

用于图像超分辨率重建的双通道残差网络

左龙¹, 张鹏², 荆树旭¹, 赵一¹, 李凡³

(1. 长安大学信息工程学院, 710054, 西安; 2. 西安电子科技大学通信工程学院, 710071, 西安)

3. 西安交通大学信息与通信工程学院, 710049, 西安)

摘要: 针对现有基于深度学习的自然图像超分辨率算法在图像高频细节重建方面的不足, 提出了一种更注重图像高频细节重建的双通道残差网络。使用带有通道注意力机制的残差结构作为网络的主通道; 为了在重建过程中更好地保留原始图像的几何结构和边缘信息, 使用自适应结构化卷积设计了网络的辅助通道, 以此构建的双通道残差网络在学习过程中会有更强的高频信息捕获能力; 为了使重建图像效果更加符合人眼的主观视觉感受, 结合使用 L_1 损失函数和多尺度结构相似度损失函数来训练网络, 使网络在训练过程中能够较好地保留图像的视觉效果。实验结果表明: 在主通道外并构基于结构化卷积的辅助通道可以使重建图像的峰值信噪比提高 2 dB; 结合使用 L_1 损失函数和多尺度结构相似度损失函数可以使重建图像的峰值信噪比提高 3 dB、结构相似性提高 0.5; 与同类网络客观定量相比, 所提网络在两个公开数据集上取得的效果更优。

关键词: 超分辨率重建; 深度学习; 通道注意力; 残差网络; 自适应结构化卷积

中图分类号: TP391.41 **文献标志码:** A

DOI: 10.7652/xjtuxb202201018 **文章编号:** 0253-987X(2022)01-0158-07



OSID 码

Dual-Channel Residual Network for Image Super-Resolution Reconstruction

ZUO Long¹, ZHANG Peng², JING Shuxu¹, ZHAO Yi¹, LI Fan³

(1. School of Information Engineering, Chang'an University, Xi'an 710054, China;

2. School of Telecommunications Engineering, Xidian University, Xi'an 710071, China;

3. School of Information and Communications Engineering, Xi'an Jiaotong University, Xi'an 710049, China)

Abstract: Aiming at the shortcomings of the existing deep learning based image super-resolution algorithms in image high-frequency detail reconstruction, a dual-channel residual network emphasizing image high-frequency detail reconstruction is proposed. Residual structure with channel attention mechanism is leveraged as the main channel of the network. Targeting on retaining more delicate geometric structure and edge information of the original image during the reconstruction, the auxiliary channel of the network is designed by an adaptive structured convolution hence the evolved dual-channel residual network has a stronger ability to capture high-frequency information during the learning process. To make the reconstructed image well coincide with the subjective visual experience of human eye, L_1 loss function is combined with the multi-scale structural similarity loss function to train the network, so that the network completely retains the visual effect of the image during the training process. Experiments on the benchmark database show that combining the auxiliary channel based on adaptive structured convolution outside the main channel can heighten the peak signal-to-noise ratio of the

reconstructed image by 2 dB. The simultaneous operation of L_1 loss function and the multi-scale structural similarity loss function can heighten the peak signal-to-noise ratio of the reconstructed image by 3 dB and the structural similarity by 0.05. The objective and quantitative comparison with the competing networks exhibits the proposed network's outstanding effectiveness on two public data sets.

Keywords: super-resolution reconstruction; deep learning; channel attention; residual networks; adaptive structured convolution

图像超分辨率(SR)是从一张或多张低分辨率(LR)图像中恢复出高分辨率(HR)图像。近些年,随着深度学习技术的快速发展,基于深度学习的SR技术也随着取得了迅猛发展,并在行人识别、安全监控、医学成像等众多领域得到了广泛的应用^[1-3]。然而,图像SR是一个不适定的逆问题,每个LR图像都可能对应多个重建方案。近些年,许多研究人员使用深度学习技术解决这种逆问题,取得了有效的成果^[4-7]。最近的很多研究成果也表明,网络的深度对于提升模型的重建性能至关重要,更深的网络能够从输入的低分辨率图像中能学习到更多信息。然而,简单地堆叠网络深度并不会对网络带来太大的提升,甚至会使训练变得更加困难,使网络更难以收敛。

