



智能体在人才研究领域的实践探索

北京人才发展战略研究院
综合服务部、人才开发部
2026 年 7 月

说 明

本课题为我院自研课题，由综合服务部、人才开发部共同承担。针对通用大模型在人才研究领域泛而不专、学术规范性差等局限，课题构建了专属智能体 Talegent，旨在打造可信可用的“数字研究员”，全面推行“人+智能体”工作模式。

研究团队系统整合政策文件、学术成果、人才榜单等多源数据，运用检索增强生成与知识图谱技术，搭建人才领域专属知识库。在此基础上，实现智能问答、科研助理、辅助决策三大功能。智能问答实现人才、科技、产业相关信息的实时检索与应答；科研助理自主完成文献综述、问卷分析、图表生成等任务，减少重复性工作人力投入；辅助决策融合多源数据，支持研究报告和参阅信息的起草撰写，形成“数据驱动+专家决策”的人机协同模式。同时，课题从交互感知、文本处理、逻辑推理三大维度，与 GPT-4、文心一言等通用大模型进行系统性对比，验证了 Talegent 在多模态支持、长周期记忆等方面的显著优势。

课题研究成果基本建成人才研究智能体，完成 PC 端和移动端智能助理系统开发，系统事实准确率超 98%，初步实现从通用工具到专属科研伙伴的跨越。随着数据成果的持续积累和模型能力的训练迭代，该智能体将成为研究院未来的核心竞争力，为人才工作数字化转型提供可复制、可推广的实践范式。

目录

1 引言	1
1.1 研究背景	1
1.2 研究目的及意义	1
1.3 研究内容	2
2 文献综述	4
2.1 智能体的概念、原理及特点	4
2.2 智能体的发展	5
2.3 智能体的应用	7
2.3.1 智能体在教育领域的应用	7
2.3.2 智能体在商业领域的应用	8
2.3.3 智能体在医疗领域的应用	10
2.3.4 智能体在社科领域的应用	11
3 人才领域知识库搭建	13
3.1 知识库的作用	13
3.2 知识库的搭建方法	13
3.2.1 知识获取与预处理技术	13
3.2.2 向量化与存储优化技术	15
3.3 知识检索与更新机制	16
3.4 知识库的管理	19
3.4.1 知识管理	19
3.4.2 用户管理	20
3.4.3 质量管理	21
3.4.4 系统管理	22
3.4.5 总体评价与用户反馈	23
4 智能体问答模块搭建	24
4.1 智能问答模块的作用	24
4.2 智能问答模块的搭建方法	24
4.2.1 深度语义理解与多轮对话技术	24
4.2.2 问答生成与验证技术	25
4.2.3 用户意图与个性化响应技术	25
4.3 问答模块的效果检验	27

4.3.1 响应速度	27
4.3.2 语义理解	28
4.3.3 多轮对话能力	30
4.3.4 结果准确性	31
4.3.5 总体评价与用户反馈	32
5 智能体科研助理模块搭建	34
5.1 科研助理模块的作用	34
5.2 科研助理模块的搭建方法	34
5.2.1 自主科研任务分解与规划	34
5.2.2 科研工具调用与集成	35
5.2.3 多智能体协作与知识共享技术	35
5.3 科研助理模块的效果检验	36
5.3.1 功能完整性检验	36
5.3.2 任务执行效率评估	38
5.3.3 结果准确性验证	39
5.3.4 用户体验评估	40
5.3.5 实际应用价值验证：典型案例分析	41
5.3.6 结论与优化方向	43
6 智能体辅助决策模块搭建	45
6.1 决策辅助模块的作用	45
6.2 辅助决策模块的搭建方法	45
6.2.1 多源数据融合与特征工程技术	45
6.2.2 动态优化与人机协同决策机制	46
6.3 辅助决策模块的效果检验	46
6.3.1 检验设计与实施流程	47
6.3.2 历史场景还原测试	47
6.3.3 综合评估	48
7 人才研究智能体与通用智能体对比分析	50
7.1 交互感知能力对比	50
7.1.1 核心对比维度与指标详解	50
7.1.2 Talegent 的场景化交互能力与实证表现	53
7.2 文本处理能力对比	55
7.2.1 核心对比维度与指标详解	55
7.2.2 Talegent 的文字处理能力与实证验证	59

7.3 逻辑推理能力对比	61
7.3.1 核心对比维度与指标详解	61
7.3.2 关键差异的场景化解析	65
7.4 智能体优势总结	67
8 结论与展望	69
8.1 主要结论	69
8.2 存在问题	70
8.3 研究展望	71
参考文献	73

