

学术前沿专栏:甲骨智能计算

【特约主持人】李雪山:河南师范大学二级教授;

刘栋:河南师范大学教授

【主持人按语】甲骨文作为中华民族珍贵的文化遗产,是研究商周社会历史与汉民族语言文字起源的重要载体.近年来,随着人工智能、大数据等新一代信息技术的迅猛发展,甲骨学研究正迎来智能计算与人文科学深度融合的重大机遇.党和国家高度重视中华优秀传统文化的传承与创新,强调要运用现代科技手段加强古籍保护与整理研究,推动中华优秀传统文化的创造性转化与创新性发展.在此背景下,积极探索甲骨文图像修复、分期断代与智能识别等前沿方向,已成为推动古文字数字化、智能化发展的重要路径.在论文《R2PixGAN:一种高效的甲骨文拓片去噪新方法》中,作者提出了一种基于生成对抗网络的拓片图像去噪算法,显著提升了甲骨拓片的清晰度与可辨识度.在论文《FEC-PVT:基于PVT架构的甲骨钻凿图像分割网络》中,作者构建了融合特征增强与上下文感知的钻凿形态分析模型,为甲骨钻凿工艺的智能识别与分类提供了有效工具.

R2PixGAN:一种高效的甲骨文拓片去噪新方法

王士斌^{a,b},王宇^c,喻琪^c,刘栋^b,闫娟^c

(河南师范大学 a.计算机与信息工程学院;b.河南省教育人工智能与个性化学习重点实验室;

c.甲骨智能计算实验室,河南 新乡 453007)

摘要:随着技术的发展,深度学习在图像去噪领域取得了显著的效果.然而,由于甲骨拓片往往同时包括各种噪声,现有的去噪模型无法适应甲骨文独特的字形和复杂的文字背景.针对上述挑战,提出了一种基于条件对抗网络(Pix2Pix)的图像去噪方法 R2PixGAN.在该方法中,生成器采用了 R2U-Net 模型,该模型不仅保留了传统 U-Net 在特征提取方面的优势,还通过引入循环神经网络(RNN)结构,进一步提升了图像重建能力,同时增强了去噪效果.此外,还将感知损失纳入模型,以更好地保留原始图像的关键细节和特征.实验结果表明,R2PixGAN 在 PSNR 和 SSIM 分数方面优于对比实验,图像去噪效果得到了明显的提升.

关键词:图像去噪;甲骨文;感知损失

中图分类号:TP399

文献标志码:A

文章编号:1000-2367(2026)01-0001-07

甲骨文拓片作为中国古代文化遗产的重要组成部分,承载着丰富的历史和文化信息.然而,由于长期的埋藏和保存过程中的各种因素,甲骨文拓片往往受到噪声的严重干扰,出现文字模糊不清,难以辨识.这不仅

收稿日期:2024-12-11;修回日期:2025-01-09.

基金项目:国家自然科学基金(62072160);河南省高等学校重点科研项目(24A520018).

作者简介:王士斌(1981-),男,河南滑县人,河南师范大学副教授,博士,研究方向为图像处理、多任务遗传算法优化,
E-mail: wangshibin@htu.edu.cn.

通信作者:刘栋, E-mail: liudong@htu.edu.cn; 闫娟, E-mail: yanjuan@htu.edu.cn.

引用本文:王士斌,王宇,喻琪,等.R2PixGAN:一种高效的甲骨文拓片去噪新方法[J].河南师范大学学报(自然科学版),
2026,54(1):1-7.(Wang Shibin, Wang Yu, Yu Qi, et al. R2PixGAN: an efficient new method for denoising oracle
bone topographies[J]. Journal of Henan Normal University(Natural Science Edition), 2026, 54(1): 1-7. DOI: 10.
16366/j.cnki.1000-2367.2024.12.11.0004.)

给甲骨文的解读和研究带来了极大的困难,也限制了其作为文化遗产的保存和传播.因此,研究破损文字图像的去噪方法,对于人类历史的保存和发展无疑具有重要意义.