基于此,一些学者将残差结构和深度网络相结合,推出了一些效果十分出众的模型,如VDSR^[8]、RCAN^[9]等。这些模型通过加入网络之间的跳跃连接,让图像的低频信息可以绕过网络,使网络的主通道更加专注于学习图像的高频细节,提高了重建的效果。然而,这些算法在图像边缘和纹理重建方面仍然不尽如人意,经常会出现几何结构变形和扭曲、边缘模糊、细节纹理缺失等现象,导致重建图像的视觉效果及客观指标评价均不佳。

RCAN网络中提出的残差中的残差(RIR)结构试图通过结合残差网络的粗粒度和细粒度信息使网络专注于高频信息的学习,但是RCAN网络简单地堆叠RIR结构,导致网络训练过慢、参数量过大。另外,VDSR、RCAN等使用残差结构的网络在训练时使用不考虑人类视觉的 L_1 损失函数,不利于网络对于图像边缘、纹理等高频信息的学习。

不同于传统超分网络使用单一通道学习图像特征,本文设计了一种用于图像超分辨率重建的双通道残差网络。使用包含跳跃连接和通道注意力模块的残差组作为网络的主通道,并在每一个残差组之外并构了一个包含自适应结构化卷积的辅助通道,为主通道提供了自适应的感受野,增强了网络学习图像结构、边缘等高频信息的能力,加快了网络的收

敛速度,极大地提高了图像的重建质量。为了加强网络学习高频信息的能力,在损失函数中引入了能保留图像边缘和细节的多尺度结构相似度损失函数。实验结果表明,在重建效果类似的情况下,本文网络的层数远低于传统残差网络的。

1 相关研究

随着深度学习的不断发展,许多学者将深度学习技术应用于图像SR领域,提出了许多优秀的算法,较传统算法有了显著的效果提升。Dong等在2014年提出了一种基于卷积神经网络的图像超分辨率重建算法(SRCNN)^[10]。这是一个3层的卷积神经网络(CNN),虽然网络较为简单、感受野小、收敛速度慢,但是相较于传统的图像超分辨算法,其在重建效果与速度方面均有质的提升。为了提升网络的图像SR重建效果,必然使用更宽或者更深的神经网络,但是简单地提升网络深度必然会带来网络难以训练与收敛、梯度爆炸等一系列问题,因此许多研究人员在卷积网络的设计中加入了残差结构。Kim等在2016年提出的20层VDSR网络首次将残差结构与卷积网络相结合应用于图像SR任务中^[8]。由于加深了网络的深度,VDSR算法的学习能力显著加强,网络重建图像的效果也更好。Tai等在2017年提出了结合局部残差连接和残差单元的递归学习算法和52层的深度递归残差网络(DRRN)^[11]。为进一步加深网络深度,Tai等还将记忆模块引入图像SR任务,提出了深度达到80层的长期记忆网络(MemNet)^[12]。Li等在2018年首次将图像多尺度特征应用于残差结构中,提出了多尺度残差网络(MSRN)^[13]。Ahn等在2018年将局部级联和全局级联的方式引入SR网络设计,提出了级联残差网络(CARN),该网络使用了多级表示和快捷连接,网络更加轻量化^[14]。He等在2019年借鉴微分动力学系统设计了一种轻量型的图像超分辨率重建网络(OISR),使用改进的前向欧拉公式设计网络的残差块,使网络更容易收敛^[15]。Zhang等

在 2018 年首次将通道注意力机制应用于图像超分辨率问题中,设计出残差通道注意力网络 RCAN,通过自适应的学习通道之间的相互依赖性使得网络专注学习重要的通道特征,从而提高网络的性能^[9]。RCAN 网络中设计了 RIR 结构,通过在残差结构中结合使用长跳连接和短跳连接,有效加快了网络的收敛速度,提高了网络的重建性能。但是,现有算法在重建图像的几何结构、纹理细节方面等高频信息方面仍然表现不尽如人意,所以本文提出了一种双通道残差网络来解决这个问题。

2 用于图像超分辨率的双通道残差网络

2.1 网络结构

本文网络由浅层特征提取模块、特征映射模块以及重建模块共 3 大模块组成,网络结构如图 1 所示。本文用 I_{LR} 表示网络的输入。在特征预提取模块中,使用了一个卷积层对输入图像进行浅层特征的提取,抽象为公式

$$F_0 = H_{SF}(I_{LR}) \tag{1}$$

式中: H_{SF} 表示提取浅层特征的卷积运算; F_0 表示从 I_{LR} 中提取的浅层特征。

在特征映射模块中,使用双通道的残差结构作为网络的基础模块,通过堆叠基础模块加速网络。这种结构增加了残差学习的感受野与速度,加强了网络的学习能力。将第 N 个残差组的输入与运算分别表示为 F_{N-1} 与 $R_N(F_{N-1})$,辅助通道的运算表示为 $A_N(F_{N-1})$,则第 N 个基础模块的输出为主通