1 引言

1.1 研究背景

自 2022 年 11 月 ChatGPT 引发“生成式 AI 革命”以来，全球人工智能（Artificial Intelligence, AI）的发展明显提速^[1]。在短短几年时间内，AI 模型的性能实现了显著提升，其功能已从最初仅能输出文本扩展为支持多模态交互，在语言理解、数学推理和编程能力等多个方面均取得了长足进步。与此同时，AI 的使用成本大幅下降。特别是在 2024 年国产大语言模型 DeepSeek 上线后，AI 的训练和使用成本迅速降至原有水平的几分之一^[2]。受上述两方面因素的共同推动，AI 行业规模迅速扩张。据估算，2025 年全球 AI 市场规模已接近 1 万亿美元，预计到 2030 年，AI 市场规模将进一步增长至 2 万亿美元^[3]。

尽管当前 AI 技术已取得显著进展，但其能力仍存在一定局限性。现有大语言模型虽然能够在人类的指令下完成多项指定任务，但其本身并不具备自主决策能力，因而应对复杂任务能力仍存在一定局限性。在此背景下，智能体，特别是通用型智能体正逐渐成为各领域竞相布局的战略重点。据统计，2023 年中国智能体市场规模约 59.81 亿元。目前，智能体已广泛应用于金融、教育、医疗、制造、交通等多个行业，展现出强大的渗透力和应用潜力。同时，由于智能体对于复杂任务的数据收集、逻辑分析、思考速度的优势，已经成为不可或缺的工具。

1.2 研究目的及意义

在人才研究领域，智能体的应用仍处于起步阶段，相关问题缺乏深入的研究。尽管智能体在人力资源管理中的某些场景（如招聘、员工培训等）已有初步应用，但针对人才研究的深层次需求，如人才画像构建、人才流动趋势研判、人才政策分析等方面，智能体的应用较为有限。这种空白不仅限制了人才研究的效率和精准度，也阻碍了智能体技术在更广泛领域中的潜力释放。

为了填补这一空白，本课题旨在探索智能体在人才研究领域的应用实践。通过结合智能体的自主学习、适应性和交互能力，构建一套能够高效支持人才研究的智能体系统，推动人才研究的智能化。这一研究不仅能够为人才研究提供全新

的技术手段，也将为智能体技术在更广泛领域的应用提供理论和实践参考。

本课题具有重要的理论价值和显著的实践意义。在理论层面，研究成果有望填补智能体在人才研究领域应用的空白，深化人工智能与人才学的交叉融合，为相关研究提供新的理论支撑；同时，通过持续挖掘智能体在人才识别、评估与发展中的新能力，推动人才理论体系从“静态描述”向“动态预测”演进，加速人才研究的数字化转型。在实践层面，成果可直接赋能人才研究机构，提升研究人员产出高质量成果的效率；为各级人才工作部门提供高效、稳健的智能辅助工具，助力其在形势分析、问题诊断和政策制定中更加精准高效；更能服务于中央及市委重大决策需求，为主管领导随时调阅关键信息、快速响应复杂局势提供有力支撑，切实提升决策的时效性、科学性与精准度。

1.3 研究内容

本课题以构建并优化适用于人才研究领域的专属智能体为目标，搭建人才领域专属知识库，实现 AI 问答、AI 助研、AI 辅助决策三大能力，探究智能体各模块的搭建方法，对智能体在人才研究领域的实践效果进行分析，并与通用大模型进行对比研究。主要研究内容：

（1）探索构建人才领域专属知识库

以人才研究院政策库、成果库、榜单库、期刊库等为信息源，依托向量数据库技术，实现智能体对人才研究领域知识的快速处理、储存、检索。借助知识图谱技术，构建结构化的人才语义知识库，迅速理清人才研究中的概念及相互关系，并查缺补漏，联网自更新知识库。

（2）探索搭建更及时、更准确的智能问答模块

依托人才领域专属知识库和链式调用能力，智能体可初步实现问答功能。优先基于专属知识库进行应答，若无相关知识，则自动联网进行检索应答，并标注信息来源。以此，提升问题响应及时性与精准度，实现人才领域知识全知道。

（3）探索搭建更高效、更全面的科研助理模块

依托跨文档检索、摘要、综述能力，解决研究工作信息搜集慢、碎片化等难题。依托智能体对文字、图像、表格等多模态数据的处理能力，完成大数据分析、检索、解读、处理以及图表绘制、文本生成等，大幅提升研究效率。

