

陈旭天

chenxutian@buaa.edu.cn | github.com/Blossom0913 | blossom0913.github.io | 17369720372

教育经历

北京航空航天大学	2025/9 – 至今
人工智能硕士	北京, 中国
· 主要课程: 强化学习、深度学习、无人集群智能、算法设计与分析等	
暨南大学	2021/9 – 2025/7
人工智能学士	珠海, 中国
· 主要课程: 自然语言处理、机器学习、人工智能原理、敏捷软件开发等	
· 教学经历: 机器学习课程助教 (2024 年秋学期)	

实习经历

美团无人车业务部	2026/3 – 至今
实习生 - 决策规划算法	北京, 中国
· 路口数据闭环: 维护并修复人驾 Daily Pipeline, 系统定位上游数据缺失、路径不一致、Proto 迁移与 Dataset 加载等问题, 补齐历史断档数据; 设计覆盖驾驶模式、转向类型、RA 结果融合、过滤与分布统计的自动化收集 Workflow。	
· 大规模数据集构建: 构建 2025.06—2026.05 路口右转自驾数据集, 完成 SQL 场景提取、RA 挖掘、质量过滤与时间戳对齐; 将 278.9 万 Case / 13,501 小时 原始数据沉淀为 108.1 万 Case / 3,446 小时 训练可用数据, 并同步构建人驾数据集。	
· 模型训练与仿真评测: 跑通 Query Decoder 的“数据重刷—特征导出—Backbone/Head 训练—ONNX 导出—仿真评测”闭环, 完成 Shared Encoder 6.3 适配; 通过回归 Case 分析轨迹避让、选道与曲率等表现, 并优化 Mviz 路口模型可视化。	

科研经历

广东省智能科学与技术研究院	2025/2 – 2025/9
实习生	珠海, 中国
· 作为前三作者, 研究成果已申请 2 项软件著作权: 物流仓库多智能体机器人路径规划系统与基于积分任务管理与账单追踪的轻量化工具平台 。	
· 设计大鼠社会与攻击性行为分类的实验框架与基准数据集, 使用 DeepLabCut 标注身体关键点协同, 并以 LightGBM 对比 LSTM、CNN、GMM 等模型的性能; 项目代码: DLC_train ; 研究成果以共同第一作者身份发表于 Scientific Data 。	
· 使用 AutoDock Tools 处理受体蛋白靶点, 完成约 430 万个小分子的分子对接 (Docking) 实验, 为药物研发的大规模计算任务, 通过结合能和状态筛选药物分子; 设计并行框架并在 4 块 RTX 2080Ti 上部署, 源代码: Dock 。	

多智能体路径规划

研究助理	2024/3 – 2024/7
珠海, 中国	
· 复现 CL-CBS 多智能体路径规划算法, 完成环境搭建、代码调试、结果可视化与实验链路打通, 形成可复现基线。论文: CL-MAPF: 具有运动学与时空约束的类车机器人多智能体路径规划 。	
· 搭建批量实验与结果分析脚本, 实现实例自动求解、视频生成、超时跳过、CSV 汇总及性能曲线绘制, 用于 CL-CBS 与 Standard CBS 对比评估。	
· 基于 profiling 定位算法瓶颈, 分析 Spatiotemporal Hybrid-State A* 的关键优化机制, 包括三重启发、解析终端扩展、时空状态建模与代价惩罚设计。	

论文发表

1. Xutian Chen*, Guangyu Li*, Zihan Zhang*, Mingkun Xu, Zuoren Wang, Qianqian Shi. A Benchmark Dataset for Rat Social and Aggressive Behavior Classification . Scientific Data, 2026. *共同第一作者 (Co-first authors).	
--	--

竞赛经历

ASC2022	2021/11 – 2022/6
团队成员	珠海, 中国

- 项目目标：在有限算力（8× Tesla V100 16GB）和功耗约束下，完成 47 亿参数“源 1.0”大语言模型的预训练，并实现 55% 的训练加速（从 45h 降至 28h）。
- 内存优化：主导应用 **ZeRO-Offload** 与 **ZeRO Stage 1** 技术，成功将 75.2GB 的模型状态（参数、梯度、优化器状态）卸载至 CPU 内存，在仅 8×16GB 显存上跑通 47 亿参数模型，解决了 CUDA OOM 瓶颈。
- 训练加速：部署 **Megatron-DeepSpeed** 框架，设计并实施 4 路张量并行 + 2 路流水线并行的策略，相较单纯 8 路张量并行，将训练吞吐从 4.08 samples/s 提升至 4.66 samples/s。引入混合精度训练（**AMP**），在 NVIDIA Tensor Core 上利用 FP16 算术加速计算，并解决精度溢出问题。用 Intel MKL 数学库编译 PyTorch，并将其与针对 CPU 卸载优化的 DeepSpeed CPU Adam 优化器结合，显著加速了 CPU 端的优化器计算步骤。

获奖经历

- 2022 年世界大学生超算竞赛 国二等奖（rank 26）
- 2023 年广东省大学生数学竞赛 省二等奖
- 2024 年广东省大学生数学竞赛 省一等奖