近年来,深度学习技术快速发展,图像去噪迎来了新的研究途径.Pix2Pix^[1]通过对抗训练和精心设计的损失函数,能够将带有噪声的图像直接映射到无噪声的真实目标图像上,从而显著提高了生成图像的质量和逼真度.通过生成对抗网络(GANs)^[2]的训练,Pix2Pix的生成器与判别器相互对抗,促使生成器不断改进图像去噪的效果,使去噪后的图像更加逼真且保留更多的细节信息.TransUNet^[3]和UFormer^[4]等方法则在传统的U型网络架构中引入了自我注意力模块,这些模型结合了卷积神经网络(CNN)的局部特征捕捉能力和Transformer的全局建模能力,通过在编码和解码过程中堆叠多头自注意力机制,有效捕捉图像的长程依赖关系和全局特征,从而在去噪过程中更好地恢复图像的细节信息.与传统的卷积神经网络不同,自我注意机制使模型能够关注图像中重要的区域,并在不同尺度下整合全局与局部信息,从而提升了去噪效果.

尽管目前已有一些去噪网络被应用于甲骨文拓片的处理,但由于甲骨文独特的字形结构以及复杂的背景纹理,使得传统的去噪技术往往难以在去除噪声和保留关键细节之间取得理想的平衡,从而影响了处理后甲骨文的清晰度、完整性和准确性.因此,现有去噪方法在处理甲骨拓片时面临两大主要问题:一是现有模型对复杂背景噪声的去除能力有限,可能会过度去噪,出现文字笔画的细节丢失;二是现有方法在保留甲骨文独特的字形结构和语义一致性方面存在不足,这可能影响去噪后图像的实际应用.为了解决这些问题,本文提出了一种新的图像生成网络结构,称为R2PixGAN.该方法将Pix2Pix^[1]中的生成器替换为R2U-Net^[5].R2U-Net^[5]基于U-Net^[6]架构,通过其多层次的特征提取、残差连接和循环机制,使其能够在不同尺度上捕捉甲骨文的关键特征.此外,为了进一步提升去噪后的甲骨文图像的质量,还引入了感知损失,引导生成器关注图像的高层语义特征和感知质量,使去噪后的甲骨拓片图像在视觉效果上更加清晰自然.本文实验在甲骨文数据集上进行了测试,并与现有先进的图像去噪模型Pix2Pix^[1]、TransUNet^[3]、UFormer^[4]、CharFormer^[7]、MaskedDenoising^[8]进行比较,结果表明本文所提出的算法在客观指标和视觉效果上都有明显的提升.

1 相关工作

1.1 图像修复

在图像修复任务领域,研究的主要目标是通过去除图像中的各种失真或损坏,从而使图像能够更加接近真实场景.随着深度学习技术的兴起,特别是卷积神经网络(CNNs)^[9]的发展,图像修复领域取得了显著进展.ZHANG等^[10]提出了一种名为SCUNet的方法,该方法利用了深度卷积网络的强大特性去除多种类型的噪声.SeNM-VAE^[11]结合了变分自编码器(VAE)和半监督学习,利用少量的标注数据来提高模型性能.LG-BPN^[12]采用了局部和全局的双边传播机制来增强去噪效果.此外,随着视觉转换器(ViT)^[13]的兴起,自注意力机制逐渐也被应用于图像处理领域.图像处理转换器(IPT)^[14]利用自注意力机制在全局特征建模方面的优势,从而在复杂的图像退化问题展现优异的性能.ZHU等^[15]提出了一个结合扩散模型和Plug-and-Play策略的图像恢复方法,能够在无需重新训练的情况下,灵活地处理各种图像恢复任务.此外,DA-CLIP网络^[16]旨在通过结合去噪机制和视觉-语言模型(CLIP)^[17],利用语言信息与图像信息的协同作用,来提高图像去噪的效果.

1.2 文本图像去噪

文本图像的去噪已经成为一个越来越受关注的话题,过去的研究比较了多种基于物理的中文字符图像降噪方法,例如最小化全变分^[18]和小波变换^[19]等技术.然而,这些方法在实际应用中的表现不佳,因为真实世界的字符图像退化比合成噪声更复杂.为了解决这个问题,研究者针对某些类型的噪声设计专门的去噪方法.例如,GANGAMMA等^[20]提出了一种基于积分的去噪方法,用于恢复退化的历史文献图像中的不均匀背景.

