

基于轻量级和笔触的超高清图像艺术风格迁移

杨 天,丁一峰,尹淑婷,李长庚

安徽理工大学 计算机科学与工程学院,安徽 淮南 232001

摘要:目的 随着现代技术逐渐发展,图像的分辨率和用户需求日益增加,而有限的 GPU 显存面对超高清分辨率图像样式迁移任务时往往会发生内存溢出,此外现有的超高清图像艺术迁移方法忽视了笔触的影响,通常使用小笔触迁移样式,为此,本文提出了一个结合补丁切割方法的轻量级模型 PCLM(Patch Cutting and Lightweight Model)。方法 为了有效减少内存消耗,PCLM 采用两步措施:首先将超高清大图切割成小补丁集合,将大作业转换为小作业,从根本上解决内存消耗问题,同时采用知识蒸馏压缩模型大小。此外,提出了补丁相关损失 PRL(Patch Relation Loss)来控制样式迁移中的笔触大小和补丁一致性。结果 通过大量定性和定量的实验证明,PCLM 能够在低显存下实现大笔触的超高分辨率风格迁移,在保证低内存占用的同时,在测试集上生成的艺术画作质量指标 SSIM、PSNR、LPIPS 和储存空间(GB)分别达到 0.495 928、11.610 026、0.546 720 和 3.373 454。结论 从理论和实验分析了模型 PCLM 能够有效解决内存消耗问题,生成大笔触下高质量的超高清艺术画作,对后续处理超高分辨率图像任务具有参考价值。

关键词:切割补丁;超高清;模型压缩;大笔触

中图分类号:TP391.4 **文献标识码:**A **doi:**10.16055/j.issn.1672-058X.2026.0004.002

Artistic Style Transfer for Ultra-High-Resolution Images Based on a Lightweight Model with Large Brushstrokes

YANG Tian, DING Yifeng, YIN Shuting, LI Changgeng

School of Computer Science and Engineering, Anhui University of Science and Technology, Huainan 232001, Anhui, China

Abstract: Objective With the continuous advancement of modern technology, image resolution and user demands continue to increase. However, limited GPU memory frequently causes out-of-memory errors during ultra-high-resolution (UHR) style transfer tasks. Furthermore, existing UHR artistic transfer methods often neglect brushstroke characteristics, predominantly relying on small brushstrokes for style transfer. To address these challenges, this paper presents a lightweight framework integrating patch cutting, termed the Patch Cutting and Lightweight Model (PCLM). **Methods** To effectively reduce memory consumption, the proposed model employed a dual optimization strategy. First, UHR images were segmented into patch collections, a process that broke down a large-scale task into smaller subtasks to fundamentally address the memory bottleneck. Simultaneously, knowledge distillation was utilized to compress model parameters. In addition, a patch relation loss (PRL) was proposed to control both brushstroke size and patch consistency during style transfer. **Results** Extensive qualitative and quantitative experiments demonstrated that PCLM successfully realized UHR style transfer with large brushstrokes under constrained memory conditions. While maintaining a minimal memory footprint, the model achieved the following evaluation metrics on the test set: 0.495 928 (SSIM), 11.610 026 dB (PSNR), 0.546 720 (LPIPS), and 3.373 454 GB (GPU memory usage). **Conclusion** Theoretical analysis and

收稿日期:2024-03-29 修回日期:2024-06-19 文章编号:1672-058X(2026)04-0010-08

基金项目:安徽省自然科学基金项目资助(2023085MF220)。

作者简介:杨天(2000—),男,安徽滁州人,硕士研究生,从事机器视觉研究。Email: yt18355042816@163.com.

引用格式:杨天,丁一峰,尹淑婷,等.基于轻量级和笔触的超高清图像艺术风格迁移[J].重庆工商大学学报(自然科学版),2026,43(4):10-17.

Yang Tian, Ding Yifeng, Yin Shuting, et al. Artistic style transfer for ultra-high-resolution images based on a lightweight model with large brushstrokes[J]. Journal of Chongqing Technology and Business University (Natural Science Edition), 2026, 43(4): 10-17.

empirical results confirm that PCLM effectively resolves the memory consumption challenge and generates high-quality UHR artistic paintings featuring large brushstrokes. This study provides a valuable reference for subsequent UHR image processing tasks.

