

基于深度学习的伪装贴图生成技术研究

赵爽¹ 张麟华^{1,2} 陈润峰² 韩嘉琪²

ZHAO Shuang ZHANG Linhua CHEN Runfeng HAN Jiaqi

摘要

在现代作战中,飞机在临时起降点停留时极易暴露于无人机实时目标检测系统,对任务安全构成显著威胁。为应对此挑战,文章提出一种基于深度学习的伪装贴图生成方法。该方法的生成框架采用 GAN 网络,DDPM 完成预训练与初始样本生成,选取均方误差 MSE 作为损失函数,最后通过进化策略的方案动态优化。研究采用高强度轻质复合材料为载体,通过数字印刷技术将生成的高精度迷彩图案印制其上,构建适用于飞机临时降落场景的便携式伪装系统。实验结果表明,采用该策略生成的伪装贴图可使飞机目标在无人机检测中的识别准确率降低 90%,显著提升战场隐蔽性,为临时起降点伪装提供了创新技术方案。

关键词

进化策略;伪装贴图;GAN;DDPM

doi: 10.3969/j.issn.1672-9528.2026.01.005

0 引言

在现代战争中,无人机凭借其出色的机动性和智能化检测技术,已经成为战场上进行目标侦察和信息收集的核心装备。特别是在动态部署和临时任务场景下,例如飞机需要在无预设保障的地点进行临时起降时,其地面停驻状态极易暴露于敌方无人机的监视之下,从而面临严峻的安全威胁。面对敌方高分辨率成像技术和基于深度学习的目标识别算法,传统的伪装技术逐渐暴露出其局限性,难以有效应对新型侦察手段带来的威胁。因此,发展能够有效应对新型智能化侦察威胁的先进伪装策略,对于提升关键装备在复杂战场环境下的生存能力具有迫切的现实意义。

近年来,深度学习领域的飞速发展,尤其是以生成对抗网络(GAN)^[1]为代表的生成模型,为解决复杂数据生成问题开辟了新的途径。GAN 通过构建生成器与判别器之间的动态博弈,能够学习并模拟真实数据的复杂分布,在图像生成、数据增强乃至艺术创作等多个领域取得了令人瞩目的成就。这一技术潜力也引起了军事领域的关注,探索将其应用于设计更逼真、更具迷惑性的伪装图案成为一个新兴的研究方向。

然而,将现有的通用生成模型直接应用于军事伪装设计时,往往面临若干挑战,如生成伪装图案的视觉效果稳定性有待提高、针对特定侦察系统(如无人机视觉检测算法)的隐蔽效能可能不足,以及模型训练和部署对计算资源要求较高等问题^[2]。这些因素在一定程度上限制了先进生成模型在军事伪装领域的直接应用和快速部署。特别是在需要生成对

抗特定检测算法、具备高度稳定性和优异隐蔽性的军事级伪装图案时,急需更为精准和高效的技术方案。

鉴于此,本文针对飞机临时起降点的伪装需求,提出一种基于 GAN 的地面伪装贴图生成方法。核心思路为融合去噪扩散概率模型(DDPM)、均方误差(MSE)监督学习与进化策略,构建可控图像生成方案。以 GAN 作为骨干网络搭建 DDPM 生成框架,借 MSE 损失函数优化去噪流程,引入进化策略动态调控扩散范围、筛选高质量样本,达成目标导向的图像生成,旨在对抗无人机光学检测系统,为提升战场目标隐蔽性与生存能力,探索数码迷彩技术应对无人机侦察威胁的新范式。