最近,一种基于残差学习的局部自适应阈值算法被提出,用于处理具有不均匀光照和低对比度的历史文献图像^[21].ZHANG等^[22]通过引入高斯噪声和椒盐噪声的建模,应用于中文书法字符图像的去噪处理.然而,这些方法的去噪效果仍然存在一些不足.主要问题在于它们难以同时满足两个关键目标:一方面是高效

去除噪声,另一方面是确保字符笔画的完整性,避免因噪声去除而导致笔画的细节丢失.

2 本文方法

2.1 模型框架

R2PixGAN 由一组生成器和判别器构成,其中生成器采用了 R2U-Net^[5] 模型,判别器使用了 Pix2Pix 中的 PatchGAN^[23] 结构,模型结构如图 1 所示.在模型的训练过程中,生成器 G 以带有噪声的甲骨文拓片图像 x 作为输入,通过多层卷积网络的特征提取与还原过程生成去噪图像 $G(x)$.在这一过程中,生成器逐步学习如何有效去除噪声,同时保留甲骨文图像中的关键细节.判别器 D 则用于评估生成器输出图片的质量,通过对输入图像(包括真实图像 y 和生成图像 $G(x)$) 的判别,判断其真实性.判别器的反馈帮助生成器逐步改进其去噪能力,目标是生成的去噪图像 $G(x)$ 足够逼真,使得经过训练的判别器 D 无法区分 $G(x)$ 与真实图像 y .这种对抗性训练策略有效提升了去噪图像的质量,确保了细节的还原和整体的真实性.

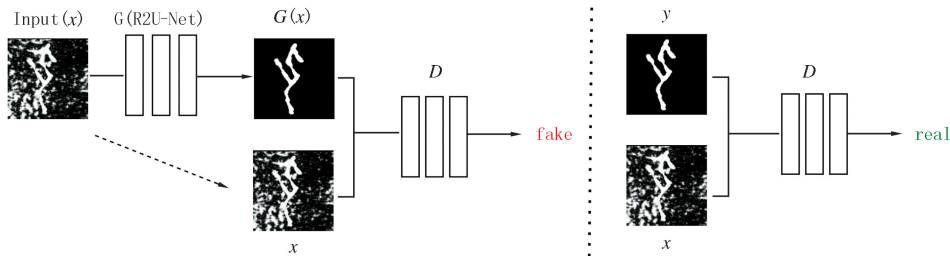


图1 本文算法模型框架

Fig.1 Algorithmic modeling framework of this paper

2.2 R2U-Net 生成器

为实现高质量的图像去噪,生成器采用了 R2U-Net^[5] 模型.如附录图 S1 所示,编码器部分从输入图像开始,通过多个卷积层逐步提取特征,此外,编码器的每一层均采用残差连接,能够缓解深层网络训练中的梯度消失问题,并增强特征的表达能力.

在解码器部分,从编码器输出的高维特征图开始,通过多层卷积逐步恢复图像细节.解码过程中,采用反卷积操作还原图像的空间尺寸,同时结合残差连接,将编码器对应层的特征图与解码器拼接在一起.这种结构不仅保留了图像的局部细节信息,还在全局特征的基础上补充了低层次的纹理细节.最后,解码器通过卷积层生成与输入尺寸相同的特征图,输出最终的去噪图像.

2.3 MicroPatch 判别器

在 GANs 训练过程中,生成器的目标是通过判别器提供的梯度反馈不断优化自身参数,以生成更真实的去噪图像.然而,当输入图像仅存在局部噪声时,传统的全局判别器无法有效识别这些噪声的存在.为了解决这个问题,在训练过程中引入了局部判别器机制.将输入图像划分为多个固定大小的局部区域,局部判别器独立地对每个小区域进行判别.与整幅图像进行整体判断相比,这些局部噪声区域的输出得分较低,提供了更具针对性的梯度信息.

如图 2 所示,判别器模型采用了 PatchGAN^[23] 结构.PatchGAN 判别器将输入图像划分为多个固定大小的局部区域,分别对每个区域的真实性进行评估,以判断其是否源自真实图像或是由生成器产生的合成图像.最终,每个局部区域都会获得一个对应的真实性评分,以所有局部区域的平均响应值作为整张图像的综合判断结果.