Keywords: patch cutting; ultra-high-resolution; model compression; large brushstroke

图像艺术风格迁移是一个长期的研究课题,旨在通过为照片注入独特的艺术风格来改造照片。在 Gatys 等^[1]的开创性工作之后,样式迁移方法已经能够捕捉艺术画的显著特征并将其应用于内容图像,同时保留原始内容图像的纹理结构。它能够在短时间内完成一个老练艺术家数天的艺术创作,吸引了大量学者的关注。Gatys 等^[1]采用 Gram 矩阵来描述图像的风格特征,内容图像通过对齐自身的 Gram 矩阵和风格图像的矩阵来学习风格图的样式纹理。然而,由于其 Gram 矩阵的学习过程是像素迭代的,生成艺术图像的迭代过程相对较慢。因此,出现了很多优化方法来加速样式迁移^[2-3],提高艺术图像的质量^[4-5],追求艺术风格的真实感^[6-7]以及实现多样化的风格迁移^[8-9]。但这些方法由于网络模型参数较多,只适用于寻常大小分辨率图像的风格迁移,大多不超过 2 K。

随着现代化拍照设备的不断发展,图像分辨率逐渐提高,以往的方法在有限的 GPU 显存上难以完成样式迁移任务。由于输入图像的分辨率不断提高,样式迁移所需的大量显存超过 GPU 的显存容量,导致 GPU 显存溢出和迁移失败。超高清风格迁移^[10-12]方法逐渐出现,其主要障碍是克服大量的显存消耗。超清分辨率风格迁移目前主流的方法有两种:压缩模型和切割补丁操作。压缩模型主要通过知识蒸馏和剪枝的操作缩减网络中的通道数来减少模型的参数量,比如 Wang 等^[10]通过压缩模型来减少内存开销,但它损害了样式迁移的质量,文献^[11]提出了小模型的风格调制策略,目前表现较好。然而,这种方法仍然面临着超分辨率图像内存爆炸的挑战,并且无法从根本上解决内存消耗的问题。而切割补丁操作是将超清大图切割为小补丁集合进行样式迁移,Chen 等^[12]提出了一种用于无约束分辨率样式转换的潜在补丁解决方案,但其方法忽略了补丁与补丁之间关系,导致块与块之间的色调不一致,且仅适用于具有实例规范化结构的模型,缺乏对内容细节的有效保留。

为了解决上述问题,本文提出了一个结合补丁切割方法的轻量级模型 PCLM (Patch Cutting and Lightweight Model),采用两步法来减少内存开销。利用了切割补丁操作将超高清分辨率图像变为小补丁集合依次迁移,

将大作业转换为小作业问题,同时利用知识蒸馏的方法压缩编码器和解码器,减少模型的规模。VGG 网络的前四层作为教师模型中编码器模块,学生模型中的编码器只使用了 VGG 网络的前三层并减少了特征通道数,解码器为编码器的对称结构。此外,本文还提出了一种补丁关联损失(PRL),通过比较补丁之间的相关性来弥补补丁操作的切割痕迹,同时调整补丁大小和损失在模型优化中的权重以实现多尺度的笔触大小。

1 相关工作

1.1 艺术风格迁移

Gatys^[1]的开创性工作开创了神经风格转移(NST)的时代。在他们开创性的工作中,图像的艺术风格是通过从预训练的深度卷积神经网络(DCNN)中提取的特征之间的相关性来表示的,特别是以 Gram 矩阵的形式表示,这在表示图像风格方面非常有效。后续工作在许多方面引入了优化,例如提出了前馈网络以提高效率,尽管其仅限于学习特定的艺术风格。近年来,任意风格转移方法以及提高艺术图像质量的诸多研究已被提出^[2-9],尽管这些方法取得了重大成就,但由于高分辨率图像需要巨大内存和计算资源,它们无法直接应用。近年来,针对这项任务提出了许多方法,但研究人员忽略了超分辨率任务中内存消耗和笔触大小的问题。

1.2 超高清风格迁移

限制超高清风格迁移发展的主要因素是高分辨率图像传输过程中巨大的 GPU 内存消耗。因此,成功的关键在于以较小的内存占用实现风格转换效果。一种方法是训练一个轻量级网络模型 ArtNet^[13],它是通过修剪模型 Google net^[14]而非使用之前的 VGG 网络来构建的。另一种方法是采用知识提炼^[10]来训练一个小型解码器,该解码器以较低的内存消耗提取图像特征。该方法首先训练一个编码器-解码器架构,然后使用经过训练的编码器作为教师来训练一个较小的编码器。最初,超分辨率图像能够在 12 GB GPU 上渲染。尽管内存消耗明显减少,但使用微型编码器可能会在一定程度上导致艺术图像质量下降。最近,Chen 等^[12]提出将输入的高分辨率图像分割成小块,在每个块上执行样式转移,然后使用缩略图实例归一化将这些块拼接

在一起。这为无约束分辨率风格转移提供了一种可能的解决方案。然而,由于确保补丁连续性的操作增加了时间成本,并且它只能应用于具有 InstanceNormal 层的模型。本文无须借助缩略图实例归一化,也不要求网络架构中设置 InstanceNormal 层,有效提升了原有方法的时间效率与适用范围。