1 生成对抗网络技术

生成对抗网络(generative adversarial networks, GAN)最初由 Goodfellow 等^[1]于 2014 年提出,广泛应用于图像生成^[3]、图像超分辨率^[4]、图像风格迁移^[5]等任务。GAN 是一种深度学习框架,由两部分组成:生成器和判别器。生成器负责生成逼真的数据样本,而判别器则负责区分生成样本与真实样本的差异。生成器和判别器通过对抗过程不断优化,使得生成器能够生成越来越真实的样本。随着 GAN 的不断发展,许多不同的变体应运而生,以应对更复杂的生成任务和优化问题。文献[6]提出了一种基于深度卷积神经网络的生成对抗网络架构(DCGAN),通过特定的网络结构约束,提升了生成图像的质量和训练的稳定性。文献[7]引入了 Wasserstein 距离作为损失函数,改善了传统 GAN 训练中的不稳定性问题,并有效缓解了模式崩溃现象。文献[8]提出了条件生成对抗网络(CGAN),通过将条件信息输入生成器和判别器,实现了对生成样本的控制。文献[5]提出了

1. 太原师范学院计算机科学与技术学院 山西太原 030619

2. 太原工业学院计算机工程系 山西太原 030008

一种条件生成对抗网络架构,用于图像到图像的转换任务,如图像修复、图像着色等。文献[9]提出了一种无监督的图像到图像转换方法,通过引入循环一致性损失,实现了源域和目标域之间的映射学习。

尽管生成对抗网络及其变体在图像生成领域取得显著成效,但在处理复杂语义控制、提升生成样本多样性与稳定性等方面仍面临挑战。随着人工智能应用场景对图像生成的可控性与精细化要求不断提高,急需探索更高效、灵活的生成策略。

2 本文方法

本文提出一种融合去噪扩散概率模型、MSE 监督学习与进化策略的可控图像生成方法。该方法以 GAN 为核心生成框架,先通过 DDPM 完成初始样本的预训练生成,再将其输入 GAN 进行对抗性精细化迭代,通过 MSE 损失函数优化去噪过程,并引入进化策略动态调整扩散范围与筛选高质量样本,实现目标导向的图像生成任务,具体如图 1 所示。

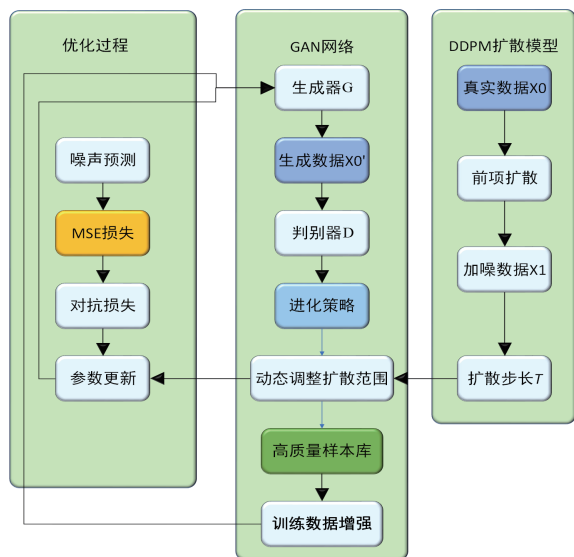


图 1 本文方法网络结构图

图 1 为本文方法的网络结构图,整体流程主要为 DDPM 生成加噪数据,GAN 经生成、判别优化,进化策略评估筛选,构建高质量样本库用于训练数据增强。下面三个小节详细介绍各部分的原理。

2.1 DDPM 生成框架

本文采用去噪扩散概率模型作为初始样本预训练与初始优化模块。DDPM 的工作原理是:在前向扩散过程中,逐步向图像添加高斯噪声,直到图像变为纯粹的随机噪声。然后,模型学习逆转这个过程,即反向过程,通过逐步去除噪声,从纯噪声中重建出清晰的图像。

DDPM 负责动态协调整个过程。噪声添加方法模拟前向扩散,接收原始图像 x_0 和时间步长 t ,生成带噪声的图像 x_t 以及添加的噪声。

$$x_t = \sqrt{\alpha_t}x_0 + \sqrt{1-\alpha_t}\varepsilon \quad (1)$$

式中: ε 是高斯噪声; α_t 是 α 值的累积乘。

模型的主要任务由 UNet 神经网络实现,它在给定时间步长 t 和带噪声图像 x_t 的情况下,预测添加到图像中的噪声 ε 。DDPM 中的 sample 方法利用训练好的 UNet 来生成新图像。从随机噪声开始,然后迭代地使用 UNet 的预测来逐步去噪,直到生成高质量的图像。

