

景昱

15835760778 | 350589792@qq.com | Agent / LLM 工程师 | 北京



天津大学人工智能与网络科学研究所硕士，曾任字节跳动豆包 AI 大模型工程师。聚焦 Agent Runtime、Memory、DeepResearch、工具调用、评测与后训练，具备从系统设计、工程落地到固定评测集验证的完整实践。

教育经历

天津大学(推免) | 人工智能与网络科学研究所 | 硕士 | 全日制 | 985 / 211 / 双一流 2022.09–2025.06
中北大学 | 软件学院 | 软件工程 | 本科 | 全日制 | GPA 4.52 / 5.00 | 综合排名 1 / 277 2018.09–2022.07

专业技能

- **Agent 与 Harness 工程**: Agent Loop、ReAct、Plan-and-Execute、Tool Use、Subagent / Multi-Agent; 自建状态机与编排; MCP 协议开发与集成。
- **Memory 与 Context Engineering**: 短期 / 长期记忆、Episodic / Semantic 多层记忆、滚动摘要、混合检索、KV / 前缀缓存、STABLE / DYNAMIC Prompt 拆分与主动压缩。
- **RAG 与评测**: GraphRAG、混合检索、RRF / Cross-Encoder Rerank、ChromaDB / Qdrant / Neo4j; Ragas、LLM-as-Judge、EM / F1 / 检索精确率及隔离测试集设计。
- **工程能力**: Go、Python、FastAPI 异步编程、Docker、Git; 熟悉 LangChain / LangGraph, 高强度使用 Claude / Cursor / Copilot, 实践 Skill-Driven Development。

工作经历

字节跳动 | 豆包 | AI 大模型工程师 | 北京 2025.07–2026.05

- **Agent Runtime**: 参与豆包新一代 Agent Runtime 核心链路建设，覆盖请求接入、上下文组装、Agent 调度、工具调用与 SSE 流式输出; 统一 ToolCall / ToolResult 协议并接入搜索、文生图、POI、路线规划和 RAG 等工具，工具调用成功率稳定在 **96%+**，失败兜底成功率达到 **97%+**。
- **稳定性与性能**: 建设 Memory Injection、SSE Stream Event 与 error_stage 分层指标，将线上问题定位时间从 **30 分钟**以上缩短至 **5 分钟**以内; 将首 token P95 从 **2.8s** 降至 **1.6s**，finish 到达率提升至 **99.7%**。

项目一 | 豆包 Agent Runtime 与自进化评测平台

- **系统设计**: 从零独立实现 Runtime、Memory、Skills 自进化、Eval 与安全五大子系统，打通「评测驱动 → 失败归因 → 技能进化 → 回归验证」的 Agent 自改进闭环。
- **Runtime 与工具调度**: 针对传统 Agent Loop 扩展性差、工具调度耦合的问题，设计 **Dispatch Table** 动态注册 + **Batch Tool Calling** + **Interrupt** 中断恢复机制，新增工具平均约 **30 行**即可接入核心循环; 在 **120 任务 × 3 轮**评测中，多工具任务 LLM 交互轮次 **↓45%**，端到端耗时 **↓38%**。
- **Memory 与 Context Engineering**: 构建 **Episodic JSONL + Semantic Vector** 双层记忆，实现 Context 四操作 (Write / Select / Compress / Isolate) 与工具结果外部化; 在 **80 组**长程任务 (平均 **40+** 轮) 评测中，单任务 Token 从 **12.3k** → **5.6k** (**↓54%**)，任务成功率从 **71%** → **76%**。
- **Eval 驱动的技能自进化**: 实现「轨迹分析 → 失败归因 → 新技能提案 → 沙箱回归 → 确认集成」闭环，以 **60 Case** 隔离测试集作为成功率门控; 经 **6 轮**迭代，任务成功率从 **52%** → **80%** (**+28pp**)，失败自动归因覆盖率 **85%**，技能提案回归通过率 **60%**。
- **安全与权限**: 实现 **三级权限模型**、**Dry-run** 副作用预览与成本守卫; 预设高危操作 **100%** 进入审批流程，单任务成本上限 **\$0.5**，实际运行中未出现越权调用。

项目二 | 豆包 DeepResearch 多智能体检索系统

- **检索与知识资产**: 构建 **NaiveRAG**、**GraphRAG** 与 **DeepResearch** 分层链路，基于 Document / Chunk 向量资产、Neo4j 实体关系图、实体索引与 Leiden 多层社区摘要，支持局部证据链、全局 Map-Reduce 与迭代深研。

- 研究 Agent：实现 问题分解、多假设分支、社区感知的 Chain of Exploration 与迭代检索；通过 动态知识图、EvidenceChainTracker、信息缺口检查与最大迭代预算控制停止条件，并在合并阶段完成矛盾消解与引用生成。
- 质量与可观测：输出 reasoning chain、证据统计、动态知识图、中心实体、矛盾列表与执行日志，使研究过程可回放、可诊断；在固定评测快照中取得 EM 0.7333、F1 0.7500、事实一致性 0.9333、检索精确率 0.7000，平均延迟为 87.88s。

项目三 | 豆包自主 Search Agent 后训练 (SFT + AgenticRL)

- 训练目标：针对多 Agent Pipeline 单条查询需多轮 Planner / Verifier / Synthesizer 调用、推理成本高的问题，以 Qwen3-4B 为底座训练 端到端 Search Agent，使模型自主完成「推理 → 工具调用 → 环境反馈 → 最终回答」。
- SFT 与 GRPO：设计 工具调用轨迹格式、HTTP Retrieval Server 与 Search-R1 风格多轮 GRPO 框架；围绕 Correctness、Faithfulness、Context Precision、hop recall 与格式约束完成 v4e-v15e 四阶段 Reward 迭代，并保证 SFT / GRPO 推理模板一致。
- 两阶段训练：在 185 条中文金融测试集上，先强化 grounding，再使用 correctness-heavy reward；v14e step30 将 Judge_C 从 SFT base 的 0.279 提升至 0.334，Faith 从 0.169 提升至 0.199，首次在两个维度同时超过 SFT 基线。
- 失败分析：验证 硬门控会导致全面退化，且 Stage 3 在 hop_pr 为 0 占 66%、faith 为 0 占 100% 时出现 Reward 信号稀疏；据此确认 检索精度需在第一阶段训练，而不能靠后置调权补救。

开源项目

AgentGuide | Agent 学习知识库与面经知识图谱（核心作者）

- 项目影响力：GitHub 5k stars，multiagent Topics Top 20，LLM Topic 全站前 0.5%，已成为国内 Agent 方向最具影响力的开源学习项目之一。
- 项目地址：github.com/adongwanai/AgentGuide
- 知识图谱构建与自动化 Pipeline：基于 OpenClaw 设计全网面经自动爬取、清洗、归并 Pipeline，累计处理 3 万 + 篇面经；通过 LLM 自动抽取高频知识点，按「原理层 / 工程层 / 面试层」三级分类构建 Agent 领域知识图谱，集成至 InterviewGuide 子站并持续更新，覆盖 RAG / Agent 训练 / 数据合成全技术栈。

荣誉奖项

- ACM-ICPC 国际大学生程序设计竞赛全国邀请赛（西安）铜牌；PAT 甲级满分；CCF-CSP 认证 300 分（全国前 3.53%）。
- 中国软件杯全国总决赛三等奖；蓝桥杯国赛三等奖；全国 IT 技能大赛国赛一等奖；团体程序设计天梯赛全国银奖。
- 校长奖章、国家奖学金、国家励志奖学金、综测特等奖等国家级、省部级奖项二十余项。