1.3 对比学习

对比学习(Contrastive Learning)是一种自监督学习方法,可用于比较不同风格图像之间的差异。通过将不同风格的图像输入模型并使用对比损失函数来比较它们的特征表示,模型可以学习如何区分不同风格的特征。这可以帮助模型在风格转换任务中更好地保留和转换输入图像的风格。在风格迁移中,对比学习也可用于保持一致性。通过使用对比损失函数比较输入

图像和生成图像的特征表示,模型可以学习如何保留输入图像的内容信息,并将其与目标样式相结合以生成一致的结果。针对上述问题,本研究旨在首次探索超分辨率风格转换中的灵活多样性问题,并引入对比学习以生成纹理结构保留更完整、视觉真实感更强的艺术风格图像。

2 方法

如图 1 所示,模型 PCLM 在设备 NVIDIA GeForce RTX 3060 上仅用 3.37 GB 显存即可生成 4 K 超高清分辨率样式迁移结果。

模型结构如图 2 所示,包括切割补丁模块、压缩的小型编码器、样式迁移模块 EHM、对称的解码器结构和 PCLM 损失约束模块。



图 1 4 K 分辨率的图像艺术画作

Fig. 1 Artistic painting with 4 K resolution

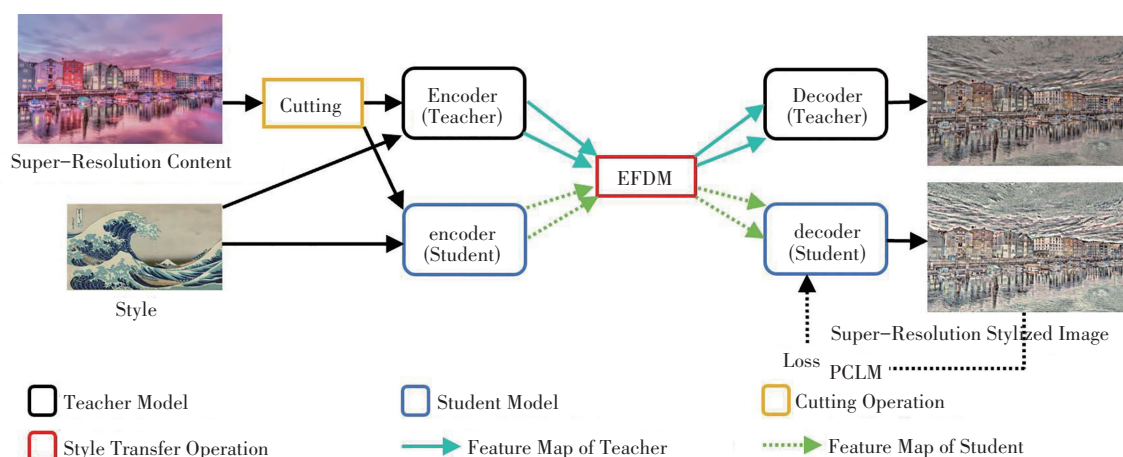


图 2 模型总体框架

Fig. 2 Overall framework of the model

首先利用切割补丁操作将超高清分辨率的图像切割为小补丁集合,通过对小补丁集合依次进行样式迁移,将大作业转换为小作业以减少内存消耗。其中教师编码器采用 VGG-19 网络前四层来提取特征表示,

其结构为对称结构。其次,知识蒸馏被用来压缩模型中的编码器和解码器模块,通过减少自身大小来减少样式迁移过程中的显存占用。将特征映射输入解码器中生成超高清艺术画作,最后通过损失约束 PCLM 来

指导解码器保证生成画作的内容一致性。

2.1 切割模块

对于超高清分辨率的图片,仅使用模型压缩虽能缓解内存消耗,但不能从根本上解决内存溢出的问题。例如 4 K 的内容图,补丁大小是可以手动调节的,模型中默认使用 500×500 大小,首先计算切割后的补丁块数量:

$$N = \text{Math. ceil}(C_w \times C_h + 2 \times S / P_w \times P_h) \quad (1)$$

式(1)中,Math. ceil 为取整函数, C_w 、 C_h 分别为超高清内容图的宽和高, P_w 、 P_h 分别为补丁代下的宽和高, S 为步长大小,使用 32 大小平滑切割痕迹。计算出切割后的数量后,再对超高清内容大图尺寸进行微调,使其