2.2 MSE 损失函数

训练 DDPM 模型时,需要衡量模型去噪能力的准确性。选择使用均方误差(MSE)作为损失函数,是一种常见的衡量预测值与真实值之间差异的方法。MSE 损失函数的公式为:

$$MSE = \frac{1}{N} \sum_{i=1}^N (T_i - P_i)^2 \quad (2)$$

式中: N 代表样本的数量; T_i 代表第 i 个样本在加噪过程中实际添加的噪声; P_i 代表模型对第 i 个样本的噪声预测。

MSE 损失函数计算了每个样本的实际噪声与预测噪声之间差异的平方,然后将所有这些平方差加起来,最后取平均值。通过对差异进行平方处理,可以确保损失值始终为正,并且较大的误差会受到更严厉的惩罚,这有助于模型更关注那些预测偏差较大的情况。通过不断最小化这个损失,模型就能逐渐学会如何准确地预测噪声,从而提升其在去噪过程中的表现。

在训练过程中,会随机选择一些图像,并为它们添加一定量的噪声。这样就有了“带噪声的图像”和“实际添加的噪声”两部分信息。接着,模型会尝试根据带噪声的图像来预测它认为的噪声。MSE 损失函数的作用就是比较模型预测的噪声与实际添加的噪声之间的差异。如果模型的预测与实际噪声非常接近,那么损失值就很小,说明模型表现良好;如果差异很大,损失值就会非常高,提示模型需要进一步学习和调整。通过不断地最小化这个损失,模型就能逐渐学会如何准确地预测噪声,从而提升其在去噪过程中的表现。

2.3 进化策略动态优化

为了让生成的图像质量更高,在 DDPM 生成初始样本后,以 GAN 为核心执行精细化优化,同时引入进化策略动态调控 GAN 的生成参数,筛选高质量伪装样本。

本文设计了一个基于目标检测性能的评估函数,该函数通过计算生成图像在目标检测任务中的 mAP 得分来量化其质量。模型会生成一批图像,然后由评估函数进行打分。接着,系统会根据这些图像的平均得分来动态调整扩散模型的一个重要参数——噪声范围。如果生成的图像得分普遍较低,这表明模型可能需要更大的探索空间,系统就会扩大噪声范围,鼓励模型去尝试更多样化的生成方式,就像让它在更广阔的

画布上自由创作。反之,如果得分很高,说明模型已经掌握了生成高质量图像的方法,系统就会缩小噪声范围,让模型进行更精细的调整和优化,就像让它专注于细节的打磨。

本文还设置了一个经验回放池,用来收集得分较高的“高质量”图像。这些被认可的图像会被保存起来,并在后续的训练中与新数据一起使用。这就像一个学习榜样,模型会从中学习成功经验,从而进一步提高生成高质量图像的能力。这种动态调整和经验积累的机制,使得模型能够像生物进化一样,不断地迭代和优化,最终生成出更令人满意的图像。

3 实验结果和分析

3.1 实验配置

训练服务器配备了 1 颗 NVIDIA GeForce RTX 3090 GPU, CPU 为 Intel Xeon(R) Silver 4210R 处理器。软件配置包括:操作系统为 Ubuntu 20.04.6 LTS、编程语言为 Python 3.8、深度学习框架为 PyTorch 2.1.1 以及 CUDA 版本为 12.0。

3.2 实验数据集

本文采用文献 [9] 中的公开数据集 MAR20 作为训练数据集,同时使用 3D 建模软件 blender 进行建模对生成的结果在仿真条件下进行验证。

MAR20 数据集是目前规模最大的遥感图像军用飞机目标识别数据集。该数据集由 3 842 张遥感图像组成,涵盖 20 种军用飞机型号,共计 22 341 个目标实例。这些图像均通过 Google Earth 从全球 60 个军用机场收集,具有较高的多样性和挑战性。