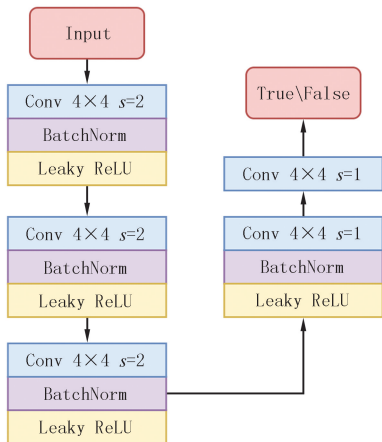


图2 判别器模型框架图

Fig.2 Discriminator modeling framework diagram

2.4 损失函数

2.4.1 对抗损失

对抗损失^[2]是 GANs 中的关键组成部分,用于训练生成器和判别器.对抗损失的设计包括以下两个方面:

a)生成器对抗损失:生成器的目标是生成尽可能逼真的图像.生成器的对抗损失可以表示为:

$$L_{\text{GAN}}^G = E_{x,y} [\ln(1 - D(x, G(x)))],$$

其中 x 代表输入图像, y 代表目标图像, $G(x)$ 是生成的图像, $D(x, G(x))$ 是判别器对生成图像的判断概率.

b)判别器对抗损失:判别器希望能够准确地区分真实图像和生成图像.判别器的对抗损失可以表示为:

$$L_{\text{GAN}}^D = E_{x,y} [\ln D(x, y)] + E_{x,z} [\ln(1 - D(x, G(x)))],$$

其中 $D(x, y)$ 是判别器对真实图像的判断概率, $D(x, G(x))$ 是判别器对生成图像的判断概率.

总对抗损失是这两个部分的总和,公式如下: $L_{\text{adv}}(G, D) = L_{\text{GAN}}^G + L_{\text{GAN}}^D$.

2.4.2 L1 损失

在图像去噪任务中,仅依赖对抗损失无法完全确保有效的去噪.为了解决这一问题, Pix2Pix^[1] 将对抗损失^[2]与 L1 损失^[24]相结合,从而在有效去除噪声的同时保留图像的细节和结构. L1 损失^[24]可以表示为:

$L_{L1}(G) = \frac{1}{N} \sum_{i=1}^N |G(x)_i - y_i|$, 其中, N 是图像中像素的总数, $G(x)_i$ 是生成图像在位置 i 的像素值, y_i 是目标图像在位置 i 的像素值, $|\cdot|$ 表示绝对值.

2.4.3 感知损失

传统的 L1 损失^[24]不足以捕捉甲骨文的复杂细节和结构.而感知损失^[25]则通过使用卷积神经网络提取的特征图来衡量图像之间的差异,使得模型在训练时更加注重结构性信息的保留.因此,在本工作中,引入了感知损失函数,旨在通过引入更具表达能力的特征级别差异度量,改善模型的细节保留能力,其计算方法如下:

$L_{\text{per}} = \frac{1}{N} \sum_{i=1}^N (F_i(x) - F_i(y))^2$, 其中, x 是输入图像, y 是目标图像, $F_i(x)$ 和 $F_i(y)$ 分别表示它们在预训练神经网络第 i 层中的特征表示, N 是特征层的数量.

2.4.4 总损失

λ_1 和 λ_2 是每个损失项的权重因子,用于平衡不同损失的影响.所有损失项,完整的目标函数被表述为:

$L_{(G,D)} = L_{\text{adv}}(G, D) + \lambda_1 L_{L1}(G) + \lambda_2 L_{\text{per}}$, 其中, λ_1 和 λ_2 是每个损失项的权重因子,用于平衡不同损失的影响.

3 实验

3.1 数据集

甲骨文字图像去噪数据集来自河南省甲骨文信息处理教育部重点实验室.该数据集包含 41 957 组甲骨文图像,每组图像包括带噪声图像及对应的无噪声图像,每张图像的宽度和高度为 130 像素.该去噪数据集被划分为训练集和测试集,从中选择了 5 000 组甲骨文图像用于网络的训练,并选择了 40 组图像用于测试.

