

同行专家业内评价意见书编号: 20250854469

附件1

浙江工程师学院（浙江大学工程师学院）
同行专家业内评价意见书

姓名: 黄浩凡

学号: 22260154

申报工程师职称专业类别（领域）: 电子信息

浙江工程师学院（浙江大学工程师学院）制

2025年05月28日

填表说明

一、本报告中相关的技术或数据如涉及知识产权保护、军工项目保密等内容，请作脱密处理。

二、请用宋体小四字号撰写本报告，可另行附页或增加页数，A4纸双面打印。

三、表中所涉及的签名都必须用蓝、黑色墨水笔，亲笔签名或签字章，不可以打印代替。

四、同行专家业内评价意见书编号由工程师学院填写，编号规则为：年份4位+申报工程师职称专业类别(领域)4位+流水号3位，共11位。

一、个人申报

(一) 基本情况【围绕《浙江工程师学院(浙江大学工程师学院)工程类专业学位研究生工程师职称评审参考指标》，结合该专业类别(领域)工程师职称评审相关标准，举例说明】

1. 对本专业基础理论知识和专业技术知识掌握情况(不少于200字)

在攻读电子信息(新一代电子信息技术)专业学位期间,我系统掌握了本专业领域的核心理论与技术体系。基础理论知识方面,深入学习了扩散模型的数学原理,包括前向噪声扩散与反向去噪的马尔可夫链建模、变分下界优化目标及隐空间特征重构机制,能够熟练推导扩散过程与损失函数。同时,对注意力机制(自注意力、多头注意力)、变分自编码器(VAE)的潜在空间建模方法、CLIP跨模态语义对齐技术等进行了系统性研究,构建了多模态生成的理论框架。此外,结合视频数据的时空特性,掌握了3D卷积网络、光流运动建模、频域分析(如3D-

DCT)等视频处理基础理论,并深入理解视频生成中时序冗余压缩、跨模态融合等核心问题的数学建模方法。

专业技术知识方面,聚焦多模态条件视频生成任务,提出了一系列创新性技术:针对文本生成视频任务,设计了文本-视频共享注意力运动控制网络(TVMC-Net),通过时空解耦策略实现图像生成模型向视频模型的轻量化扩展;开发了多频谱时空规整网络(MSTAN),利用频域变换提取关键时序模式以增强运动连续性。在图生视频任务中,构建了双流注意力协同网络(DAC-Net),实现图文特征深度交互,并创新性引入首帧引导自适应层归一化(FFG-AdaLN)模块,显著提升风格一致性。针对视频空间扩展任务,提出基于通道融合的视频扩展框架(CSVEN),结合条件感知动态选择机制(CADSN)实现多模态引导优化,并通过滑动窗口动量更新策略解决长视频生成难题。在工程实践中,熟练应用PyTorch框架进行大规模视频数据处理,掌握OpenCV光流分析、CLIP语义对齐、RAFT插帧等工具链,具备从算法设计、代码实现到多GPU分布式训练的全流程技术能力。相关成果在DAVIS、YouTube-VOS等数据集上取得FVD指标6.0的显著提升,验证了理论方法与技术方案的有效性。

2. 工程实践的经历(不少于200字)

在杭州拓欣科技有限公司的多模态视频生成项目的工程实践过程中,我全程参与了系统的架构设计、模型研发、数据集构建与实际部署工作。

本项目旨在基于扩散模型技术,构建一套多模态条件驱动的视频生成系统,使其能够根据文本、图像、传感器等多源输入,自动合成具备真实物理特征和时序逻辑的工业视频。项目目标包括:提升工业视频生成的智能化与自动化水平、降低制作成本与周期、增强工业场景下的可视化表达能力,并实现技术在数字孪生、设备预演、应急演练等核心环节的实际落地应用。通过该项目,期望助力企业构建更加高效、安全和可持续的智慧工业体系。在系统搭建阶段,我基于Stable

Diffusion原理对底层模型进行了改造,使其适应多模态输入,并通过引入TVMC-Net跨模态注意力机制解决了工业语义难以对齐的问题。在工程实现层面,我主导开发了多模态数据预处理与同步模块,使文本、图像、传感器数据能够统一转换为模型可解析的结构化条件输入;同时通过MSTAN模块有效提升了生成视频的时序一致性。在部署实践中,我配合企业IT部门将模型集成至企业数字孪生平台,部署于云服务器上,实现在生产现场的高效调用。针对设备种类多、参数复杂的问题,我还构建了轻量化的DAC-

