

詹俊

✉ jzhan24@m.fudan.edu.cn · ☎ +86 198 2167 0042 · 🏠 Homepage

🎓 教育背景

复旦大学, 上海	2024.09 – 2027.06 (预计)
博士研究生 人工智能, 导师: 邱锡鹏教授	
复旦大学, 上海	2022.09 – 2024.06
硕士 电子信息, 导师: 邱锡鹏教授	
华中科技大学, 武汉	2018.09 – 2022.06
学士 软件工程	

🔪 研究兴趣

多模态交互, 全模态基础模型

📄 论文

- **OmniVAE: An Audio-Video VAE with Cross-Modal Alignment for Joint Generation**
Jun Zhan*, Chen Yang*, Yitian Gong*, et al.
arXiv preprint, 项目负责人, 共同一作 [Paper] [Code]
- **AnyGPT: Unified Multimodal LLM with Discrete Sequence Modeling**
Jun Zhan*, Junqi Dai*, Jiasheng Ye*, Yunhua Zhou*, et al.
ACL 2024, 项目负责人, 共同一作 [Paper] [Code]
- **VStyle: A Benchmark for Voice Style Adaptation with Spoken Instructions**
Jun Zhan*, Mingyang Han*, Yuxuan Xie*, et al.
ICASSP 2026, 项目负责人, 共同一作 [Paper] [Code]
- **SpeechGPT: Empowering Large Language Models with Intrinsic Cross-Modal Conversational Abilities**
Dong Zhang, Shimin Li, Xin Zhang, **Jun Zhan**, Pengyu Wang, Yaqian Zhou, Xipeng Qiu
EMNLP 2023 Findings [Paper] [Code]

* 表示共同一作

🧩 研究/项目经历

OmniVAE	2025 年 12 月 – 2026 年 5 月
项目负责人	
• 提出联合训练的音视频 VAE, 通过单模态语义蒸馏和片段级音视频对比学习, 提升单模态与跨模态表征的可学习性 • 主导音视频训练数据管线建设, 定义音视频概念体系, 自动生成检索词并采集视频, 完成数据质量与语义一致性过滤 • 在基本保持音频与视频重建质量的同时, 提升下游文本到音视频生成质量与跨模态对齐水平	
AnyGPT	2023 年 9 月 – 2024 年 2 月
项目负责人	
• 基于离散表示构建端到端全模态模型, 统一建模图像、语音、音乐和文本的理解与生成 • 构建并发布首个大规模 Any-to-Any 多模态指令数据集, 实现任意模态输入输出的全模态对话 • 开源模型权重、代码、数据, GitHub Stars 800+, 引用数 300+, 模型和数据下载次数八万余次	

SpeechGPT 系列

2023 年 4 月 – 2025 年 1 月

组长

- 参与 SpeechGPT 系列模型研发：SpeechGPT首次实现大语言模型的内生语音对话，SpeechGPT-Gen实现音色可控的语音生成，SpeechGPT 2.0-preview实现端到端音频建模及多情感、多风格、多音色对话
- 负责 TTS 模型训练，支持笑声、重读、停顿、指令风格控制和多人对话
- 负责语音对话数据的构建
- 开源模型与代码累计获得 GitHub Stars 1.7k+，系列工作累计引用 800+

第二届粤港澳大湾区（黄埔）国际算法算例大赛（百万奖金赛题）2023 年 10 月 – 2024 年 4 月

出题负责人

- 设计“多模态大模型学科能力综合强化”赛题，聚焦模型的多模态学科知识理解与推理能力
- 基于 MOSS2 构建并发布 LLaVA-like 基线模型，为参赛队伍提供训练与评测基准
- 收集并处理近十年高考模拟题，构建覆盖多学科的多模态问答数据集与评测基准

MOSS

2023 年 1 月 – 2023 年 3 月

语音模块成员

- 参与研发国内首个类 ChatGPT 开源大语言模型 MOSS
- 负责语音合成模块，支持零样本音色克隆，扩展模型的语音生成能力

🧑 实习 & 工作

Meta AI (FAIR)

2026 年 6 月 – 至今

研究型实习生 法国巴黎

- 研发统一建模语音、动作与表情的端到端多模态交互模型，提升虚拟角色交互的自然度与跨模态一致性
- 研发支持身体、双手与面部表情流式编解码的多模态 VAE，并负责交互模型的训练优化

OPPO

2025 年 9 月 – 2025 年 11 月

校企合作实习生 中国北京

- 基于 OPPO 自有语音数据，负责预训练数据清洗与过滤、指令数据构建及训练框架构建
- 训练端到端语音对话模型，支持多风格、多情感和多音色语音生成

阿里巴巴未来生活实验室

2025 年 5 月 – 2025 年 8 月

研究型实习生 中国北京

- 训练端到端语音对话模型，提升模型对自然语言风格指令和隐式情感表达的响应能力
- 构建语音风格自适应评测基准 VStyle，覆盖声学属性、自然语言指令、角色扮演和隐式共情，并基于大型音频语言模型实现自动化评测