3.2 实验设置

R2PixGAN 模型在深度学习平台 PyTorch 上实现,并使用 A100 GPU 和 Python 3.8 进行训练.训练时采用 Adam 优化器,并运行 200 个 epoch.初始学习率设置为 0.000 1,批量大小设置为 32.训练使用的数据集图像尺寸为 128×128 .在以下实验中,总损失函数的超参数设置为 $\lambda_1 = 100, \lambda_2 = 5$.

3.3 对比分析

在本节中,与其他图像去噪模型进行了对比分析,附录图 S2 展示了可视化结果.从生成的图像可以看出, Pix2Pix^[1] 去噪后依然存在大量噪声; TransUNet^[3] 在一些笔画较为复杂或背景噪声较强的区域依旧残留较多噪声; MaskedDenoising^[8] 在准确修复甲骨文的字体笔画方面仍然存在较大不足.相比之下,本文的模型通过引入 R2U-Net 生成器和感知损失,不仅能够有效去除复杂噪声,还能更准确地恢复甲骨文的细节和结构特征.

表 1 展示了各模型在 PSNR 和 SSIM 指标下的具体评估结果. Pix2Pix^[1] 在笔画细节完整性和图像真实

性上效果较差;而 TransUNet^[3] 在处理复杂背景噪声时效果有限,PSNR 和 SSIM 分数仍不理想.UFormer^[4] 和 CharFormer^[7] 的指标也有提高,但在笔画的连续性和完整性方面依旧存在较大不足.

表 1 不同图像去噪模型的客观指标评价结果

Tab. 1 Evaluation results of different image denoising models with objective metrics

Methods	Pix2Pix	TransUNet	UFormer	CharFormer	Masked	Ours
PSNR ↑	16.66	17.98	18.90	19.87	19.88	21.40
SSIM ↑	0.829 153	0.813 623	0.588 915	0.685 689	0.764 513	0.924 612

3.4 效率分析

为了评估 R2pixGAN 的效率,从计算复杂度(FLOPs)、参数量(Params)以及推理时间(Runtime)等多个维度对其进行全面分析.表 2 显示了不同模型在这些指标上的对比结果.虽然 R2pixGAN 的参数量相比部分模型略有增加,但其显著的去噪效果充分证明了计算开销的合理性和必要性.

表 2 不同模型之间的效率比较

Tab. 2 Comparison of efficiency between different models

Metrics	Pix2Pix	TransUNet	UFormer	CharFormer	Masked	Ours
FLOPS/G	5.28	63.74	10.27	24.56	73.70	39.00
Params/M	44.59	22.75	20.60	13.93	3.30	41.85
Runtime/s	0.010	0.035	0.027	0.037	0.042	0.030

3.5 消融实验

为了评估所提方法中每个组件的贡献,进行了以下消融实验,结果如表 3 所示.可以看出,每个组件的引入或移除都会影响去噪效果.完整的模型在去噪任务中取得了最佳的结果.

表 3 在甲骨文数据集上的消融实验

Tab. 3 Ablation experiments on the oracle dataset

Variants	Base	V1	V2	V3
Perceptual Loss	w/o	✓	✓	w/o
R2U-Net	w/o	w/o	✓	✓
PSNR ↑	16.66	17.02	21.40	20.82
SSIM ↑	0.829 153	0.838 868	0.924 612	0.891 416

在局部消融实验中,图 3 显示了不同配置下生成的去噪图像.R2U-Net 的引入显著增强了去噪性能.在加入感知损失后,去噪效果得到了明显改善.这表明感知损失能够引导网络关注目标图像的高级特征,从而使生成的去噪结果更接近于真实图像,并显著提升了细节的保留效果.

3.6 泛化实验

为了验证 R2pixGAN 模型的泛化能力,从 OBI-241 数据集里的扫描字里随机挑选了 100 张甲骨文图片进行测试.实验结果如图 4 所示.

相比其他模型,R2pixGAN 不仅有效抑制了复杂背景噪声,还较好地保留和还原了字符的细节与结构完整性.尤其是在绿色框标注的区域,R2pixGAN 显示出更清晰、自然的字符边缘,充分展现了其在复杂噪声环境下的泛化能力.