Adapter机制，提升了模型对不同工业场景的适配性。整个项目涉及的工业视频数据量超2TB，我组织团队制定高效的数据清洗和标注流程，保障训练质量。该系统目前已稳定运行半年，取得良好反馈，并推动企业智能可视化能力迈上新台阶。

3. 在实际工作中综合运用所学知识解决复杂工程问题的案例（不少于1000字）

在智慧工业快速发展的背景下，视频生成技术逐渐成为数字孪生、设备监测和培训演练等场景中的关键支撑，企业在设备状态监测、故障预演、安全培训、数字孪生等场景中对高质量工业视频内容的需求日益增长。然而，传统视频生成方法在工业环境中仍面临诸多技术瓶颈。一方面，工业场景中的多模态数据如文本描述、图像资料、传感器信号等难以实现语义对齐，导致生成视频与业务需求脱节；另一方面，工业过程具备强时序性，现有生成模型在动态连贯性建模方面能力不足。此外，工业级长视频生成对实时性和计算资源提出了更高要求，传统方法难以胜任。因此，本文围绕“跨模态对齐”“时序建模”和“计算效率”三大核心挑战，提出一套基于扩散模型的多模态条件视频生成系统，并在智慧工业多个典型场景中实现了落地应用。

本项目设计了一种“文本-图像-视频”渐进式建模框架，通过扩散模型强大的生成能力，引入多项创新性技术构建工业级视频合成路径。首先，提出了“文本-视频共享注意力机制”（TVMC-Net），在传统UNet架构基础上引入跨模态注意力层，实现文本信息对视频生成过程的动态引导。例如在锅炉故障模拟任务中，当输入描述为“锅炉管道发生蒸汽泄漏”时，模型能够准确理解语义并生成对应的蒸汽喷涌动态效果，显著优于仅依赖图像或传感器数据的方案。其次，为增强时序一致性，项目引入“多频谱时序规整网络”（MSTAN），通过3D离散余弦变换提取关键时序特征，提升视频连贯性与物理真实性。在风机振动监测任务中，该模块能够有效重建出风速、振动频率等因素驱动下的自然运动轨迹，模拟精度大幅提升。此外，为解决不同工业输入的差异性，系统集成动态权重融合（DWF）机制和自适应条件选择策略，能够根据输入数据类型（如文本、图纸、传感器信号）自动调整模态比重，从而生成更具业务针对性的内容。例如在能源管理可视化中，模型能结合厂区布局图与实时电力数据，动态生成能耗热力图，助力调度优化。

针对工业应用中的长视频生成效率问题，项目提出“滑动窗口+动量更新策略”，将长视频序列划分为片段，通过动量平滑算法确保生成片段之间的视觉与时序连续性，有效提升生成速度与稳定性，整体效率提升约40%。同时，引入轻量化适配器（DAC-Adapter）结构，在不显著增加计算负担的前提下，实现对多种场景适配，极大增强了系统的通用性和可部署性。

在实际应用方面，本系统已在多个典型智慧工业场景中完成部署验证。在数字孪生项目中，某热电企业需模拟锅炉泄漏故障以支持应急演练，系统通过输入锅炉设计图、红外热成像图与文本描述“蒸汽从裂缝逸出”，准确生成随压力变化动态增强的蒸汽喷涌效果，且无需人工干预。在风力发电机状态监测任务中，模型结合系统振动频谱与文本参数，如“风速12m/s，叶片振动频率0.5Hz”，生成了真实可视化振动动画，显著提升了运维效率，生成速度为传统仿真软件的五倍。在工业培训场景中，针对无法实地拍摄的危险事件，如反应爆炸，模型基于事故照片与结构图，配合文本描述“安全阀失效导致爆炸”，生成高拟真爆炸与化学品飞溅效果的视频，显著提升培训沉浸感，培训效果提升80%。