测试集使用的为自己创建的 3D 建模场景,分为两个不同的场景,每个场景放置若干飞机。在此测试集中,模拟了 3D 仿真的飞机临时停靠场景来测试本文生成图案的效果。

3.3 实验评估指标

为了全面评估生成模型的伪装效果,实验选取了 $mAP@50-95$ 作为评估指标。 $mAP@50-95$ 是一个全面的目标检测精度评估指标,它综合考虑了召回率 (Recall)、精确率 (Precision) 和置信度。

召回率用于衡量正样本被正确识别的概率,其计算公式为:

$$\text{Recall} = \frac{TP}{TP + FN} \quad (3)$$

式中: TP 表示被准确判定为正类的样本数量; FN 表示那些实际上是正类样本,但被错误地判断为负类的样本数量。

精确率反映在所有被识别出来的样本中,正样本所占的比例。其计算公式为:

$$\text{Precision} = \frac{TP}{TP + FN} \quad (4)$$

平均精度用于评估在不同召回率水平下,精确率的平均表现情况。其计算公式为:

$$AP = \int_0^1 P(R) dR \quad (5)$$

各类别平均精度,在所有类别上分别计算平均精度后,再对 AP 值求平均值。其计算公式为:

$$mAP = \frac{1}{N} \sum_{n \in N} AP(n) \quad (6)$$

式中: N 表示类别的总数; $AP(n)$ 表示第 n 个类别的平均精度。

$$mAP@50-95 = \frac{1}{10} \sum_{IoU=0.5}^{IoU=0.95} mAP(IoU) \quad (7)$$

通过公式计算 $mAP@50-95$,其设定了一系列交并比 (IoU) 阈值,从 0.50 逐步递增至 0.95,每次增加 0.05, $mAP@50-95$ 指的是在这组 IoU 阈值下,分别计算各个类别在不同阈值对应的平均精度 (AP),并取平均值。 $mAP@50-95$ 能够更全面地反映模型在不同类别和多个阈值下的整体性能,因此本文使用此参数来表示检测精度。

$$ACC_r = \frac{ACC_a - ACC_b}{ACC_a} \times 100\% \quad (8)$$

式中: ACC_a 为初始识别准确率,即未做攻击时的识别率; ACC_b 为最终识别准确率,即攻击后的识别率; ACC_r 为最终降低的识别准确率。通过公式降低识别准确率。

3.4 对比实验

为了更全面地验证本文方法的效果,将其与 GAN 和 CGAN 两个主流的代表性模型进行了对比实验。具体结果如表 1 所示。

表 1 对比实验结果

实验模型	mAP_{clean}	mAP_{adv}	$ACC_r/\%$
GAN	0.52	0.17	67
CGAN	0.52	0.14	73
Ours	0.52	0.05	90

表 1 中, mAP_{clean} 为无攻击时检测得到的 $mAP@50-95$ 值, mAP_{adv} 为特定模型作为攻击模型攻击后检测得到的 $mAP@50-95$ 值。可以看出, GAN 模型作为攻击模型时, mAP_{adv} 值为 0.17, 识别准确率降低了 67%, CGAN 模型作为攻击模型时, mAP_{adv} 为 0.14, 识别准确率降低了 73%, 而使用本文方法作为攻击模型时, mAP_{adv} 仅为 0.05, 识别准确率降低了 90%。

3.5 可视化结果与分析

为了评估本文方法所得飞机伪装贴图的效果,对目标检测模型在受攻击前后的检测结果进行了可视化对比实验。使用 YOLOv3 模型在自建的 3D 仿真环境中开展测试,将两个