道输出与辅助通道输出的点积,表示为

$$F_N = R_N(F_{N-1}) \otimes A_N(F_{N-1}) \tag{2}$$

式中: F_N 表示第 N 个基础模块的输出结果; $\otimes$ 表示矩阵的点乘运算。特征映射模块的输出 F_{FM} 视作图像的深层特征, F_{FM} 由第 G 个基础模块的输出 F_G 和浅层特征得到

$$F_{FM} = F_0 + H_C(F_G) \tag{3}$$

式中 H_C 表示特征映射模块最后的卷积层运算。将 F_{FM} 通过重建模块即一个上采样块和一个卷积层后,输出网络重建的图像

$$\hat{I}_{SR} = H_{RM}(F_{FM}) = H_C(H_{UP}(F_{FM})) \tag{4}$$

式中: H_{RM} 和 H_{UP} 分别表示重建模块和重建模块中的上采样模块函数; $\hat{I}_{SR}$ 是网络的重建图像输出。

2.2 辅助通道

现有的图像 SR 中,重建算法经常会使自然图像中所包含的几何结构以及边缘信息出现断裂、弯曲、模糊等现象,导致重建图像难以获得良好的视觉效果。为了在重建过程中更好地保留原始图像的几何结构和边缘信息,本文在主通道之外并构了一个辅助通道,增大了网络的宽度。在辅助通道中加入了自适应结构化卷积,本文改进了传统的膨胀卷积,将输入图像按尺寸比例划分为不同的块,在图像的不同块中采用不同的膨胀率,膨胀卷积可以适应图像局部结构信息的变化,也使辅助通道有了自适应的感受野,让基础模块可以更加专注于图像的高频特征的提取。第 N 个残差组的辅助通道运算可表示为

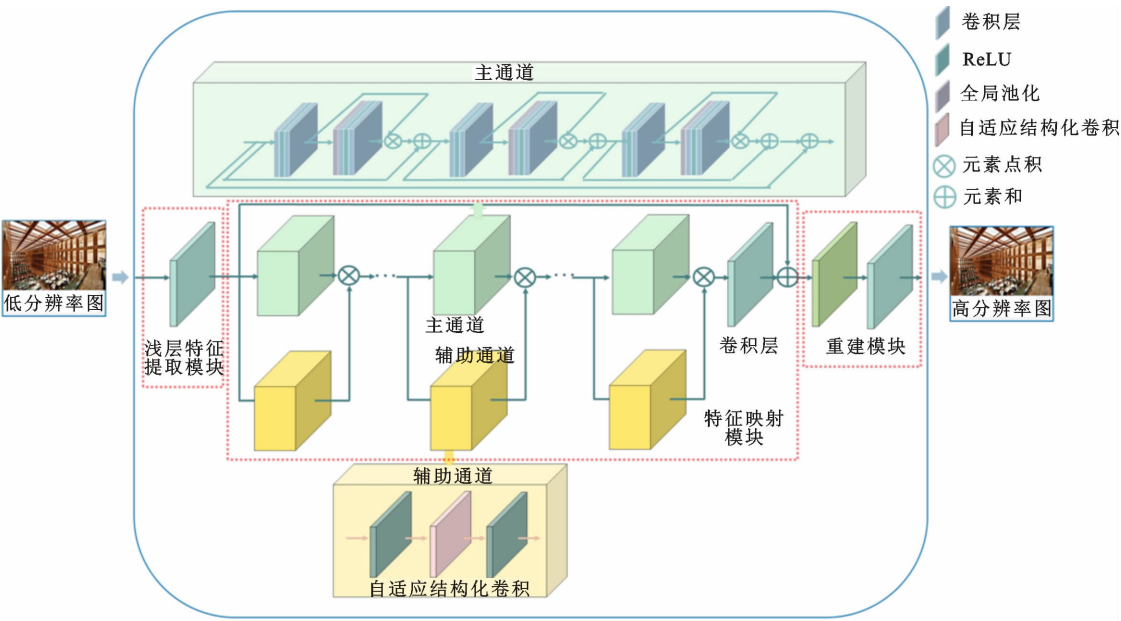


图 1 网络结构

Fig. 1 Network structure

$$\mathbf{A}_N(\mathbf{F}_{N-1}) = \mathbf{H}_C(\mathbf{H}_{IC}(\mathbf{H}_C(\mathbf{F}_{N-1}))) \quad (5)$$