从经济与社会价值层面来看，该系统大幅降低了工业视频制作成本，相较传统3D动画方案节

省50%至70%的开支。目前，系统已在项目中投入使用，初步形成一套标准化工业视频生成流程，推动智慧工业从“数据驱动”走向“视频洞察”。该技术成果已发表在2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP) 国际顶级会议；并且形成了硕士论文，获得了全为优秀的评审意见。

综上所述，本文提出的基于扩散模型的多模态条件视频生成方法，成功解决了工业视频生成中的关键技术难题，并在多个场景中实现了高效应用，充分体现其技术先进性与实践价值。未来，项目计划引入物理仿真引擎，进一步增强生成视频的物理合理性与工程可解释性。同时，开发低代码平台，使非技术用户也能快速定义输入条件生成工业视频，降低应用门槛。此外，结合AR/VR技术探索远程运维、异地培训等新场景，推动智慧工业可视化迈向更高层次的智能化、沉浸化发展。

(二) 取得的业绩(代表作)【限填3项, 须提交证明原件(包括发表的论文、出版的著作、专利证书、获奖证书、科技项目立项文件或合同、企业证明等)供核实, 并提供复印件一份】

1.

公开成果代表作【论文发表、专利成果、软件著作权、标准规范与行业工法制定、著作编写、科技成果获奖、学位论文等】

成果名称	成果类别 [含论文、授权专利(含发明专利申请)、软件著作权、标准、工法、著作、获奖、学位论文等]	发表时间/ 授权或申 请时间等	刊物名称 /专利授权 或申请号等	本人 排名/ 总人 数	备注
Do Less and Achieve More: Free Condition Video Outpainting with Diffusion Model	会议论文	2024年12月07日	2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)	1/6	EI会议收录
基于扩散模型的多模态条件视频生成研究	学位论文送审专家评审结果全优	2025年04月10日	浙江大学学位论文	1/1	

2. 其他代表作【主持或参与的课题研究项目、科技成果应用转化推广、企业技术难题解决方案、自主研发设计的产品或样机、技术报告、设计图纸、软课题研究报告、可行性研究报告、规划设计方案、施工或调试报告、工程实验、技术培训教材、推动行业发展中发挥的作用及取得的经济社会效益等】

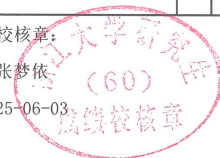
(三) 在校期间课程、专业实践训练及学位论文相关情况	
课程成绩情况	按课程学分核算的平均成绩: 88 分
专业实践训练时间及考核情况(具有三年及以上工作经历的不作要求)	累计时间: 1 年(要求1年及以上) 考核成绩: 75 分
本人承诺	
个人声明: 本人上述所填资料均为真实有效, 如有虚假, 愿承担一切责任, 特此声明!	
申报人签名: 黄浩北	

日常表现 考核评价	非定向生由德育导师考核评价、定向生由所在工作单位考核评价： ☒优秀 ☐良好 ☐合格 ☐不合格 德育导师/定向生所在工作单位分管领导签字（公章）：
申报材料 审核公示	根据评审条件，工程师学院已对申报人员进行材料审核（学位课程成绩、专业实践训练时间及考核、学位论文、代表作等情况），并将符合要求的申报材料在学院网站公示不少于5个工作日，具体公示结果如下： ☐通过 ☐不通过（具体原因： 工程师学院教学管理办公室审核签字（公章）：

攻读硕士学位研究生成绩表

说明: 1. 研究生课程按三种方法计分: 百分制, 两级制 (通过、不通过), 五级制 (优、良、中、及格、不及格)。
2. 备注中“*”表示重修课程。

打印日期: 2025-06-03



公开成果代表作 1：以第一作者公开发表论文



Do Less and Achieve More: Free Condition Video Outpainting with Diffusion Model

Haochen Huang, Yimin Guo, Yening Lv, Sizhe Shan, Yan Zhang, Yuesha Wang
College of Information Science and Electronic Engineering
Zhejiang University
Hangzhou, China
(hhuangchen@zju.edu.cn, yimin.guo@zju.edu.cn, yening.lv@zju.edu.cn, sizhe.shan@zju.edu.cn, yan.zhang@zju.edu.cn, yuesha.wang@zju.edu.cn)