4 结 论

本研究提出了一种基于条件对抗网络的图像去噪模型 R2PixGAN.该模型由 R2U-Net 生成器和 Patch-GAN 的判别器组成.由于其独特的编码器-解码器结构、残差连接、有效的特征学习以及对不同类型噪声的鲁棒性,R2U-Net 在去噪效果、细节保留、适应性和性能评估方面表现出色.此外,在 R2PixGAN 中加入了感

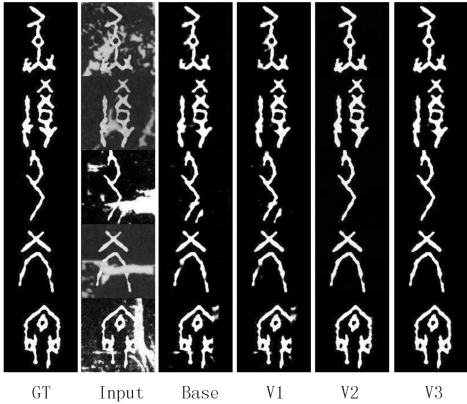


图3 消融实验的可视化结果

Fig.3 Visualization results of ablation experiments

知损失,以减少图像中的失真和伪影,从而提高生成图像的质量.实验结果表明,本文提出的方法在甲骨拓片去噪方面优于其他去噪方法.

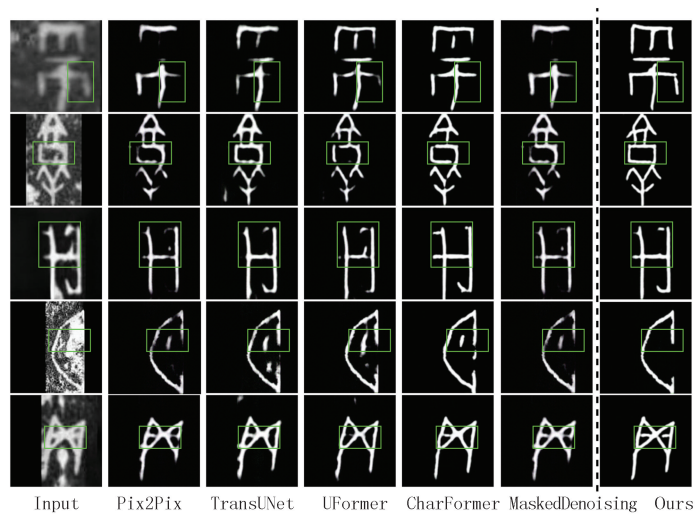


图4 泛化实验的可视化结果

Fig.4 Visualization results of the generalization experiments

附录见电子版(DOI:10.16366/j.cnki.1000-2367.2024.12.11.0004).

参 考 文 献

[1] ISOLA P,ZHU J-Y,ZHOU T H,et al.Image-to-image translation with conditional adversarial networks[C]//2017 IEEE Conference on Computer Vision and Pattern Recognition(CVPR).Honolulu:IEEE,2017:5967-5976.

[2] GOODFELLOW I,POUGET-ABADIE J,MIRZA M,et al.Generative adversarial networks[J].Communications of the ACM,2020,63(11):139-144.

[3] CHEN J N,LU Y Y,YU Q H,et al.TransUNet:transformers make strong encoders for medical image segmentation[EB/OL].[2024-12-06].<https://arxiv.org/abs/2102.04306>.

[4] WANG Z D,CUN X D,BAO J M,et al.Uformer;a general U-shaped transformer for image restoration[C]//2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition(CVPR).New Orleans:IEEE,2022:17662-17672.

[5] ALOM M Z,HASAN M,YAKOPCIC C,et al.Recurrent residual convolutional neural network based on U-Net(R2U-Net)for medical image segmentation[EB/OL].[2024-12-06].<https://arxiv.org/abs/1802.06955>.

[6] RONNEBERGER O,FISCHER P,BROX T.U-Net:convolutional networks for biomedical image segmentation[C]//Medical Image Computing and Computer-Assisted Intervention-MICCAI 2015.Cham:Springer,2015:234-241.

[7] SHI D Q,DIAO X L,SHI L D,et al.CharFormer:a glyph fusion based attentive framework for high-precision character image denoising [C]//Proceedings of the 30th ACM International Conference on Multimedia.Lisboa:ACM,2022:1147-1155.