能够被完整切割。

2.2 模型压缩

模型压缩分为两步,首先压缩编码器模块,第二步压缩解码器模块。

(1) 编码模块压缩。如图 3 所示,教师编码模块采用 VGG-19 网络的前四层提取内容特征和风格特征,随着层数增加,通道数逐步增多,但一些通道提取的特征信息冗余。因此,学生编码模块学习使用更少的通道数来代表特征信息,并针对内容特征需要更深层的信息来代表特征结构,内容特征取到 relu4-4,而风格信息只训练到 relu4-1。

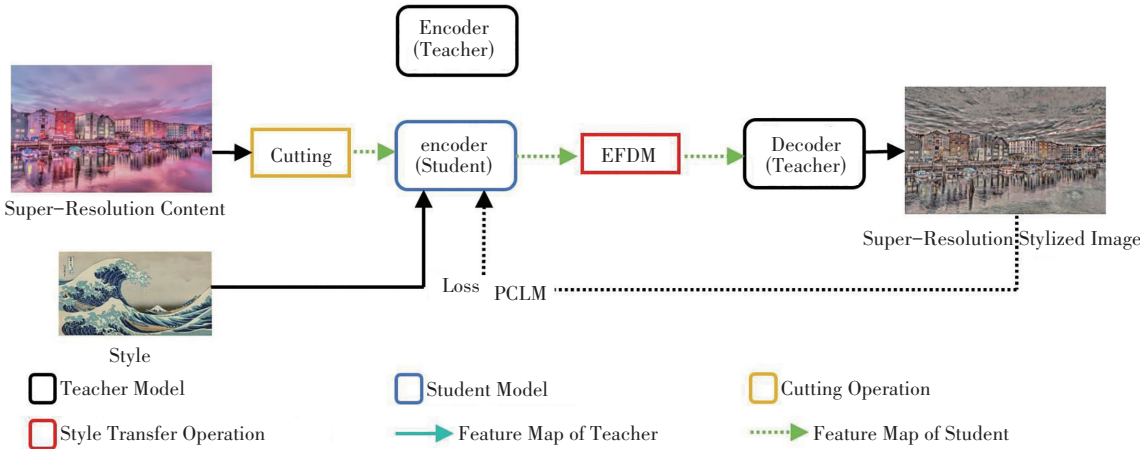


图 3 训练学生编码模块

Fig. 3 Training of the student encoder module

首先固定教师解码器结构,教师编码器换为学生编码器重新进行训练。训练过程根据每层通道计算的参与度,对模型通道中参与较少的通道进行删除,对于没有删除通道的卷积层,依据损失约束模块 PRL 进行参数预训练。

(2) 解码模块压缩。如图 4 所示,在学生解码器模

块学习教师特征提取模块之后,固定学生解码器,并将教师解码器换为学生解码器。学生解码器结构搭建为训练好的学生编码器对称结构,以用于还原图像。训练过程中仍然使用 PCML 损失和 Gram 约束模块来指导学生解码器的参数更新。

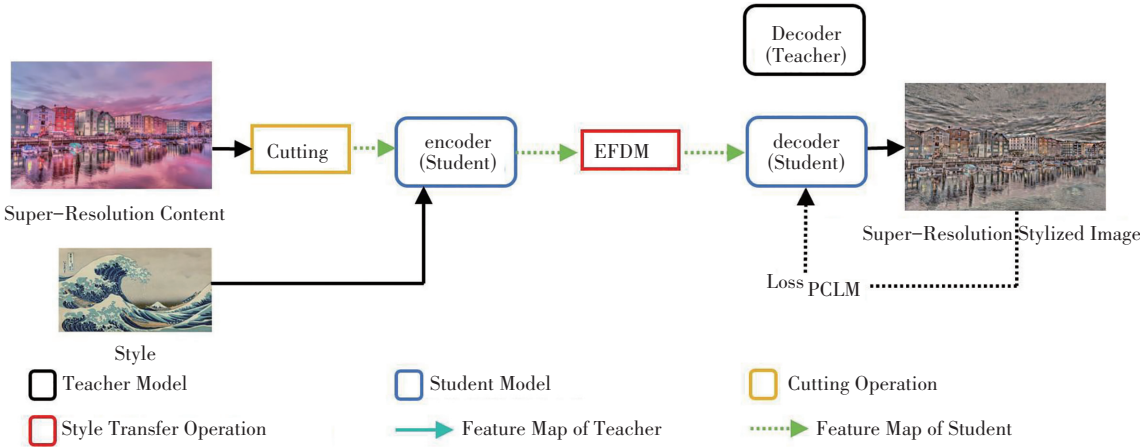


图 4 训练学生解码模块

Fig. 4 Training of the student decoder module

2.3 样式迁移模块 EHM

模型中样式迁移的方法使用 Zhang 等^[15]提出的精确特征分布匹配 EFDM (Exact Feature Distribution Matching)。传统样式迁移方法假设特征分布符合高斯分布,通常将内容图像与风格图像的均值和标准差相匹配。然而,真实数据的特征分布通常比高斯分布复杂得多,仅使用一阶和二阶统计量无法实现精确的分布匹配,而在计算上使用高阶统计量进行分布匹配又不可行。