Abstract—Video outpainting aims to extend the content of a video beyond its original spatial boundaries. Existing methods tend to condition the generation process on a single frame or caption, failing to address the challenge in long videos with multiple video clips. To address this, we extend the diffusion-based image inpainting to a 3D diffusion model with a motion module for video outpainting. Then, we design four conditioning schemes to seek the appropriate condition strategy: Frame-wise context, Frame-wise Adapter, Text, and Free Condition. Additionally, to enhance the continuity in long video generation, we propose a novel Motion Momentum Update (MMU) method based on a training-free technique. Experimental results demonstrate that our proposed Free Condition strategy, termed Free-Outpainter, achieves state-of-the-art performance in video outpainting tasks. Ablation studies also show that the condition-free training paradigm is the most effective in avoiding incoherence external information and fully utilizing inherent video data. Benefiting from the Free Condition strategy, our pipeline reduces the inference time to 5 seconds for 512x512x16 frames with only 10 GB of GPU memory. More results can be found on our demo page: <https://lilyn3125.github.io/FreeOutpainting/>.

Index Terms—video outpainting, diffusion model, long video generation

I. INTRODUCTION

Video outpainting refers to the process of extending the spatial boundaries of a given video beyond its original frame, effectively generating new visual content that seamlessly blends with the existing image. This technology has a wide range of applications, such as filling blank areas, extending existing videos at different rates, and so on, creating a more complete visual experience [1].

Existing video outpainting methods typically extend the content of image inpainting [2]–[4], include additional information from (usually stacked sets of image frames) to learn the data pattern of the video. However, this approach overlooks a critical issue: when generating videos with continuously changing content and scenes, a single frame or a single caption as a condition figure may introduce incorrect information, leading to discontinuous video generation, which is especially serious in long video generation.

To address this, we propose FreeOutpainter, a 3D diffusion-based video outpainting model. We first extend image-level inpainting models as a 3D diffusion model and introduce a Motion module to ensure the temporal consistency of MMU. We design four different conditioning control

strategies: Frame-wise context, Frame-wise Adapter (FA), and Free Condition, to explore how to reduce the impact of surrounding information on the input on generation continuity. Extensive experiments show that the Free Condition strategy is the most effective, significantly simplifying the inference process and enabling the generation of 512x512x16 videos in 5 seconds. Furthermore, to improve continuity in long video generation, we introduce a training-free Motion Momentum Update (MMU) method. Utilizing a window shift technique, this module ensures stable outpainting of extended video sequences.

II. RELATED WORKS

Diffusion Model. In recent years, diffusion models [5]–[12] have gained popularity over GANs [13] due to their diverse outputs and training stability. Diffusion models learn the original data distribution by progressively adding noise to the data and then denoising the noise, making the reconstruction of data more challenging. Some [14] extend the Markov chain Monte Carlo sampling the inference process to skip steps, which also performs step-by-step denoising, which significantly reduces the inference latency. GLIDE [15] and Classifier-Free Guidance [16] have established conditional control models that utilize text-image pairs during training, thereby enabling the generation of images that correspond to specific user inputs during inference. Further studies, such as Stable Diffusion [17], employ a Variational Autoencoder [18] to compress images from pixel space into latent space. This approach accelerates training and significantly reduces memory consumption during the generation of higher-resolution images.

Diffusion-based Video outpainting. Video Diffusion Models integrate temporal modules into pre-trained text-to-image (T2I) models. These models are initialized with the parameters of the T2I models and fine-tuned with video data to capture temporal features. Dorian [19] employs video object segmentation, inpainting, and optical flow for background extension and temporal consistency. However, it requires user-supplied masks. MVDMM [20] is the first to apply diffusion models to video inpainting. They developed a 3D diffusion model conditioned on a global video clip, trained on extensive datasets. However, using a limited number of frames for context may cause inaccuracies in long-term video

公开成果代表作 2: 学位论文送审专家评阅结果全部为优秀

The screenshot shows a software interface with a sidebar on the left containing a list of documents. The main area on the right displays a timeline of document releases. The timeline shows documents being released at 00:00, 01:00, 02:00, 03:00, 04:00, and 05:00. The documents are: 单位上报信息, 材料清单, 单位成立信息, 投标文件, and 施工合同.