式中 $\mathbf{H}_{IC}$ 为自适应结构化卷积运算,公式为

$$\mathbf{H}_{IC}(\mathbf{P}) = \sum_{i=1}^I \sum_{j=1}^J \mathbf{H}_{DC}(\mathbf{P}_{ij}, \mathbf{r}) \quad (6)$$

其中, I, J 为图像横纵的分块上限, $\mathbf{r}$ 为膨胀卷积的膨胀率, $\mathbf{H}_{DC}$ 表示膨胀卷积。

2.3 残差通道注意力块

每一个主通道模块都由 K 个残差通道注意力块构成,在第 g 个主通道模块中,第 b 个残差通道注意力块的输出可表示为

$$\mathbf{F}_{g,b} = \mathbf{F}_{g,b-1} + \mathbf{H}_{CA}(\mathbf{X}_{g,b}) \quad (7)$$

式中 $\mathbf{X}_{g,b}$ 和通道注意力函数 $\mathbf{H}_{CA}$ 分别为

$$\mathbf{X}_{g,b} = \mathbf{H}_C(\delta(\mathbf{H}_C(\mathbf{F}_{g,b-1}))) \quad (8)$$

$$\mathbf{H}_{CA}(\mathbf{X}_{g,b}) = \mathbf{X}_{g,b} \otimes \mathbf{S}(\mathbf{H}_C(\delta(\delta_C(\mathbf{H}_G(\mathbf{X}_{g,b})))) \quad (9)$$

其中 $\delta, \mathbf{S}$ 和 $\mathbf{H}_G$ 分别表示 ReLU 激活函数、Sigmoid 激活函数和全局平均池化函数。

2.4 损失函数

重建图像 $\hat{\mathbf{I}}_{SR}$ 在结构与细节等高频信息方面应该与 $\mathbf{I}_{LR}$ 所对应的高分辨率图像 $\mathbf{I}_{HR}$ 尽量相接近。所以,研究人员在损失函数中引入了多尺度结构相似度(MS_SSIM)损失^[16],但是仅以 MS_SSIM 损失作为参考容易导致重建图像的颜色和亮度产生偏差。因此,本文结合 L_1 损失与多尺度结构相似性损失 L_{MS_SSIM} 为总的损失函数,公式为

$$L_{total} = (1 - \alpha)G_{gu}L_1 + \alpha L_{MS_SSIM} \quad (10)$$

式中: α 为平衡参数; G_{gu} 为高斯分布变量。 L_1 损失函数和 L_{MS_SSIM} 损失函数的定义如下

$$L_1 = \frac{1}{k} \sum_{i=1}^k \|\hat{\mathbf{I}}_{SR} - \mathbf{I}_{HR}\|_1 \quad (11)$$

$$L_{MS_SSIM} = 1 - MS(\hat{\mathbf{I}}_{SR}, \mathbf{I}_{HR}) \quad (12)$$

式中: k 为训练图像数量; $MS^{[17]}$ 为多尺度结构相似性运算。

3 实验结果与分析

3.1 数据集

本文的训练数据集为 DIV2K,测试数据集为 Set5^[18]、Set14^[19]、BSD100^[20]、Urban100^[21]。网络的输入通过对数据集中的原始图像进行双 3 次插值下采样得到。

3.2 参数设置

特征预提取模块使用 64 个尺寸为 3×3 像素的卷积核。特征映射模块的主通道由 7 个主通道模块和一个额外的卷积层构成。每个主通道模块为 3 个

相同的残差通道注意力结构的叠加,其中的卷积层均采用尺寸为 3×3 像素的卷积核和 0 填充,激活函数使用 ReLU,池化层采用全局池化。特征映射模块尾部的卷积层使用 3×3 像素的卷积核。辅助通道深度为 3,第二层使用自适应的膨胀卷积。重建模块中的最后一个卷积层使用 3 个 3×3 像素的卷积核完成重建图像输出。平衡参数 α 取 0.84,自适应膨胀卷积的膨胀率上限设为 4,图像横纵的分块上限 I, J 设为 10。网络的搭建、训练与测试均在 Ubuntu16.04 系统上使用 Pytorch 深度学习框架完成,实验硬件平台如下:CPU 型号为 Intel(R)Core (TM)i9-9900K CPU @ 3.60 GHz,内存容量为 32 GB,GPU 使用两块 RTX2080Ti。