EFDM 通过精确匹配图像特征的经验累积分布函数(eCDFs)来实现精确特征分布匹配(EFDM),这可以通过在图像特征空间中应用精确直方图匹配(EHM)来实现。

$$\begin{aligned} x: \tau &= (\tau_1, \tau_2, \tau_3, \dots, \tau_5) \\ y: k &= (k_1, k_2, k_3, \dots, k_5) \end{aligned} \quad (2)$$

式(2)中, $\{x_{\tau_i}\}_{i=1}^n$ 和 $\{y_{k_i}\}_{i=1}^n$ 是 x 和 y 值的升序排序。根据式(3)中的定义,使用其第 τ_i -th 个元素 o_{τ_i} 的排序匹配输出为

$$o_{\tau_i} = y_{k_i} \quad (3)$$



图 5 大小笔触下的艺术画作

Fig. 5 Artistic paintings with large and small strokes

3 仿真实验与结果分析

在本节中,对模型进行了定性和定量的比较及消融研究,以验证各组件在模型中的重要性。所有比较方法使用相同的硬件环境和参数配置,均在 NVIDIA GeForce RTX 3060(12 GB)GPU 上进行。

3.1 实验数据集与参数设置

(1) 实验设计。使用 Adam 优化器对网络进行训练,学习率为 0.000 1,批量大小为 16,迭代次数为

2.4 补丁相关损失

提出了一种新的补丁相关损失函数 PRL,用于比较补丁与补丁之间的特征信息,以保持缝合艺术化补丁后生成的超高清艺术画作色调一致性。

$$L_{\text{relation}} = \frac{\exp((S_{p_i})^T(P_i)/N)}{\exp((S_{p_j})^T(P_j)/N) + \exp((S_{p_i})^T(P_i)/N)} \quad (4)$$

式(4)中, S_{p_i} 为第 i 个风格化的补丁, S_{p_j} 为超高清切割补丁集合中的其他补丁, N 为补丁数量。

2.5 笔触大小

在风格迁移中,笔触大小可以影响生成图像的外观和质感,也可以作为一种重要的参数来控制图像的外观和风格。其效果如图 5 所示。通常来说,较大的笔触大小会产生更粗糙、更抽象的效果,而较小的笔触大小则会产生更精细、更细腻的效果。这种效果的差异可以通过在生成过程中调整笔触大小来实现。模型中可通过调节损失 PRL 在约束中的权重来实现大小笔触生成艺术画作。

160 000 次^[16]。使用 MS-COCO^[17] 和 WikiArt^[18] 分别作为内容数据集和风格数据集,将每次迭代的 8 幅图像的小批量输入到网络模型中。所有后续实验均在 NVIDIA GeForce RTX 3060(12 GB)GPU 上进行。

(2) 数据可用性。支持本研究结果的数据集可在 MS-COCO 上公开获取,网址为 <http://images.cocodataset.org/zips/train2014.zip> 和 WikiArt<https://www.kaggle.com/c/painter-2014.zip>。

3.2 定性比较

图 6 显示了生成的艺术画作在质量方面的比较结果。一些现有的风格转换方法,如第 2 列 AdaIN_UR^[19]生成的艺术画作通过整体对齐风格图像的均值和方差忽略了局部样式,第 3 列 LST_UR^[20]的艺术画作无法维护内容结构,导致内容细节部分发生泄露。Collab-Distil^[10]的方法虽然使用较少的内存空间,但会牺牲艺术画的质量。目前先进的方法之一 MicroAST^[11]能够生成较为符合人类感知的艺术画,但局部内容结构仍

有所丢失。本文思想产生的艺术画经过残差链接、判别器、局部风格预测和对比损失等多方面来保护内容结构,如图 3 最右侧,内容图的细节纹理结构得到了有效保存。

此外进行了用户调研,使用 10 张超高清内容图像和 15 张风格图,分别用图 6 中的方法产生 150 张艺术画作让用户选择。让用户根据偏爱率和欺骗率分别选出最喜爱的和最真实的艺术画作,对调研的结果整理如表 1 所示。

