

# 基于并行残差卷积网络的图像超分辨重建

杨伟铭<sup>1</sup>, 张 钰<sup>2</sup>

(1. 军事科学院系统工程研究院后勤科学与技术研究所, 北京, 100800; 2. 陕西师范大学计算机科学学院, 西安, 710119)

**摘要** 针对 VDSR 模型卷积核单一和 DRRN 模型不能全局利用的问题, 提出了基于并行残差卷积神经网络的联合卷积图像超分辨重建模型。模型首先利用原始卷积层和扩张卷积层融合, 建立联合卷积层, 然后利用跳跃链接, 将多种抽象层次的特征进行融合, 最后完成整个超分辨网络的模型构建。提出的模型具有以下优点: ①扩张卷积神经网络与原始卷积神经网络融合, 在计算复杂度不变的情况下, 可以获取更多尺度的信息, 因此具有更强的表达能力; ②跳跃链接方式, 将抽象层次较低与较高抽象层次的信息融合, 获取更多的信息, 使得模型具有更强的学习能力。通过在多个数据集上进行实验, 模型在大多数任务中与 VDSR、DRRN 和 SRCNN 等先进模型相比, IFC 值取得了大于 0.1 的提升。

**关键词** 卷积神经网络; 图像超分辨率; 扩张神经网络; 跳跃链接; 深度学习

**DOI** 10.3969/j.issn.1009-3516.2019.04.013

**中图分类号** TP311 **文献标志码** A **文章编号** 1009-3516(2019)04-0084-06

## Image Super-Resolution in Combination with Convolution Neural Network

YANG Weiming<sup>1</sup>, ZHANG Yu<sup>2</sup>

(1. Institute of Logistics Science and Technology, Academy of System Engineering,  
Academy of Military Science of Chinese PLA, Beijing 100166, China;

2. School of Computer Science, Shaanxi Normal University, Xi'an 710119, China)

**Abstract:** Aimed at the problems that the VDSR model convolution kernel is single and the DRRN model fails to take advantages of global features, a combined convolution image super-resolution model is proposed based on parallel residual convolution neural networks. Firstly, the combined convolution neural layer is structured by the original convolution layer and dilated convolution layer, and the skip connection approach is employed to connect the different layers to take advantage of different level features, completing super-resolution network. There are two advantages of this model: ①Combination of dilated convolution neural layers and original convolution layers can capture multi-scale features without computation-consuming. Based on this approach, the network can get more presentation capacity. ②Skip connection approach fuse low-level information and high-level information. From this approach, different level features can be learned. This means that stronger learning ability can be obtained. Based on the experiment results on multiple data sets, more than 0.1 IFC improvement is achieved, compared with the state-of-the-art models VDSR, DRRN, SRCNN in most tasks.

**Key words:** convolution neural network; image super-resolution; dilated convolutional network; skip connection; deep learning

收稿日期: 2018-09-17

基金项目: 国家自然科学基金(41671409)

作者简介: 杨伟铭(1982—), 男, 江西鹰潭人, 工程师, 主要从事地理信息系统、物联网技术及应用研究。E-mail: damoyc@163.com

**引用格式:** 杨伟铭, 张钰. 基于并行残差卷积网络的图像超分辨重建[J]. 空军工程大学学报(自然科学版), 2019, 20(4): 84-89. YANG Weiming, ZHANG Yu. Image Super-Resolution in Combination with Convolution Neural Network[J]. Journal of Air Force Engineering University (Natural Science Edition), 2019, 20(4): 84-89.

图像分辨率的高低取决于成像的硬件系统和成像的方法,其中图像的超分辨(Image Super-Resolution, SR)是软件方面的一项关键技术,此项技术克服图像传感器等成像硬件的固有限制,提高了图像像素密度。其方法主要有2个途径:①仅利用单张低分辨率(Low-Resolution, LR)图像重构成高分辨率(High-Resolution, HR)图像,即单帧图像超分辨率(Single Image Super-Resolution, SISR)<sup>[1-2]</sup>;②利用给定的同一场景,多张图像序列重构场景高分辨率图像,即多帧图像超分辨率(Multiple Image Super-Resolution, MISR)<sup>[3]</sup>。