[8] CHEN H Y,GU J J,LIU Y H,et al.Masked image training for generalizable deep image denoising[C]//2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition(CVPR).Vancouver:IEEE,2023:1692-1703.

[9] ZEILER M D,FERGUS R.Visualizing and understanding convolutional networks[C]//Computer Vision-ECCV 2014.Cham:Springer,2014:818-833.

[10] ZHANG K,LI Y W,LIANG J Y,et al.Practical blind image denoising via swin-conv-UNet and data synthesis[J].Machine Intelligence Research,2023,20(6):822-836.

[11] ZHENG D H,ZOU Y H,ZHANG X W,et al.SeNM-VAE:semi-supervised noise modeling with hierarchical variational autoencoder [C]//2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition(CVPR).Seattle:IEEE,2024:25889-25899.

[12] WANG Z C,FU Y,LIU J,et al.LG-BPN:local and global blind-patch network for self-supervised real-world denoising[C]//2023 IEEE/ CVF Conference on Computer Vision and Pattern Recognition(CVPR).Vancouver:IEEE,2023:18156-18165.

[13] DOSOVITSKIY A,BEYER L,KOLESNIKOV A,et al.An image is worth 16x16 words:transformers for image recognition at scale[EB/ OL].[2024-12-06].<https://arxiv.org/abs/2010.11929>.

[14] CHEN H T,WANG Y H,GUO T Y,et al.Pre-trained image processing transformer[C]//2021 IEEE/CVF Conference on Computer Vi- sion and Pattern Recognition(CVPR).Nashville:IEEE,2021:12294-12305.

[15] ZHU Y Z,ZHANG K,LIANG J Y,et al.Denoising diffusion models for plug-and-play image restoration[C]//2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops(CVPRW).Vancouver:IEEE,2023:1219-1229.

[16] LUO Z W,FREDRIK K,ZHENG Z,et al.Controlling Vision-Language Models for Universal Image Restoration[J].Conference on Computer Vision and Pattern Recognition(CVPR),[S.l.:s.n.],2023:208-214.

[17] RADFORD A,KIM J W,HALLACY C,et al.Learning transferable visual models from natural language supervision[C]//International Conference on Machine Learning.[S.l.:s.n.],2021.

[18] CHAMBOLLE A.An algorithm for total variation minimization and applications[J].Journal of Mathematical Imaging and Vision,2004,20(1):89-97.

[19] LI J C,CHENG B D,CHEN Y,et al.EWT:Efficient Wavelet-Transformer for single image denoising[J].Neural Networks,2024,177:106378.

[20] GANGAMMA B,MURTHY S,SINGHA V.Restoration of Degraded Historical Document Image[J].Journal of Emerging Trends in Computing & Information Sciences,2012,21(4):57-67.

[21] HUANG Z K,WANG Z N,XI J M,et al.Chinese rubbing image binarization based on deep learning for image denoising[C]//Proceedings of the 2nd International Conference on Control and Computer Vision.[S.l.]:ACM,2019:46-50.

[22] ZHANG J L,GUO M T,FAN J P.A novel generative adversarial net for calligraphic tablet images denoising[J].Multimedia Tools and Applications,2020,79(1):119-140.

[23] ZHU J Y,PARK T,ISOLA P,et al.Unpaired image-to-image translation using cycle-consistent adversarial networks[C]//2017 IEEE International Conference on Computer Vision(ICCV).Venice:IEEE,2017:2242-2251.

[24] RAJPUT V.Robustness of different loss functions and their impact on network's learning[J].Neural Networks,2021,16(6):5697-5711.

[25] JOHNSON J,ALAHÍ A,LI F F.Perceptual Losses for Real-Time Style Transfer and Super-Resolution[M].Berlin:Springer International Publishing,2016:694-711.