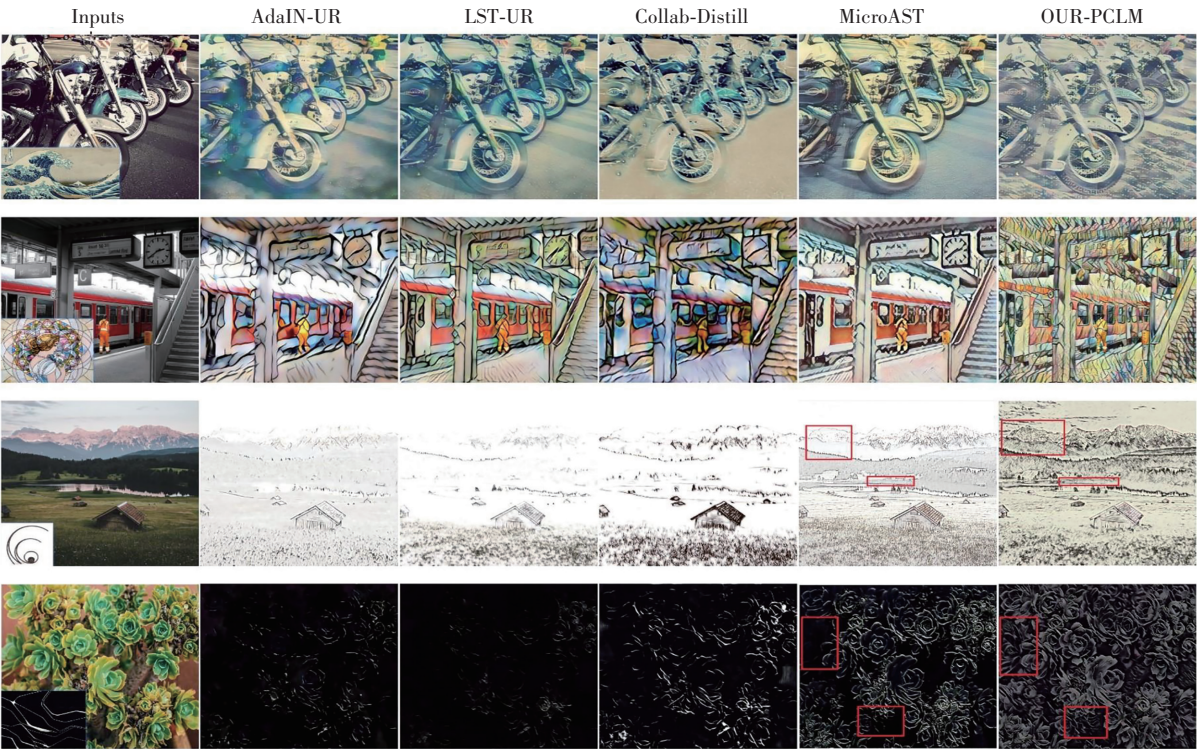


图 6 与当前先进方法的比较

Fig. 6 Comparison with state-of-the-art methods

表 1 用户对不同方法的学习得分 (越高越好)

Table 1 Users' learning scores of different methods (higher is better)

Methods	AdaIN_UR ^[19]	LST_UR ^[20]	Collab-Distil ^[10]	MicroAST ^[11]	OUR-PCLM
Preference ↑	0. 155±0. 23	0. 207±0. 21	0. 180±0. 25	0. 224±0. 22	0. 234±0. 11
Deception ↑	0. 455±0. 22	0. 503±0. 17	0. 399±0. 23	0. 512±0. 23	0. 517±0. 15

3.3 定量比较

除了人工感官分析和调查外,还采用了两种广泛使用的评价指标来验证方法的有效性:峰值信噪比和学习感知图像块相似度。峰值信噪比(Peak Signal-to-Noise Ratio, PSNR)是衡量图像质量的常用指标之一。它用于评估图像经过压缩或处理后与原始图像之间的相似程度,常用于图像处理、图像压缩和图像恢复等领域。学习感知图像块相似度(Learning Perceptual Image Patch Similation, LPIPS)是一种用于衡量图像之间相似度的方

法,其灵感来自人类视觉系统对图像的感知方式。这种方法不仅考虑了图像的低层次特征,如颜色和纹理,还考虑了高层次语义信息,如对象和场景的内容。LPIPS 比较如表 2 所示。较低的 LPIPS 值表明风格化图像和内容图像之间的感知差异较小,这意味着重建图像在感知方面与原始图像更相似。可以看出,模型即使在压缩后仍然可以达到最先进的水平。此外,在同样 4 K 大小的超高清内容图样式迁移内存占用比较中,可以看出 PCLM 模型使用的显存最少。