本文主要关注单幅图像重建,即 SISR 问题。一般来说, SISR 问题的目标是利用高分辨率退化的低分辨率图像中恢复高分辨率的图像,可以将其表述为:

$$y_{hr} = \Phi(x_{lr}) \quad (1)$$

式中:  $x_{lr}$  和  $y_{hr}$  分别表示低分辨率输入图像和高分辨率输出图像;  $\Phi$  是重构过程的表示函数。对于单张图像的超分辨重建问题的研究已取得了很大的成功<sup>[4-5]</sup>。基重建的图像超分辨算法利用低分辨率图像先验知识进行图像重建,此类方法简单且计算量低,但无法处理复杂结构的图像。基于学习的图像超分辨重建,此类方法假设低分辨率的图像已经拥有预测其所对应的高分辨率部分的信息,通过对大量样本的学习建立低分辨率图像和高分辨率图像之间的映射关系。典型的基于学习的方法如 Chang 等人提出的邻域嵌入(Neighbor Embedding)算法<sup>[6]</sup>, Yang 等人提出的稀疏编码(Sparse Coding)算法<sup>[7]</sup>,以及基于卷积神经网络(Convolutional Neural Network, CNN)的算法<sup>[8-9]</sup>等。

目前,基于深度学习的超分辨重建(SISR)因其良好的学习能力,表达能力成为图像超分辨领域的主流研究方向。Dong 等首次将深度学习方法(Deep Learning, DL)用于图像超分辨重建,证明了卷积网络通过学习可以有效的建立起低分辨率图像(LR)与高分辨率图像(HR)之间的映射关系<sup>[10]</sup>。Kim 等<sup>[11]</sup>进一步改进深度学习模型,利用增加卷积神经网络深度的方法扩大卷积神经网络的感受野(Receptive Field),提出了一个包含20层的网络结构,即深度超分辨率卷积网络(Very Deep Super-Resolution Network, VDSR)。与此前的超分辨率卷积网络(Super-Resolution Convolution Neural Network, SRCNN)相比, VDSR 具有更大的感知范围,因此也获得了更好的表现效果。最近,在 VDSR 的基础上引入了残差神经网络(Residual Neural Network),此方法通过残差神经网络,引入共享权

重机制,减少了模型的参数,提高了模型的性能。(Deep Recursive Residual Network, DRRN)模型通过多路径方法充分利用不同卷积信息,提高了重建模型的性能<sup>[12]</sup>。

本文基于深度学习的方法为研究对象,在深入研究 VDSR 方法和 DRRN 方法的基础上,考虑到 VDSR 模型仅仅考虑单一卷积核;而 DRRN 模型每一层输入仅仅来自上一层信息,没有全局考虑其他层信息输入。本文提出构建基于输入残差的深度神经网络来解决超分辨重建任务。本文所提出的联合卷积残差网络主要有以下优势:

1)联合卷积解决了 VDSR 中单一卷积核的单调性,通过不同尺度的卷积核,能够提取出不同等级的特征,从而达到更充分地特征提取的目的。

2)充分考虑到深度学习网络会随着网络模型的加深,输入信息会在不断的卷积过程,低频信息会被丢弃,导致深层网络只输出图像中的高频信息。本文引入了首尾相联结的结构来保证在图像残差重建时充分地考虑了低层次特征信息。

## 1 相关工作

国内外已经有学者在超分辨率重构领域做了很多工作。李民等<sup>[13]</sup>基于稀疏字典编码在这一方面做了尝试;而干宗良等<sup>[14]</sup>也针对该问题采用了相似的方法。黄婧等<sup>[15]</sup>则采用全局运动模型配准策略应对运动中图像的模糊问题。国内研究更多的停留于传统算法,而国外同行们则更加倾向于采用智能的卷积神经网络解决这一问题。

卷积神经网络是一种处理二维输入数据设计的多层人工神经网络<sup>[10]</sup>。卷积神经网络的图像超分辨率算法利用卷积神经网络在图像处理方面的优势:①多维数据可以直接输入网络进行训练;②局部权值共享特点可以有效地减少训练参数数量;③强大的特征学习能力和建模能力。本节主要讨论3种基于卷积神经网络的图像超分辨算法: SRCNN, VDSR, DRRN。