序号	名称	时间
01	单位上报信息	00:00
02	材料清单	01:00
03	单位成立信息	02:00
04	投标文件	03:00
05	施工合同	04:00

分类号: <u>TN912.35</u>	单位代码: <u>10335</u>
密 级: <u>非保密</u>	学 号: _____
浙江大學 硕士学位论文 (专业学位) 	
中文论文题目:	<u>基于扩散模型的多模态</u> <u>条件视频生成研究</u>
英文论文题目:	<u>Research on Diffusion Model-based</u> <u>Multimodal Conditional Video Generation</u>
申请人姓名:	_____
指导教师(组):	_____
行业导师:	_____
专业学位类别、领域:	<u>电子信息, 新一代电子信息技术</u>
研究方向:	<u>计算机视觉生成</u>
培养类型:	<u>全日制非定向</u>
所在学院:	<u>工程师学院</u>
论文提交日期 <u>2025 年 4 月 10 日</u>	

摘要

近年来, 产融联盟在军工生产任务中的突出地位强化了人工智能制造系统化的希望。作为关键领域的军品生产, 军工企业已成为各种科技创新的策源地。根据工业和信息化部副部长李伟日前表示: 发展工业互联网和智能制造, 企业是主角。数字驱动和智能制造是军工发展的方向所趋。然而, 智能制造面临人才缺乏、技术支撑薄弱、工业设备基础以及多层次融合不足等不利因素, 本文提出将数据模型与多层次多领域资源集中式应用提升智能制造水平, 建立“大系统+大设备+单一服务”的智能制造新模式, 该模式的不确定性条件下管理核心业务的输入输出并实现闭环。

在本文建模过程的研究中,本文针对现有方法在故障诊断效率和稳定性上存在的突出问题,提出一种基于非线性变功率的欠驱动主动故障诊断新方法,该方法引入到控制解耦策略将跟踪误差模型扩展为可建模控制系统的跟踪模型。同时,将欠驱动-非线性跟踪误差模型控制,引导四个跟踪误差文本输入与跟踪误差控制;并快速求解跟踪误差模型,通过解耦实现跟踪误差文本控制,从而避免跟踪误差模型与跟踪误差。本文方法设计了一款跟踪模型,并进行了仿真,验证了该模型的有效性,并进行了跟踪模型的跟踪模型。

在國際森林管理研究中,本文以混交林組合方式類型、概念與實施方法為主題,提出混交林政策落實執行機制與政策實施策略,將原本以國策制定及地區之林務政策立法與行政層面的政策調整,轉化、設計成可直接運用於一般化的管理與經營時空條件,明確規劃出政策執行時空和計畫在空間的維度分析,提出不同政策組合與實施在空間上的一條河、一片森林,以 10^2 年時間尺度與 0.05 km² 為管理單位,驗證了政策與實施在空間上

在基于视频事件的外部应用的研究中,针对视频区域内容提取尚不完善,并需综合考虑一下关于视频识别困难的问题,提出基于通用视频识别模型扩展模型,该方式通过引入嵌入模型提高模型,模型中并考虑视频内容内部影响,提升模型识别区域内容的特征一致性与公平合理,同时,设计一种考虑外部信息影响的模型,对模型中特定信息输出进行模型识别,并设计了一种通用模型识别模型,主要模型识别模型,在 DAVIS 数据集上,该方法在 FVD 指标上取得有最佳方法下降低 0.0,在模型鲁棒性和识别准确性方面取得更优结果。

综上所述,本文针对多模态数据语义任务中存在的建模难题,提出两两针对性网络结构图

计,揭示了模型的全点输入与可解性,为后续研究提供了方法指导与理论支撑。

关键词：人工智能平台；深度学习；数据挖掘；计算机视觉；决策支持系统

目录

摘要	III
Abstract	IV
目录	V
目目录	VI
表目录	VII
1 绪论	1
1.1 研究背景及意义	1
1.2 国内外研究现状	2
1.2.1 图像生成现状	2
1.2.2 文本生成和生成现状	3
1.2.3 图像生成和生成现状	4
1.2.4 视觉知识生成现状	5
1.3 研究目标与内容	6
1.4 论文组织结构	6
2 相关知识理论与数据工作	7
2.1 扩散模型	7
2.2 注意力机制	11
2.2.1 自注意力机制	11
2.2.2 非自注意力机制	12
2.3 交互与编码	21
2.4 文本编码器	22
2.5 视觉数据的清洗和标注	23
2.5.1 视觉数据清洗	23
2.5.2 视觉数据标注	24
2.6 本章小结	25
3 基于非注意力文本条件生成视觉生成	26
3.1 基于文本的图像生成迁移到视觉生成的问题分解	26

