的视觉效果进行可视化分析。

从图 11 和图 12 中可以看出,双三次插值算法重建图像较为模糊,且图像中包含噪点,主观重建效果较差。RFDN-L 算法重建图像视觉效果虽然优于双三次插值图像,但是对于图像边缘重建效果一般,如图 11(c)中的座椅边缘部分存在模糊;部分图像存在将噪点当成纹理执行重建的情况,从而导致伪影的出现,如图 12(c)中的斜拉桥主干上出现了部分伪影。本文算法重建图像边缘清晰且伪影较少,对于

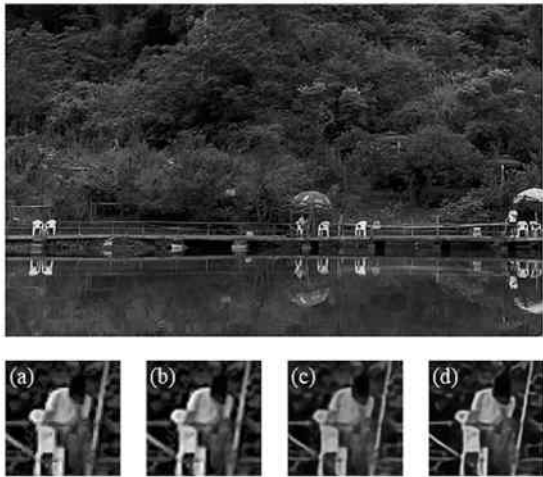


图 11 iPhone 11 Pro 拍摄图像重建效果对比:
(a) LR; (b) 双三次插值; (c) RFDN-L; (d) 本文算法
Fig. 11 Comparison of image reconstruction effects taken by iPhone 11 Pro: (a) LR; (b) Bicubic; (c) RFDN-L; (d) Ours

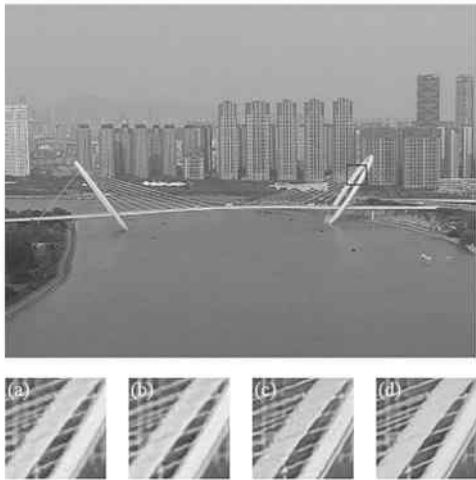


图 12 DJI Spark 拍摄图像重建效果对比:
(a) LR; (b) 双三次插值; (c) RFDN-L; (d) 本文算法
Fig. 12 Comparison of image reconstruction effects taken by DJI Spark: (a) LR; (b) Bicubic; (c) RFDN-L; (d) Ours

图像纹理部分的重建较为准确。虽然部分图像区域较为平滑,存在涂抹感,如图 12(d)中斜拉桥的绳索部分,其由于基于卷积神经网络的图像重建算法的局限性所致,即无法在保证重建准确性的同时兼顾主观感知质量^[23]。综上所述,本文算法总体视觉效果优于双三次插值与 RFDN-L 算法,在不同设备拍摄的图像中均有较好的重建效果。

3. 4. 4 算法效率实验

为了测试本文算法的计算量,在不同放大系数下使用不同尺寸的图像作为输入,实验结果如表 6 所示。

从表 6 的实验结果可以看出,本文算法由于采用了文献[17]的自适应上采样层,因此参数量不随放大系数的改变而改变。当输入图像尺寸相同时,算法的计算量随着放大系数的增加而逐渐增加。在放大系数不变时,算法的计算量随着输入图像的尺寸的增加而急剧增大,因为当输入图像的长和宽各增加一倍时,总像素数量变为原来的 4 倍,所以其计算量也约为原来的 4 倍。所以当输入尺寸增加时,计算量呈平方速度增长。

为了测试本文算法的重建时间,在 3 倍放大系数下将常用视频制式分辨率作为输入图像的尺寸,分别使用 Tesla M40 GPU 和 RTX2080 GPU 对其进行重建并统计时间,实验结果如表 7 所示。

从表 7 的实验结果可以看出,采用较新的 RTX 2080 GPU 的重建时间短于 Tesla M40 GPU。在使

用 RTX 2080 GPU 重建的前提下,当图像分辨率为 128×96 时,算法重建时间为 0.03 s,即每秒可以处理 33.33 张图像,在该分辨率下满足电影等每秒 24 帧视频的实时性要求;当分辨率为 176×144 时,其重建时间为 0.06 s,即每秒可以处理 16.67 张图像,在该分辨率下满足监控视频等每秒 15 帧视频的实时性要求。