SRCNN 是第1个在超分辨率重建领域采用卷积神经网络的模型。它由3个部分组成:特征提取和特征表示层,非线性映射层和重构层。第1层作用是从 LR 输入中提取高维向量;第2层的目的是将原始向量映射到另一个高维向量;最后的重建层作用是生成 HR 图像。在这些神经网络层中,分别使用的卷积核大小分别为  $9 \times 9$ ,  $1 \times 1$  和  $5 \times 5$ 。虽然 SRCNN 在超分辨率重建领域取得了巨大成功,但它仍然存在着架构简单导致的感受野有限,模型

还有巨大的改进空间。

为了解决感受野太小带来的问题, Kim 等人基于模型越深, 感受野越大的特点, 基于残差网络改进了 SRCNN 模型, 提出了 VDSR 模型。基于增加网络深度可以显著提升 SISR 的性能, 他们利用残余结构使网更深(20 层)。与 SRCNN 相比, VDSR 具有更大的接受范围( $41 \times 41$ )。作者的主要贡献是: ①利用残差层调节梯度下降率, 残差的提出有效的解决了梯度消失问题; ②残差层的引入可以有效的简化训练, 降低计算复杂度。基于 VDSR 的实验表明, 相等于 SRCNN, VDSR 显示出更好的性能。

DRRN<sup>[12]</sup> 的结构类似于 ResNet<sup>[9]</sup>。原始卷积层和多个残差块(RB)叠加被用于提取图像的特征。值得注意的是, DRRN 模型与 ResNet 模型有相同的缺点: 特征映射是由第 1 个卷积层生成的。“预激活”结构在 DRRN 模型中被使用, 模型主要通过残差进行学习。DRRN 利用多路径模式来促进学习, 并防止过度拟合。与 VDSR 相比, DRRN 可以用更少的参数获得更好的性能注意。

## 2 联合卷积神经网络

### 2.1 联合卷积

联合卷积中的“联合”一词在这里有 2 层含义。首先, 在本文提出的方法中, 用普通卷积和扩张卷积联合提取特征及进行特征的非线性映射操作。

如图 1 所示, 利用普通卷积神经网络与扩张卷积神经网络的联合, 获取相邻像素之间的信息。图中蓝色色块代表由普通卷积提取出的特征, 绿色色块代表的是由扩张卷积提取出的特征。把这 2 种特征沿着特征通道方向联合拼接, 并在后续计算时把它们当作一个整体。在计算时使用不同的 2 种卷积操作, 既得到了像素相邻之间的上下文信息, 同时也获得了彼此不相邻但相距较近的像素间的上下文信息。这样既加强了对输入的信息提取能力, 又提高了特征的利用率, 具体效果体现在下文运用该思想构建的网络组织及超分辨率重建实验中。

使用联合卷积方法的超分辨率重构网络在标准数据集中有更好的表现。

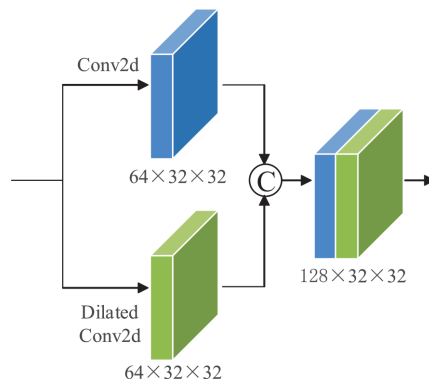


图 1 联合卷积

### 2.2 多层网络之间“联合”

此外, 此处的“联合”还有另一层意思。通过对 VDSR 模型的观察发现该模型仍存在一定的不足。简单级联的深层卷积网络使得浅层卷积层所输出的特征层在最后计算输入残差时无法被充分地考虑进来。因此, 本文在所提出的模型中也加入了跳跃连接。

### 2.3 网络结构

本文所提出的网络的主要结构如图 2 所示。通过跳跃链接, 提高网络对不同抽象程度特征的利用率。其中浅蓝色方块表示由图 1 中的联合卷积步骤得到的特征层。为了加速网络训练, 首先通过 Bicubic 插值把输入放缩到目标大小。将其作为网络的输入。整个网络通过学习来修正 Bicubic 插值得到的结果的各个像素点, 即网络学习的是 Bicubic 插值输入的残差。在此, 广义地把它分为 3 个模块, 分别是特征提取模块, 特征映射模块及残差重建模块。特征提取模块是一个联合卷积结构, 该结构能够极大地提取输出的语义特征和上下文信息。特征映射模块由一个带有跳跃连接的级联联合卷积结构组成, 与普通的级联结构相比, 它能够充分地利用各联合卷积结构的输出。最后, 残差重建模块则是一个普通的卷积层表示。值得一提的是, 在每个卷积层之前均有批正则层和 ReLU 依次相连。