场景下原始图像和添加对抗伪装后图像的检测结果进行可视化呈现,具体结果如图2所示。

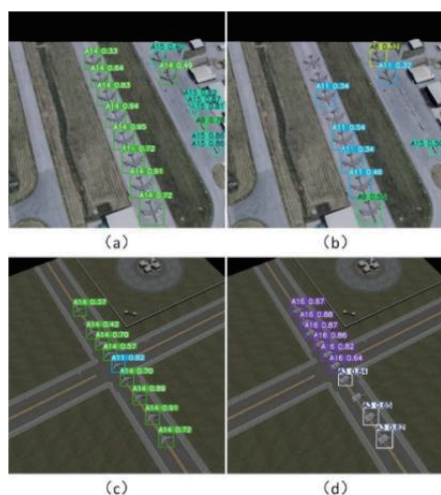


图2 可视化示例图

图2中,(a)和(c)为两个场景未做攻击检测图像,YOLOv3模型可准确识别出九架飞机目标,清晰输出目标边界框与类别置信度,仅1架飞机识别错误,展现模型原始识别能力;(b)和(d)为将飞机放置在部分攻击贴图上方,可以明显看出,攻击效果显著,图(b)中,有3架飞机未被检测到,而七架飞机检测结果错误,图(d)中,有1架飞机未被检测到,九架飞机检测结果错误。直观体现飞机伪装贴图可有效降低模型对特定目标的识别精度,造成漏检、误检,验证了静态伪装手段对目标检测模型的对抗干扰效果。

4 结论

针对现有方法生成伪装图案的视觉效果稳定性较低、在特定侦察系统下的隐蔽效能不足等问题,本文以GAN作为骨干网络构建DDPM生成框架,通过MSE损失函数优化去噪过程,通过进化策略动态调整扩散范围与筛选高质量样本,最终生成了视觉效果稳定的伪装图案,并且可以在YOLO检测系统下达到很好的隐蔽效果。实验结果表明,本文方法优于GAN和CGAN等主流生成方法,且可以在多个场景下拥有良好的攻击效果。

参考文献:

- [1]GOODFELLOW I J, POUGET-ABADIE J, MIRZA M, et al. Generative adversarial nets[C]//NIPS'14: Proceedings of the 28th International Conference on Neural Information Processing Systems. New York: ACM, 2014: 2672-2680.
- [2]李春彦,王珍,罗毅,等.迷彩伪装中基于背景的轮廓生成技术[J].信息记录材料,2019,20(12):74-75.
- [3]RADFORD A, METZ L, CHINTALA S. Unsupervised representation learning with deep convolutional generative

adversarial networks[EB/OL].(2016-01-07)[2025-12-02].
https://doi.org/10.48550/arXiv.1511.06434.

- [4]LEDIG C, THEIS L, HUSZÁR F, et al. Photo-realistic single image super-resolution using a generative adversarial network[C]// 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). Piscataway: IEEE, 2017: 4681-4690.
- [5]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). Piscataway: IEEE, 2017[2025-12-02].
https://ieeexplore.ieee.org/document/8237506.DOI: 10.1109/ICCV.2017.244.
- [6]ARJOVSKY M, CHINTALA S, BOTTOU L. Wasserstein generative adversarial networks[C]//ICML'17: Proceedings of the 34th International Conference on Machine Learning. New York: ACM, 2017: 214-223.
- [7]MIRZA M, OSINDERO S. Conditional generative adversarial nets[EB/OL].(2014-11-06)[2025-12-02].https://doi.org/10.48550/arXiv.1411.1784.
- [8]ISOLA P, ZHU J Y, ZHOU T H, et al. Image-to-image translation with conditional adversarial networks[C]//Proceedings of the IEEE conference on computer vision and pattern recognition. Piscataway: IEEE, 2017: 1125-1134.
- [9]YU W Q, CHENG G, WANG M J, et al. MAR20: a benchmark for military aircraft recognition in remote sensing images[J]. Journal of remote sensing, 2023, 27(12): 2688-2696.

【作者简介】

赵爽(2000—),男,山西文水人,硕士研究生,研究方向:机器视觉, email: m18403587717@163.com。

张麟华(1982—),男,山西阳曲人,硕士,教授,硕士生导师,研究方向:机器视觉。

陈润峰(2004—),男,辽宁大连人,本科,研究方向:机器视觉。

韩嘉琪(2004—),女,山西太原人,本科,研究方向:机器视觉。

(收稿日期:2025-07-04 修回日期:2025-12-31)