3.2 基于共享注意力的文本条件视频生成模型	29
3.2.1 预训练文生图模型介绍	29
3.2.2 基于预训练文生图模型的时间维度建模	31
3.2.3 文本-视频共享注意力运动控制网络	32
3.2.4 时序自适应调整网络	35
3.2.5 训练与推理过程	37
3.3 实验结果与分析	39
3.3.1 数据集介绍	39
3.3.2 评价指标	40
3.3.3 实验环境及参数设置	42
3.3.4 定量实验分析	44
3.3.5 定性实验分析	45
3.3.6 消融实验	49
3.4 本章小结	53
4 基于图文协同的多模态条件视频生成	55
4.1 基于图文协同的多模态条件视频生成问题分析	55
4.2 基于图文多模态协同视频生成模型	57
4.2.1 双流注意力协同网络	57
4.2.2 自适应动态权重融合	60
4.2.3 首帧引导自适应归一化网络	62
4.2.4 训练与推理过程	63
4.3 实验结果与分析	64
4.3.1 数据集介绍	64
4.3.2 评价指标	65
4.3.3 实验环境及参数设置	66
4.3.4 定量实验分析	67
4.3.5 定性实验分析	68
4.3.6 消融实验	70
4.4 本章小结	75

5 基于视频条件的解空间扩展内容生成	77
5.1 视频模态作为驱动条件的视频生成问题分析	77
5.2 基于通道融合的解空间扩展内容生成模型	79
5.2.1 基于通道融合的解空间扩展网络	79
5.2.2 条件感知动态选择网络	81
5.2.3 基于滑动窗口无监督训练的长视频生成	83
5.2.4 训练与推理过程	85
5.3 实验结果与分析	86
5.3.1 数据集介绍	86
5.3.2 评价指标	87
5.3.3 实验环境及参数设置	88
5.3.4 定量实验分析	89
5.3.5 定性实验分析	91
5.3.6 消融实验	93
5.4 本章小结	95
6 结论与展望	97
6.1 工作总结	97
6.2 展望	98
参考文献	101

6.2 展望

尽管基于扩散模型的多模态生成取得了显著进展，但该领域仍然面临着诸多挑战，同时也蕴含着巨大的发展潜力。

首先是生成内容的可控性和个性化。如何提高生成内容的可控性和个性化程度，以满足用户多样化的需求，是一个重要的研究方向。当前的多模态生成模型，在控制生成内容风格、布局、细节等方面仍然存在一定的局限性。未来研究可以探索更精细的控制机制，例如，引入更多种类的条件输入，如草图、姿势、语义分割图等、可学习的控制模块以及基于用户反馈的迭代生成等技术，提高生成内容的可控性和个性化程度。

然后是计算效率和模型轻量化。扩散模型生成过程的计算成本较高，限制了其在实际应用中的部署和推广。如何降低扩散模型的计算成本，使用轻量化模型，提高生成效率，也是一个重要的研究方向。未来研究可以探索更高效的采样算法、模型压缩和加速技术，如模型剪枝、量化、知识蒸馏以及基于硬件加速的优化方法，提高扩散模型的计算效率和模型轻量化程度。

98

最后是多模态生成模型的评估。多模态生成模型的评估仍然是一个具有挑战性的问题。现有的评估指标在衡量多模态生成内容的质量和语义一致性方面仍然存在一定的局限性，很多工作依然会使用人类偏好度评分以证明模型的有效性。未来研究可以探索更全面、更客观、更符合人类感知的多模态生成模型评估方法，例如，引入多模态一致性指标、人类偏好评估以及任务驱动的评估指标，更准确地评估多模态生成模型的性能。

未来，基于扩散模型的多模态生成技术将在更广泛的领域得到应用，例如，内容创作、虚拟现实、人机交互、教育娱乐等。随着研究的深入和技术的进步，多模态生成技术将为人类社会带来更多的创新和惊喜。