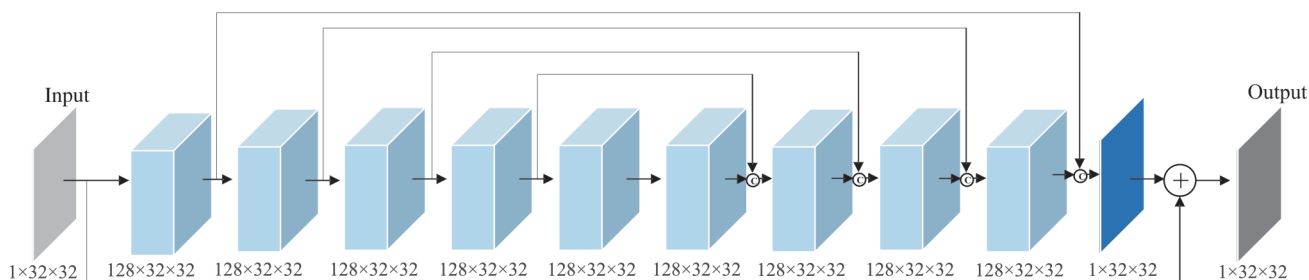


图 2 单张图像超分辨重建框架图

### 3 实验参数

在本节中将详细说明实验参数及实现细节。与此同时,将本文的结果和已有的方法进行充分地实验对比。

#### 3.1 网络结构

与文献[11~12],[16~17]保持一致,本文也采用了相同的 291 张图片作为训练集。其中的 91 张图片来自文献[18],其余图片来自文献[19]。文献[11]中指出超分辨率重建的最佳方法先把 RGB 转为 YCbCr 格式然后只对 Luminance 层进行超分辨率重建,本文沿用了这个方法。首先把所有图片都转为 YCbCr 的格式,然后提取  $L$  层。然后把图片分别旋转  $90^\circ, 180^\circ, 270^\circ$  及对称变换做数据增强操作。然后把图片切成  $32 \times 32$  的像素块,作为训练集。此外,为了增强当单个模型对各种缩放因子的鲁棒性,本文把不同缩放因子的训练集混合在一起训练。

Set5<sup>[20]</sup>和 Set14<sup>[18]</sup>经常作为超分辨率重建的标准测试集。另外,BSD100和 Urban100<sup>[21]</sup>也被用来作为额外的实验对比测试集。文中将用这 4 个数据集作为实验对比的数据集。

#### 3.2 模型训练

在训练较深的网络结构时,容易出现梯度爆炸和梯度弥散的问题。除了上文中提到的批正则化之

外,本文同样使用了梯度截断的方法避免这个问题。同时,梯度截断允许以一个较大的学习率(0.1)对模型进行训练。具体的说,每次梯度回传时,梯度都被截取在区间 $[-s/r, s/r]$ 之内,其中  $s$  是预设常量(0.01), $r$  表示当前学习率。

#### 3.3 损失函数

本文中用到的损失函数是均方差损失函数,与 VDSR 方法<sup>[11]</sup>及 DRRN 方法<sup>[12]</sup>一致,形式如下:

$$\lambda_{\text{MSE}} = \frac{1}{n} \sum_{i=1}^N \|R(y_i) - x_i\|_2^2 \quad (2)$$

式中: $y_i$  低分辨率图像的输入值; $x_i$  为高分辨率图像的目标值; $R(y)$  为深度学习层的非线性映射, $x_i$  与  $y_i$  之间的关系满足  $R(y) = x$ 。

### 4 实验对比

实验中,在一个配有 Nvidia Titan XP 显卡的服务器上使用 caffe 训练模型。每次训练 32 个样本。初始学习率为 0.1,每迭代 10 次学习率降为当前的 1/10。同时本文选择随机梯度下降(Stochastic Gradient Descent,SGD)作为该模型的优化器。整个模型的训练时间将近 2 d。