3.3 实验结果及分析

3.3.1 视觉效果 将本文网络与现有同类网络 SRCNN^[10]、VDSR^[8]、DRRN^[11]、MemNet^[12]、MSRN^[13]、CARN^[14]、OISR^[15] 进行了对比,结果如图 2 所示。可以看出:SRCNN 网络重建图像的几何结构效果最差,重建二幅图像中的框架形状均出现严重扭曲变形;MSRN 网络及 OISR 网络恢复图像的视觉效果较好,但是重建图像的几何结构仍然有部分失真;本文网络可以准确恢复两幅图像中的框架结构,重建图像具有最佳的视觉效果,而且在纹理细节视觉效果上也优于其他网络的。这是由于本文网络在主通道之外并构了一个辅助通道和使用 L_1 与 L_{MS_SSIM} 相结合的损失函数。辅助通道使得网络对于图像高频信息的学习能力更强,让网络保留了更多图像的几何结构信息。实验的视觉效果证明,本文网络学习结构信息的能力更强,重建图像的效果更为逼真。

3.3.2 定量分析 图像超分辨重建质量的评价指标分为主观和客观两种。主观评价一般采用人眼判定图像质量的方法,这种方式较为主观,评价较为片面,因此一般将主观评价作为评价图像质量的辅助指标。早期人们一般采用峰值信噪比(PSNR,用符号 I_{PSNR} 表示)作为客观评价指标,其为信号最大功率与噪声功率的比值,在本文的计算中取图像最大像素值的平方与图像均方误差比值的对数。然而,峰值信噪比是对重建像素点误差敏感的评价指标,没有考虑到人眼的视觉特性。因此,常常会出现 PSNR 评价与人的主观评价不一致的情况。所以,学界提出了另一个客观评价指标——结构相似性(SSIM,用符号 I_{SSIM} 表示)^[22]。该指标分别从亮度、对比度、结构共 3 个方面对比与原始图像的相似度,