表 2 不同方法的定量比较结果

Table 2 Quantitative comparison results of different methods

Methods	AdaIN_UR ^[19]	LST_UR ^[20]	Collab-Distil ^[10]	MicroAST ^[11]	OUR-PCLM
SSIM ↑	0.314 596	0.414 525	0.257 689	0.493 021	0.495 928
PSNR ↑	11.121 608	11.317 836	10.721 265	11.612 536	11.610 026
LPIPS ↓	0.572 540	0.531 787	0.620 315	0.526 144	0.546 720
Storage(GB) ↓	4.398 541	5.875 335	7.741 343	3.451 454	3.373 454

3.4 消融实验

为了验证该方法的有效性,对所提出的笔触大小和约束模块 PRL 进行了消融研究,以验证其对整体模型影响。结果表明:它们对模型生成的超高清风格化画像的品质起着重要作用。

从图 7 可以观察到,如果没有笔触控制,第 3 列的风格化图像以小笔触生成具有许多细节冗余信息,整体嘈杂。第 4 列的艺术画作中,由于没有约束模块 PRL 指导模型,出现了局部色调与整体色调不一致,补丁之间差异大的问题,导致补丁缝合起来的超高清分辨率艺术画作

块与块之间差异明显。模型 PCLM 综合考虑了笔触与超高清分辨率内容图补丁之间的差异,从而使得模型输出大笔触下的色调一致的超高清分辨率艺术画作。

此外视觉上的比较不足以说明质量下降,消融模块的定量研究是必要的。如表 3 所示,当模型中缺少笔触控制时,在以小笔触样式进行迁移时,图像迁移了过多的风格信息,导致生成的艺术作品内容泄露,从而质量下降。当模型缺少补丁对比损失 PRL 时,超高清艺术画作的局部之间色调出现不一致,内容信息得不到保护,如表 3 第 4 列所示。

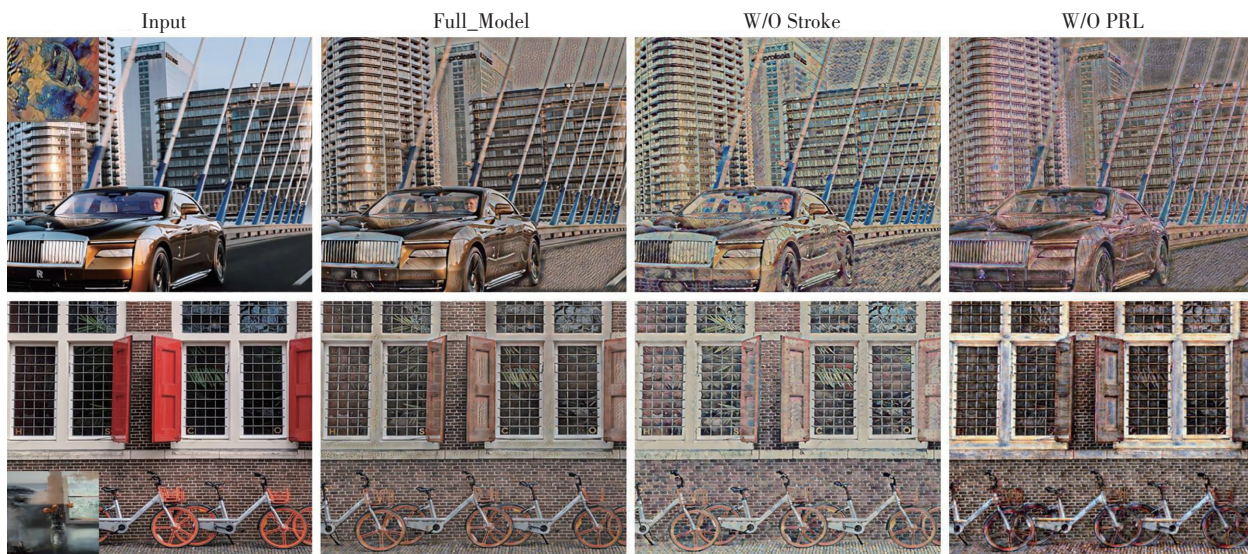


图 7 对各个模块的消融实验结果图

Fig. 7 Ablation experiment results for each module

表 3 消融方法的定量比较结果

Table 3 Quantitative comparison results of ablation methods

Methods	Full_model	W/O Stroke	W/O PRL
SSIM ↑	0.495 928	0.433 721	0.405 688
PSNR ↑	11.610 026	11.085 863	10.857 357
LPIPS ↓	0.546 720	0.591 935	0.615 662

4 结论与展望

针对现代化社会不断发展,超高清分辨率内容图像在进行样式迁移时带来的 GPU 显存挑战,本文提出

了一个结合补丁切割方法的轻量级模型 PCLM。模型采用两步来减少样式迁移中的显存消耗,首先利用切割补丁操作将超高清分辨率内容图像切割成小补丁集合,将大作业转换为小作业任务,从根本上解决了内存开销问题,其次采用知识蒸馏对模型进行压缩,使得压缩后的学生模型能够代替教师模型提取的内容图像和风格图像的特征映射。此外,提出了一个补丁相关损失 PRL 来调节模型中补丁的色调一致性和笔触大小,笔触大小在超高清样式迁移中尤为重要,大笔触样式迁移能使得生成的艺术画作整体简洁,艺术效果更加真实。