本文用训练之后的模型在上文提到的 4 个训练集上与前人的超分辨率经典模型进行详尽地对比,具体的统计数据见表 1。

表 1 PSNR/SSIM 分数

| Dataset  | Scale      | Bicubic       | SRCNN         | VDSR          | DRCN          | DRRNBIU9      | Present study |
|----------|------------|---------------|---------------|---------------|---------------|---------------|---------------|
| Set5     | $\times 2$ | 33.66/0.929 9 | 36.66/0.954 2 | 37.53/0.958 7 | 37.63/0.958 8 | 37.66/0.958 9 | 37.67/0.960 0 |
|          | $\times 3$ | 30.39/0.868 2 | 32.75/0.909 0 | 33.66/0.921 3 | 33.82/0.922 6 | 33.93/0.923 4 | 33.93/0.924 5 |
|          | $\times 4$ | 28.42/0.810 4 | 30.48/0.862 8 | 31.35/0.883 8 | 31.53/0.885 4 | 31.58/0.886 4 | 31.56/0.890 0 |
| Set14    | $\times 2$ | 30.24/0.868 8 | 32.45/0.906 7 | 33.03/0.912 4 | 33.04/0.911 8 | 33.19/0.913 3 | 33.22/0.913 9 |
|          | $\times 3$ | 27.55/0.774 2 | 29.30/0.821 5 | 29.77/0.831 4 | 29.76/0.831 1 | 29.94/0.833 9 | 29.92/0.834 2 |
|          | $\times 4$ | 26.00/0.702 7 | 27.50/0.751 3 | 28.01/0.767 4 | 28.02/0.767 0 | 28.18/0.770 1 | 28.15/0.770 9 |
| BSD100   | $\times 2$ | 29.56/0.843 1 | 31.36/0.887 9 | 31.90/0.896 0 | 31.85/0.894 2 | 32.01/0.896 9 | 32.03/0.897 2 |
|          | $\times 3$ | 27.21/0.738 5 | 28.41/0.786 3 | 28.82/0.797 6 | 28.80/0.796 3 | 28.91/0.799 2 | 28.93/0.800 7 |
|          | $\times 4$ | 25.96/0.667 5 | 26.90/0.710 1 | 27.29/0.725 1 | 27.23/0.723 3 | 27.35/0.726 2 | 27.37/0.727 5 |
| Urban100 | $\times 2$ | 26.88/0.840 3 | 29.50/0.894 6 | 30.76/0.914 0 | 30.75/0.913 3 | 31.02/0.916 4 | 31.10/0.917 8 |
|          | $\times 3$ | 24.46/0.734 9 | 26.24/0.798 9 | 27.14/0.827 9 | 27.15/0.827 6 | 27.38/0.833 1 | 27.40/0.834 2 |
|          | $\times 4$ | 23.14/0.657 7 | 24.52/0.7221  | 25.18/0.752 4 | 25.14/0.751 0 | 25.35/0.757 6 | 25.42/0.761 4 |

结果表明:在几乎所有类型的超分辨率重构任务中,本文模型均获得了较高的 PSNR/SSIM 分数,仅在  $\times 3$  与  $\times 4$  任务中相较 DRRNB1U9 稍有落后。

除了常见的评价指标峰值信噪比(Peak Signal to Noise Ratio,PSNR)和结构相似性(Structural Similarity Index,SSIM)之外,本文沿用了文献[12]中用到的额外的评价指标:信息保真度准则

(Information Fidelity Criteria, IFC)。

从上表可以看出,文中所提出的方法在 4 个标准测试集上的评价分数均达到了相当有竞争力的水平。本文方法的 PSNR 评分在大部分情况下都取得了最高的分数。值得一提的是,在相同网络深度等条件下,本文方法在 SSIM 的评分上均优于 2017 年国际计算机视觉与模式识别会议(Conference on



Computer Vision and Pattern Recognition, CVPR) 中提出的方法 DRRN。

表 2 详细地记录了各模型间的 IFC 得分情况。通过该表可以看出,本文提出的方法在 Set5, Set14 和 Urban100 上都大大超过了现有方法的最佳评分。具体地,本文的方法在 Set5 上 2 倍、3 倍和 4 倍超分辨率重建的 IFC 评分分别超过了 DRRN 0.228, 0.081 和 0.093。通过观察可知,该方法在 2

倍超分辨率重建下的 IFC 评分与之前的方法相比均有明显提升。该方法在 Set14 和 Urban100 上的 2 倍超分辨率重建 IFC 平均分别超过了 DRRN 0.29 和 0.323。结果表明:文中所提出的联合卷积残差网络结构能够在相同的网络深度下更好地提取并利用输入及网络内部的特征与上下文关系,进一步优化了超分辨率重建的结果。