表 6 不同输入下本文算法的参数量与计算量

Tab. 6 The amount of parameters and computation of proposed algorithm under different inputs

Input size	Scale	Parameters / $\times 10^6$	Flops / $\times 10^9$
128×128	2	3.53	51.43
256×256	2	3.53	205.74
512×512	2	3.53	822.94
128×128	3	3.53	54.60
256×256	3	3.53	218.40
512×512	3	3.53	873.61
128×128	4	3.53	63.87
256×256	4	3.53	255.49
512×512	4	3.53	1021.99

表 7 常用分辨率下本文算法的重建时间

Tab. 7 The reconstruction time of proposed algorithm under common resolution

Standard	Input size	Reconstructed time/s	
		Tesla M40	RTX2080
SQCIF	128×96	0.11	0.03
QCIF	176×144	0.22	0.06
SCIF	256×192	0.41	0.13
QVGA	320×240	0.63	0.20
CIF	352×288	0.83	0.27
HVGA	480×360	1.41	0.46
VGA	640×480	2.51	0.82
PAL	768×576	3.61	1.19

4 结 论

为了在保证算法重建性能的前提下解决真实图像的 SR 重建算法存在的参数量过大的问题,本文提出一种轻量级的真实图像 SR 重建算法。该算法基于深度可分离卷积和复用卷积提出了轻量级的多尺度特征提取模块,引入残差特征融合操作以减短残差路径并更好地利用局部残差特征,使用增强型注意力模块从空间和通道两个维度对特征图进行自适应调整,从而提升算法性能。最后通过自适应上采样模块获得 SR 图像。实验结果表明,本文算法在真实图像上的重建性能超越了如 SRResNet、ESR-

GAN、EDSR、CDC 等主流算法,而其参数量平均仅为以上主流算法的 8.82%,在保证重建性能的同时满足了轻量化特性。在复杂退化函数下仍然有良好的重建性能,具有较强的泛化性。未来将针对算法特征提取模块和损失函数进行优化,在保证重建性能的同时进一步提升重建图像的视觉效果。

参考文献:

- [1] WANG Z, CHEN J, HOI S C H. Deep learning for image super-resolution: a survey[J]. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2021, 43(10): 3365-3387.
- [2] CAO X M, LIU Z H, LIU J B. Superresolution reconstruction method of remote sensing image based on projection network[J]. Journal of Optoelectronics · Laser, 2020, 31(11): 1149-1156.
曹晓敏, 刘志红, 柳锦宝. 基于投影网络的遥感影像超分辨率重建方法[J]. 光电子·激光, 2020, 31(11): 1149-1156.
- [3] FAN P P, DONG X C, LI T, et al. Super-resolution reconstruction of depth map based on non-local means constraint[J]. Journal of Computer-Aided Design & Computer Graphics, 2020, 32(10): 1671-1678.
范佩佩, 董秀成, 李滔, 等. 基于非局部均值约束的深度图像超分辨率重建[J]. 计算机辅助设计与图形学学报, 2020, 32(10): 1671-1678.
- [4] LEDIG C, THEIS L, HUSZAR F, et al. Photo-realistic single image super-resolution using a generative adversarial network[C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, July 21-26, 2017, Honolulu, HI, USA. Los Alamitos: IEEE Computer Society Press, 2017: 4681-4690.
- [5] LIM B, SON S, KIM H, et al. Enhanced deep residual networks for single image super-resolution[C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops, July 21-26, 2017, Honolulu. Los Alamitos: IEEE Computer Society Press, 2017: 136-144.
- [6] LI J, FANG F, MEI K, et al. Multi-scale residual network for image super-resolution[C]//Proceedings of the European Conference on Computer Vision, September 8-14, 2018, Munich, Germany. Heidelberg: Springer, 2018: 517-532.
- [7] LIU J, ZHANG W, TANG Y, et al. Residual feature aggregation network for image super-resolution[C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, June 16-21, 2020, Seattle, USA. Los