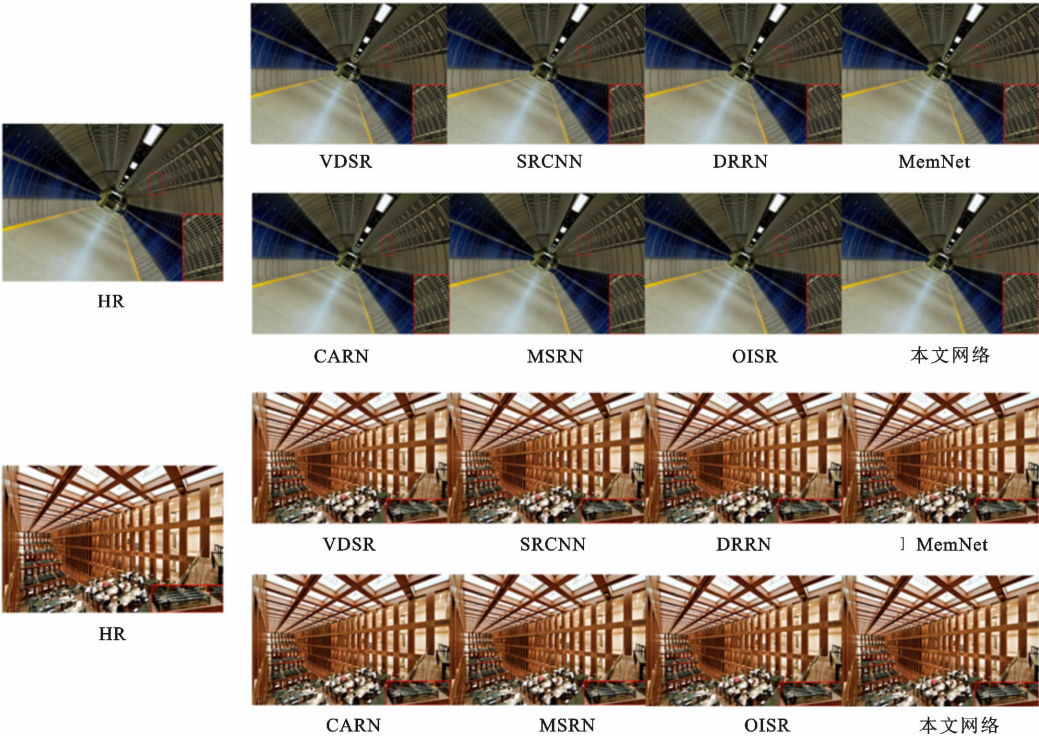


图 2 自然图像上不同网络的性能比较

Fig. 2 Performance comparison of different networks on natural images

更注重人的主观感受。将两种指标结合可以更客观全面地评估重建图像的质量。

本文网络与现有 7 个网络在 4 个数据集上进行

了 2、3、4 倍放大重建,PSNR 与 SSIM 如表 1 所示。可以看出:只采用了 3 个卷积层的 SRCNN 网络的重建结果客观评价指标最差;MSRN 网络将全局特

表 1 不同网络的峰值信噪比和结构相似性

Table 1 PSNR and SSIM of different networks

网络	放大倍数	Set5		Set14		BSD100		Urban100	
		I_{PSNR}/dB	I_{SSIM}	I_{PSNR}/dB	I_{SSIM}	I_{PSNR}/dB	I_{SSIM}	I_{PSNR}/dB	I_{SSIM}
SRCNN	2	36.66	0.954 2	32.45	0.906 7	31.36	0.887 9	29.50	0.894 6
	3	32.75	0.909 0	29.30	0.821 5	28.41	0.786 3	26.24	0.798 9
	4	30.48	0.862 8	27.50	0.751 3	26.90	0.710 1	24.52	0.722 1
VDSR	2	37.53	0.959 0	33.05	0.913 0	31.90	0.896 0	30.77	0.914 0
	3	33.67	0.921 0	29.78	0.832 0	28.83	0.799 0	27.14	0.829 0
	4	31.35	0.883 0	28.02	0.768 0	27.29	0.725 1	25.18	0.754 0
DRRN	2	37.74	0.959 1	33.23	0.913 6	32.05	0.897 3	31.23	0.918 8
	3	34.03	0.924 4	29.96	0.834 9	28.95	0.800 4	27.53	0.837 8
	4	31.68	0.888 8	28.21	0.772 1	27.38	0.728 4	25.44	0.763 8
MemNet	2	37.78	0.959 7	33.28	0.914 2	32.08	0.897 8	31.31	0.919 5
	3	34.09	0.924 8	30.00	0.835 0	28.96	0.800 1	27.56	0.837 6
	4	31.74	0.889 3	28.26	0.772 3	27.40	0.728 1	25.50	0.763 0
MSRN	2	38.08	0.960 5	33.74	0.917 0	32.23	0.901 3	32.22	0.932 6
	3	34.38	0.926 2	30.34	0.839 5	29.08	0.804 1	28.08	0.855 4
	4	32.07	0.890 3	28.60	0.775 1	27.52	0.727 3	26.07	0.789 6
CARN	2	37.76	0.959 0	33.52	0.916 6	32.09	0.897 8	31.92	0.925 6
	3	34.29	0.925 5	30.29	0.840 7	29.06	0.803 4	28.06	0.849 3
	4	32.13	0.893 7	28.60	0.780 6	27.58	0.734 6	26.07	0.783 7

续表

网络	放大倍数	Set5		Set14		BSD100		Urban100	
		$I_{\text{PSNR}}/\text{dB}$	I_{SSIM}	$I_{\text{PSNR}}/\text{dB}$	I_{SSIM}	$I_{\text{PSNR}}/\text{dB}$	I_{SSIM}	$I_{\text{PSNR}}/\text{dB}$	I_{SSIM}
OISR	2	38.12	0.960 9	33.78	0.919 6	32.26	0.900 7	32.52	0.932 0
	3	34.56	0.928 4	30.46	0.845 0	29.20	0.807 7	28.56	0.860 6
	4	32.33	0.896 8	28.73	0.784 5	27.66	0.738 9	26.38	0.795 3
本文网络	2	38.14	0.961 2	33.93	0.919 7	32.31	0.901 0	32.87	0.933 4
	3	34.60	0.928 5	30.53	0.845 8	29.23	0.809 2	28.68	0.862 9
	4	32.43	0.896 9	28.74	0.784 7	27.67	0.741 3	26.46	0.798 2