表 2 IFC 分数

| Dataset  | Scale | Bicubic | SRCNN | PSyCo | VDSR  | DRRNBIU9 | Present study |
|----------|-------|---------|-------|-------|-------|----------|---------------|
| Set5     | ×2    | 6.083   | 8.036 | 8.642 | 8.569 | 8.583    | 8.811         |
|          | ×3    | 3.580   | 4.658 | 5.083 | 5.221 | 5.241    | 5.322         |
|          | ×4    | 2.329   | 2.991 | 3.379 | 3.547 | 3.581    | 3.674         |
| Set14    | ×2    | 6.105   | 7.784 | 8.280 | 8.178 | 8.181    | 8.471         |
|          | ×3    | 3.473   | 4.338 | 4.660 | 4.730 | 4.732    | 4.820         |
|          | ×4    | 2.237   | 2.751 | 3.055 | 3.133 | 3.147    | 3.235         |
| Urban100 | ×2    | 6.245   | 7.989 | 8.589 | 8.645 | 8.653    | 8.976         |
|          | ×3    | 3.620   | 4.584 | 5.031 | 5.194 | 5.259    | 5.368         |
|          | ×4    | 2.361   | 2.963 | 3.351 | 3.496 | 3.536    | 3.649         |

图 3 列出了 2 组测试集的对比图,其中第 1 组图片来自 Set14,第 2 组图片来自 BSD100。每组图像的第 1 张是使用 Bicubic 插值的图像,第 2~4 张使用已有模型,第 5 张是使用本文模型重建的 2 倍超分辨率图像,最后一张为缩小前原图像(Ground Thruth,GT)。其中每副图片下方标有 A、B 的图片分别为上方图像中 A、B 的放大图。通过对比发现,基于深度学习的方法所重建出的结果从视觉上看均

好于普通的插值方法。对比发现,基于深度学习的 4 种方法所重建出的结果从视觉上看均好于普通的插值方法。从峰值信噪比、结构相似性以及信息保真度准则得分上来看本文所提出的方法均优于现有方法的重建水平。该结果一定程度上证实了联合卷积方法在图像超分辨率重构过程中对输入信息的提取能力及特征利用率的提升。



图 3 不同测试集图像的 PSNR/SSIM/IFC 对比图

## 5 结语

本文利用普通卷积神经网络与扩张卷积神经网络的联合,提出了联合卷积层结构;在此结构的基础上,采用跳跃链接方式,将不同的网络层进行联合。该方法在层内部利用扩张卷积方法,在多个尺度上充分获取本层输入信息,同时在层之间利用跳跃链接,将不同抽象程度的信息充分融合表达。本模型在不提高计算复杂度的情况下,较好地改进了 VS-DR 和 DRRN 模型。实验结果表明本文所提出的方法在 PSNR、SSIM、IFC 这 3 项评价指标和重建图像的视觉效果方面皆优于 VSDR 和 DRRN 模型。