- Alamitos: IEEE Computer Society Press, 2020: 2359-2368.
- [8] LIU J, TANG J, WU G. Residual feature distillation network for lightweight image super-resolution [C]//Proceedings of the European Conference on Computer Vision, August 23-28, 2020, Edinburgh, UK. Heidelberg: Springer, 2020: 41-55.
- [9] KONG X, ZHAO H, QIAO Y, et al. Classsr: a general framework to accelerate super-resolution networks by data characteristic [C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, June 19-25, 2021, Online. Los Alamitos: IEEE Computer Society Press, 2021: 12016-12025.
- [10] CHEN C, XIONG Z, TIAN X, et al. Camera lens super-resolution [C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, June 15-21, 2019, Long Beach, USA. Los Alamitos: IEEE Computer Society Press, 2019: 1652-1660.
- [11] WEI P, XIE Z, LU H, et al. Component divide-and-conquer for real-world image super-resolution [C]//Proceedings of the European Conference on Computer Vision, August 23-28, 2020, Edinburgh, UK. Heidelberg: Springer, 2020: 101-117.
- [12] SANDLER M, HOWARD A, ZHU M, et al. Mobilenetv2: inverted residuals and linear bottlenecks [C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, June 18-21, 2018, Salt Lake City, USA. Los Alamitos: IEEE Computer Society Press, 2018: 4510-4520.
- [13] LIU Y, WANG Y, LI N, et al. An attention-based approach for single image super-resolution [C]//Proceedings of the 24th International Conference on Pattern Recognition, August 20-24, 2018, Beijing, China. Los Alamitos: IEEE Computer Society Press, 2018: 2777-2784.
- [14] ZHANG Y, LI K, LI K, et al. Image super-resolution using very deep residual channel attention networks [C]//Proceedings of the European Conference on Computer Vision, September 8-14, 2018, Munich, Germany. Heidelberg: Springer, 2018: 286-301.
- [15] CHEN Z H, WU H B, PEI H D, et al. Image super-resolution reconstruction method based on self-attention deep network [J]. Laser & Optoelectronics Progress, 2021, 58(4): 0410013.
- 陈子涵, 吴浩博, 裴浩东, 等. 基于自注意力深度网络的图像超分辨率重建方法 [J]. 激光与光电子学进展, 2021, 58(4): 0410013.
- [16] WOO S, PARK J, LEE J Y, et al. CBAM: convolutional block attention module [C]//Proceedings of the European Conference on Computer Vision, September 8-14, 2018, Munich, Germany. Heidelberg: Springer, 2018: 3-19.
- [17] LV J, XU P C. Image super-resolution reconstruction algorithm based on adaptive up-sampling of multi-scale [J/OL]. Journal of Frontiers of Computer Science and Technology: 1-16 (2021-10-20) [2022-01-29]. <http://kns.cnki.net/kcms/detail/11.5602.TP.20211020.1013.002.html>.
- 吕佳, 许鹏程. 多尺度自适应上采样的图像超分辨率重建算法 [J/OL]. 计算机科学与探索: 1-16 (2021-10-20) [2022-01-29]. <http://kns.cnki.net/kcms/detail/11.5602.TP.20211020.1013.002.html>.
- [18] HAN B, NIU G, YU X, et al. Sigua: Forgetting may make learning with noisy labels more robust [C]//Proceedings of the 37th International Conference on Machine Learning, July 13-18, 2020 Vienna, Austria. Lille: Proceedings of Machine Learning Research, 2020: 4006-4016.
- [19] CAI J, ZENG H, YONG H, et al. Toward real-world single image super-resolution: a new benchmark and a new model [C]//Proceedings of the IEEE International Conference on Computer Vision, October 27-November 3, 2019, Seoul, South Korea. Los Alamitos: IEEE Computer Society Press, 2019: 3086-3095.
- [20] WANG X, YU K, WU S, et al. ESRGAN: enhanced super-resolution generative adversarial networks [C]//Proceedings of the European Conference on Computer Vision Workshops, September 8-14, 2018, Munich, Germany. Heidelberg: Springer, 2018: 63-79.
- [21] ZHANG Y, LI K, LI K, et al. Image super-resolution using very deep residual channel attention networks [C]//Proceedings of the European Conference on Computer Vision, September 8-14, 2018, Munich, Germany. Heidelberg: Springer, 2018: 286-301.
- [22] WANG X, XIE L, DONG C, et al. REAL-ESRGAN: training real-world blind super-resolution with pure synthetic data [C]//Proceedings of the IEEE International Conference on Computer Vision, October 10-17, 2021, Montreal, Canada. Los Alamitos: IEEE Computer Society Press, 2021: 1905-1914.
- [23] BLAU Y, MICHAELI T. The perception-distortion tradeoff [C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, June 18-21, 2018, Salt Lake City, USA. Los Alamitos: IEEE Computer Society Press, 2018: 6228-6237.

作者简介:

吕佳 (1978—), 女, 博士, 教授, 硕士生导师, 主要从事机器学习、数据挖掘及其在医学图像处理等方面的研究。