虽然切割补丁操作从根本上解决了内存开销,但其带来的小作业堆叠问题还有待解决,时间效率有待提高,未来会继续深入研究。

参考文献(References):

- [1] Gatys L A, Ecker A S, Bethge M. Image style transfer using convolutional neural networks[C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE Press, 2016: 2414–2423.
- [2] Wu X, Hu Z, Sheng L, et al. StyleFormer: real-time arbitrary style transfer via parametric style composition[C]//Proceedings of the IEEE/CVF International Conference on Computer Vision. Piscataway: IEEE Press, 2021: 14598–14607.
- [3] Wang Z, Zhao L, Chen H, et al. Texture reformer: towards fast and universal interactive texture transfer[J]. Proceedings of the AAAI Conference on Artificial Intelligence, 2022, 36(3): 2624–2632.
- [4] Deng Y, Tang F, Dong W, et al. Arbitrary style transfer via multi-adaptation network[C]//Proceedings of the 28th ACM International Conference on Multimedia. New York: ACM, 2020: 2719–2727.
- [5] Park D Y, Lee K H. Arbitrary style transfer with style-attentional networks[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE Press, 2019: 5873–5881.
- [6] Yang S, Hwang H, Ye J C. Zero-shot contrastive loss for text-guided diffusion image style transfer[C]//Proceedings of the IEEE/CVF International Conference on Computer Vision. Piscataway: IEEE Press, 2023: 22816–22825.
- [7] Wang Z, Zhao L, Xing W. StyleDiffusion: controllable disentangled style transfer via diffusion models[C]//Proceedings of the IEEE/CVF International Conference on Computer Vision. Piscataway: IEEE Press, 2023: 7643–7655.
- [8] Yu X, Zhou G. Arbitrary style transfer via content consistency and style consistency[J]. The Visual Computer, 2024, 40(3): 1369–1382.
- [9] Wang Z, Zhao L, Chen H, et al. Diversified arbitrary style transfer via deep feature perturbation[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE Press, 2020: 7786–7795.
- [10] Wang H, Li Y, Wang Y, et al. Collaborative distillation for ultra-resolution universal style transfer[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE Press, 2020: 1857–1866.
- [11] Wang Z, Zhao L, Zuo Z, et al. MicroAST: towards super-fast ultra-resolution arbitrary style transfer[J]. Proceedings of the AAAI Conference on Artificial Intelligence, 2023, 37(3): 2742–2750.
- [12] Chen Z, Wang W, Xie E, et al. Towards ultra-resolution neural style transfer via thumbnail instance normalization[J]. Proceedings of the AAAI Conference on Artificial Intelligence, 2022, 36(1): 393–400.
- [13] Wu K, Tang F, Liu N, et al. Lighting image/video style transfer methods by iterative channel pruning[C]//Proceedings of the ICASSP 2024—2024 IEEE International Conference on Acoustics, Speech and Signal Processing. Piscataway: IEEE Press, 2024: 3800–3804.
- [14] Szegedy C, Liu W, Jia Y, et al. Going deeper with convolutions[C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE Press, 2015: 1–9.
- [15] Zhang Y, Li M, Li R, et al. Exact feature distribution matching for arbitrary style transfer and domain generalization[C]//Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE Press, 2022: 8025–8035.
- [16] Wang X, Liu Z, Gao Y, et al. A near-optimal protocol for the grouping problem in RFID systems[J]. IEEE Transactions on Mobile Computing, 2021, 20(4): 1257–1272.
- [17] Lin T Y, Maire M, Belongie S, et al. MicrosoftCOCO: common objects in context[C]//European Conference on Computer Vision. Cham: Springer, 2014: 740–755.
- [18] Phillips F, Mackintosh B. Wiki Art Gallery, Inc.: a case for critical thinking[J]. Issues in Accounting Education, 2021(26): 593–608.
- [19] Huang X, Belongie S. Arbitrary style transfer in real-time with adaptive instance normalization[C]//Proceedings of the IEEE International Conference on Computer Vision. Piscataway: IEEE Press, 2017: 1510–1519.
- [20] Bharathi A, Sanku R, Sridevi M, et al. Real-time human action prediction using pose estimation with attention-based LSTM network[J]. Signal, Image and Video Processing, 2024, 18(4): 3255–3264.

责任编辑:陈 芳