征与局部多尺度特征相结合,使得网络对于特征的提取更加高效;OISR 网络改进了残差块的结构,使得网络的学习能力更强,MSRN 和 OISR 这两种网络的客观评价指标更优;本文网络的两种客观指标均为最高,在几何结构内容较多的 Urban 数据集上进行放大倍数为 2 的重建时,PSNR 以及 SSIM 分别为 32.87 dB 和0.933 4,均优于现有结果最好的 OISR 算法的,在自然景色占比更多的 BSD100 数据集上进行放大倍数为 2 的重建时,PSNR 以及 SSIM 分别为 32.31 dB 和0.901 0。实验结果表明,本文网络对于图像几何结构的学习能力更强。

为了证明本文网络设计的有效性,进行了放大倍数为 2 的消融实验,结果如表 2 所示。可以看出:在去掉辅助通道的情况下,在 Set5 数据集上进行重建,相较于双通道网络的 PSNR 以及 SSIM 分别下降了 5.86 dB 和0.068 1,在 Urban 数据集上进行重建,相较于双通道网络的 PSNR 以及 SSIM 分别下降了 6.7 dB 和0.141 4;在仅使用 L_1 损失函数的情况下,在 Set5 数据集上重建,相较于使用 L_{total} 损失函数的 PSNR 以及 SSIM 分别下降了 2.98 dB 和 0.048 7,在 Urban 数据集上进行重建,相较于使用 L_{total} 损失函数的 PSNR 以及 SSIM 分别下降了 4.18 dB 和 0.112。消融实验证明,本文设计的辅助通道和损失函数能够有效提升网络的学习能力。

表 2 消融实验结果

Table 2 Ablation experiment results

数据集	双通道仅使用 L_1					
	主通道		损失函数训练		双通道	
	$I_{\text{PSNR}}/\text{dB}$	I_{SSIM}	$I_{\text{PSNR}}/\text{dB}$	I_{SSIM}	$I_{\text{PSNR}}/\text{dB}$	I_{SSIM}
Set5	32.28	0.893 1	35.16	0.912 5	38.14	0.961 2
Urban100	26.17	0.792 0	28.69	0.821 4	32.87	0.933 4

4 结 论

(1)本文提出了一种用于图像超分辨重建的双通道残差网络。使用带有跳跃连接和通道注意力模

块的残差组作为主通道,并加入使用自适应结构化卷积的辅助通道,组成了双通道的残差块,使得低频信息得以更好地绕过网络,加强了高频信息在网络中的作用,为模块增加了自适应的感受野,获得了更佳的重建效果。

(2)本文使用 L_1 损失函数和多尺度结构相似度损失函数,在训练中即能够较好地保持图像的颜色和亮度,也能够保留图像的边缘、纹理细节等高频信息。

(3)实验结果证明,本文网络在自然图像 SR 重建领域的客观评价指标与视觉效果均优于现有网络的。本文网络重建的自然图像具有更完整、准确的几何结构。

参考文献:

[1] 张哲,齐春,张钊强,等. 人脸超分辨率重建中投影空间的选择方法 [J]. 西安交通大学学报,2018,52(8): 43-48.
ZHANG Zhe, QI Chun, ZHANG Zhaoqiang, et al. A selection method of projection space for face super-resolution reconstruction [J]. Journal of Xi'an Jiaotong University, 2018, 52(8): 43-48.

[2] 宋长明,王赞. 融合低秩和稀疏表示的图像超分辨率重建算法 [J]. 西安交通大学学报,2018,52(7): 18-24.
SONG Changming, WANG Yun. Super resolution reconstruction algorithm combined low rank with sparse representation [J]. Journal of Xi'an Jiaotong University, 2018, 52(7): 18-24.

[3] 李玉花,齐春. 利用位置字典对的人脸图像超分辨率方法 [J]. 西安交通大学学报,2012,46(6): 7-11.
LI Yuhua, QI Chun. A super-resolution method for face images using position-dictionary pairs [J]. Journal of Xi'an Jiaotong University, 2012, 46(6): 7-11.