R2PixGAN:an efficient new method for denoising oracle bone topographies

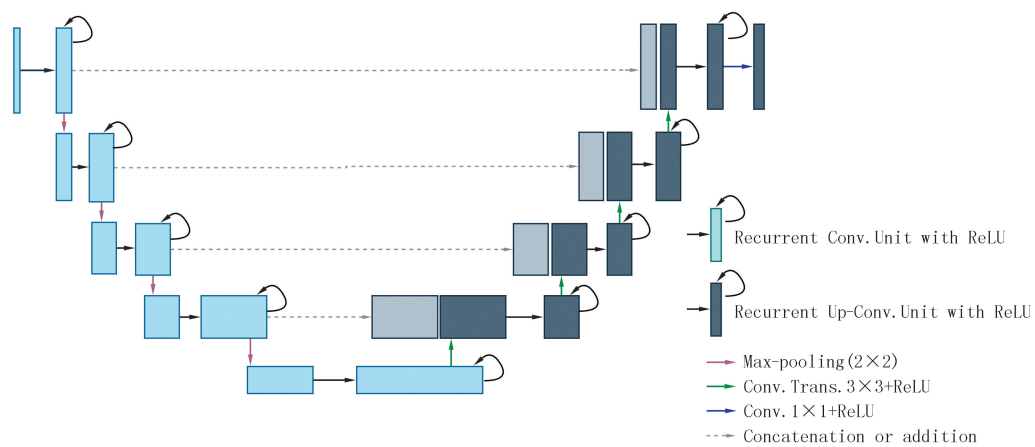
Wang Shibin^{a,b}, Wang Yu^c, Yu Qi^c, Liu Dong^b, Yan Juan^c

(a. School of Computer and Information Engineering; b. Henan Provincial Key Laboratory of Artificial Intelligence and Personalized Learning in Education Henan; c. Oracle Intelligent Computing Laboratory, Henan Normal University, Xinxiang 453007, China)

Abstract: Deep learning has shown strong effects in image denoising. However, existing models often fail to handle oracle bone rubbings which contain complex noise and unique character structures. To solve this, we propose R2PixGAN, a denoising method based on Pix2Pix. It uses an R2U-Net as the generator which keeps the benefits of U-Net and adds RNN to improve image rebuilding and denoising. We also include a perceptual loss to keep key details. Tests show that R2PixGAN gets higher PSNR and SSIM scores than other methods, proving its better performance.

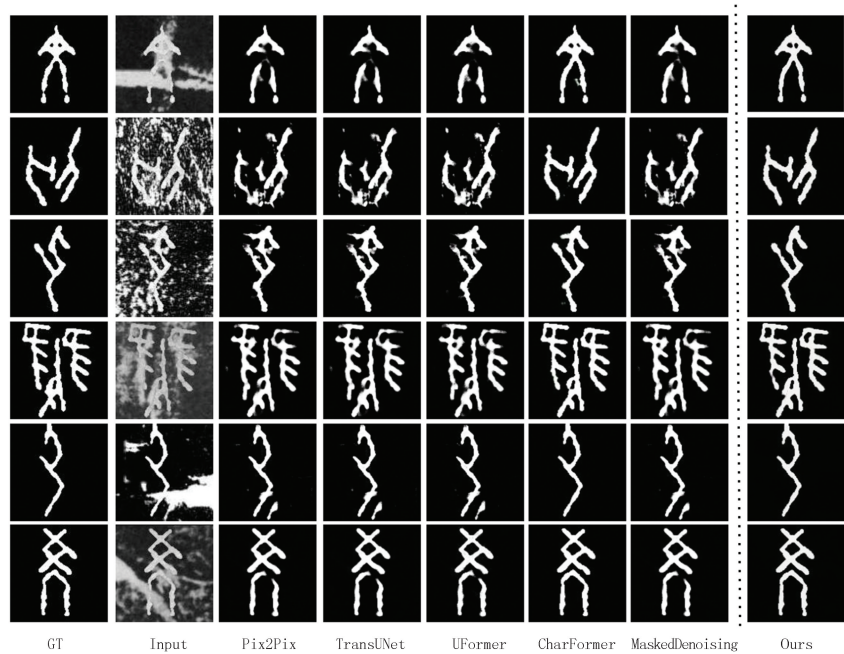
Keywords: image denoising; oracle bone inscriptions; perceptual loss

[责任编辑 陈留院 杨浦]



图S1 R2U-Net生成器的框架图

Fig.S1 Framework diagram of the R2U-Net generator



图S2 不同去噪模型的可视化结果

Fig..S2 Visualization results of different denoising models