### 参考文献(References):

- [1] ZHANG K, TAO D, GAO X, et al. Coarse-to-Fine Learning for Single-Image Super-Resolution[J]. IEEE Transition on Neural Networks and Learning Systems, 2017, 28(5):1109-1122.
- [2] ZENG K, YU J, WANG R, et al. Coupled Deep Autoencoder for Single Image Super-Resolution[J]. IEEE Transition on Cybernetics, 2017, 47(1):27-37.
- [3] CHEN C, LIANG H, ZHAO S, et al. A Novel Multi-Image Super-Resolution Reconstruction Method Using Anisotropic Fractional Order Adaptive Norm[J]. The Visual Computer, 2015, 31(9):1217-1231.
- [4] FRANCESC A. A Nonlinear Algorithm for Monotone Piecewise Bicubic Interpolation[J]. Applied Mathematics and Computation, 2016, 272: 100-113.
- [5] YANG X, ZHANG Y, ZHOU D, et al. An Improved Iterative Back Projection Algorithm Based on Ringing Artifacts Suppression[J]. Neurocomputing, 2015, 162:171-179.
- [6] FANG B, HUANG Z, LI Y, et al.  $v$ -Support Vector Machine Based on Discriminant Sparse Neighborhood Preserving Embedding[J]. Pattern Analysis and Applications, 2017, 20(4):1077-1089.
- [7] YANG J, WANG Z, LIN Z, et al. Coupled Dictionary Training for Image Super-Resolution[J]. IEEE Transactions on Image Processing, 2012, 21(8):3467-3478.
- [8] DONG C, LOY CC, HE K, et al. Image Super-Resolution Using Deep Convolutional Networks[J]. Proceedings of the 2016 IEEE Transition on Pattern Analysis and Machine Intelligence, 38(2):295-307.
- [9] HE K, ZHANG X, REN S, et al. Deep Residual Learning for Image Recognition[C]//Proceedings of the 2016 IEEE Conference on Computer Vision and Pattern Recognition. Washington, DC: IEEE Computer Society, 2016:770-778.
- [10] HE K, ZHANG X, REN S, et al. Identity Mappings in Deep Residual Networks[C]//Proceedings of the 2016 European Conference on Computer Vision. Berlin: Springer, 2016:630-645.
- [11] JIWON K, LEE J K, LEE K M. Accurate Image Super-Resolution Using Very Deep Convolutional Networks[C]//Proceedings of the 2016 IEEE Conference on Computer Vision and Pattern Recognition. Washington, DC: IEEE Computer Society, 2016:1646-1654.
- [12] TAI Y, YANG J, LIU X. Image Super-Resolution via Deep Recursive Residual Network[C]//Proceedings of the 2017 IEEE Conference on Computer Vision and Pattern Recognition. Piscataway, NJ: IEEE, 2017: 2790-2798.
- [13] 李民, 程建, 乐翔, 等. 稀疏字典编码的超分辨率重建[J]. 软件学报, 2012, 34(5):1315-1324.  
LI M, CHENG J, LE X, et al. Super-Resolution Based on Sparse Dictionary Coding[J]. Journal of Software, 2012, 34(5):1315-1324. (in Chinese)
- [14] 干宗良, 梁秀聚. 自适应稀疏约束图像超分辨率重建方法[J]. 电视技术, 2012, 36(14):19-23.  
GAN Z L, LIANG X J. Adaptive Sparse Constraint Image Super-Resolution Reconstruction Method[J]. Video Engineering, 2012, 36(14):19-23. (in Chinese)
- [15] 黄婧, 李金宗. 基于全局运动模型配准的图像超分辨率重建[J]. 中国图象图形学报, 2007, 12(8):1354-1358.  
HUANG J, LI J Z. The Image Super-Resolution Reconstruction Based on the Global Motion Registration[J]. Journal of Image and Graphics, 2007, 12(8): 1354-1358. (in Chinese)
- [16] SAMUEL S, LEISTNER C, BISCHOF H. Fast and Accurate Image Upscaling with Super-Resolution Forests[C]//Proceedings of the 2015 IEEE Conference on Computer Vision and Pattern Recognition. Piscataway, NJ: IEEE, 2015:3761-3799.
- [17] YANG J, WRIGHT J, HUANG T S, et al. Image Super-Resolution Via Sparse Representation[J]. IEEE Transactions on Image Processing, 2010, 19(11): 2861-2873.
- [18] MARTIN D, FOWLKES C, TAL D, et al. A Database of Human Segmented Natural Images and Its Application to Evaluating Segmentation Algorithms and Measuring Ecological Statistics[C]//Proceedings of the 2001 IEEE International Conference. Piscataway, NJ: IEEE, 2001: 416-423.
- [19] BEVILACQUA M, ROUMY A, GUILLEMOT C, et al. Low-Complexity Single Image Super-Resolution Based on Nonnegative Neighbor Embedding[C]//Proceedings of the 2012 British Machine Vision Conference. Surrey, UK: BMVC, 2012:1-10.
- [20] ZEYDE R, ELAD M, PROTTER M. On Single Image Scale-up Using Sparse-Representations[C]//Proceedings of the 2010 International Conference on Curves and Surfaces. Berlin: Springer, 2010:711-730.
- [21] HUANG J B, SINGH A, AHUJA N. Single Image Super-Resolution from Transformed Self-Exemplars[C]//Proceedings of the 2015 IEEE Conference on Computer Vision and Pattern Recognition. Piscataway, NJ: IEEE, 2015:5197-5206.

(编辑:徐楠楠)