[4] 马祥,齐春. 全局重建和位置块残差补偿的人脸图像超分辨率算法 [J]. 西安交通大学学报,2010,44(4): 9-12.
MA Xiang, QI Chun. A face image super-resolution

- algorithm based on global reconstruction and position-patch residue compensation [J]. Journal of Xi'an Jiaotong University, 2010, 44(4): 9-12.
- [5] 付利华, 孙晓威, 赵宇, 等. 基于运动特征融合的快速视频超分辨率重构方法 [J]. 模式识别与人工智能, 2019, 32(11): 1022-1031.
- FU Lihua, SUN Xiaowei, ZHAO Yu, et al. Fast video super-resolution reconstruction method based on motion feature fusion [J]. Pattern Recognition and Artificial Intelligence, 2019, 32(11): 1022-1031.
- [6] LEDIG C, THEIS L, HUSZÁR F, et al. Photo-realistic single image super-resolution using a generative adversarial network [C] // Proceedings of the 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). Piscataway, NJ, USA: IEEE, 2017: 105-114.
- [7] SHI Wenzhe, CABALLERO J, LEDIG C, et al. Cardiac image super-resolution with global correspondence using multi-atlas patchmatch [C] // Proceedings of the 2013 Medical Image Computing and Computer-Assisted Intervention (MICCAI). Cham, Germany: Springer, 2013: 9-16.
- [8] KIM J, 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 (CVPR). Piscataway, NJ, USA: IEEE, 2016: 1646-1654.
- [9] ZHANG Yulun, LI Kunpeng, LI Kai, et al. Image super-resolution using very deep residual channel attention networks [C] // Proceedings of the 2018 European Conference on Computer Vision (ECCV). Cham, Germany: Springer, 2018: 294-310.
- [10] DONG Chao, LOY C C, HE Kaiming, et al. Learning a deep convolutional network for image super-resolution [C] // Proceedings of the 2014 European Conference on Computer Vision (ECCV). Cham, Germany: Springer, 2014: 184-199.
- [11] TAI Ying, YANG Jian, LIU Xiaoming. Image super-resolution via deep recursive residual network [C] // Proceedings of the 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). Piscataway, NJ, USA: IEEE, 2017: 2790-2798.
- [12] TAI Ying, YANG Jian, LIU Xiaoming, et al. MemNet: a persistent memory network for image restoration [C] // Proceedings of the 2017 IEEE International Conference on Computer Vision (ICCV). Piscataway, NJ, USA: IEEE, 2017: 4549-4557.
- [13] LI Juncheng, FANG Faming, MEI Kangfu, et al. Multi-scale residual network for image super-resolution [C] // Proceedings of the 2018 European Conference on Computer Vision (ECCV). Cham, Germany: Springer, 2018: 527-542.
- [14] AHN N, KANG B, SOHN K A. Fast, accurate, and lightweight super-resolution with cascading residual network [C] // Proceedings of the 2018 European Conference on Computer Vision (ECCV). Cham, Germany: Springer, 2018: 256-272.
- [15] HE Xiangyu, MO Zitao, WANG Peisong, et al. ODE-inspired network design for single image super-resolution [C] // Proceedings of the 2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). Piscataway, NJ, USA: IEEE, 2019: 1732-1741.
- [16] ZHAO Hang, GALLO O, FROSIO I, et al. Loss functions for image restoration with neural networks [J]. IEEE Transactions on Computational Imaging, 2017, 3(1): 47-57.
- [17] WANG Z, SIMONCELLI E P, BOVIK A C. Multi-scale structural similarity for image quality assessment [C] // Proceedings of the 37th Asilomar Conference on Signals, Systems & Computers. Piscataway, NJ, USA: IEEE, 2003: 1398-1402.
- [18] BEVILACQUA M, ROUMY A, GUILLEMOT C, et al. Low-complexity single-image super-resolution based on nonnegative neighbor embedding [C] // Proceedings of the British Machine Vision Conference. Guildford, UK: BMVA, 2012: 135.
- [19] DE SOUZA FARIA G, KIM H Y. Differential audio analysis: a new side-channel attack on PIN pads [J]. International Journal of Information Security, 2019, 18(1): 73-84.
- [20] 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 8th IEEE International Conference on Computer Vision (ICCV). Piscataway, NJ, USA: IEEE, 2001: 416-423.
- [21] HUANG Jiabin, 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 (CVPR). Piscataway, NJ, USA: IEEE, 2015: 5197-5206.
- [22] WANG Zhou, BOVIK A C, SHEIKH H R, et al. Image quality assessment: from error visibility to structural similarity [J]. IEEE Transactions on Image Processing, 2004, 13(4): 600-612.

(编辑 陶晴)