（4）探索搭建更科学、更专业的辅助决策模块

通过智能体分析海量数据并根据需求调用不同大模型，对复杂问题进行系统性推理与判断，并呈现整个研判过程。同时，整合政策文本、人才数据、时事新闻等多模态信息，生成研究报告、参阅信息等供科学决策。

（5）开展人才研究智能体与通用大模型的系统性对比研究

从交互感知、文本处理、逻辑推理三大核心维度，选取多模态支持、上下文记忆、学术格式规范、检索精准度、多步推理能力、量化计算精度等关键指标，将智能体与 GPT-4、Claude 3、文心一言等通用大模型进行全面对比，明确智能体的核心竞争力，为智能体在人才研究领域的优化迭代与推广应用提供实证依据。

2 文献综述

2.1 智能体的概念、原理及特点

智能体（Agent）是指能够自主感知环境、作出决策并执行行动的智能单元，其核心架构包括大模型、工具模块和记忆模块。大模型作为“大脑”负责推理规划，工具模块（如搜索引擎、API）实现环境交互功能，记忆模块则存储历史信息以支持个性化服务。其基本原理是基于人工智能技术构建的自主化系统，具备环境感知、动态决策、任务执行与持续学习的核心能力，通过传感器或数据接口获取信息（如文本、图像、结构化数据），利用机器学习、自然语言处理、知识图谱等技术进行推理与规划，最终通过执行器（如自动化程序、机器人或交互界面等）完成特定目标。和传统 AI 相比，智能体主要具有如下特点：

（1）自主性

自主性是智能体最基本的特性之一，指的是智能体能够独立地感知环境、做出决策并执行行动，而无需持续的人类干预或指导。自主性使得智能体能够在动态且不可预测的环境中独立工作，适应变化并调整其行为。例如，自动驾驶汽车就是一个具有高度自主性的智能体，它能够在复杂的交通环境中感知周围车辆和行人，自主规划路径、控制速度和做出避障决策。自主性不仅减少了对人类监督的依赖，也使得智能体能够在需要实时反应的任务中保持高效性和可靠性。

（2）反应性

反应性是指智能体能够迅速感知环境变化并及时做出响应的能力。这种特性使得智能体能够在面对突发事件或紧急情况时做出快速而有效的反应。反应性对于实时系统和动态环境中的智能体至关重要，例如在机器人控制中，智能体需要即时感知障碍物的出现，并立即调整其路径以避免碰撞。虽然反应性通常意味着对当前状态的即时响应，但高级智能体还可以结合历史数据和预测信息，使得反应更加智能和灵活。

（3）主动性

主动性是智能体能够主动设定目标、规划行动并采取措施实现这些目标的能力，而不仅仅是对环境的变化做出反应。主动性使得智能体不仅仅局限于被动应

对外界刺激，而是能够根据其内在目标和动机采取积极行动。例如，一个智能家居系统可以主动学习用户的日常习惯，提前调节室内温度或照明，以提高用户的舒适度。具有主动性的智能体能够在环境中自主探索、发现问题并提出解决方案，从而在实现长期目标的过程中展现出更大的灵活性和创造力。

（4）社会性

社会性指的是智能体与其他智能体或人类之间进行互动、协作和交流的能力。具有社会性的智能体能够理解和遵循社会规范，与其他个体协调行动，以共同完成复杂任务。例如，在多智能体系统中，各个智能体需要通过通信协议分享信息、分配任务，并通过协作实现团队目标。社会性还体现在人机交互中，如智能语音助手能够理解用户的指令，并通过对话形式提供反馈和建议。通过增强社会性，智能体能够在团队工作、群体决策和协作环境中表现出更高的效率和有效性。

（5）进化性

进化性是指智能体通过学习和适应，在长期运行中不断提高自身能力的特性。具有进化性的智能体能够在面对新的环境或任务时，通过自我调整和优化，逐步提升其性能。这种特性通常与机器学习、进化算法或强化学习相结合，使得智能体能够在不断变化的环境中保持竞争力。例如，强化学习智能体通过与环境的持续交互，不断调整其策略以最大化长期收益。进化性使得智能体具备应对不确定性和复杂性的能力，使其在长期任务或未知环境中表现出色，并能够随着时间推移变得更加智能和高效。

2.2 智能体的发展

智能体的概念最早由马文·明斯基（Marvin Minsky）在他 1986 年出版的《思维的社会》一书中提出，其将思维描述为由大量相互作用的智能体构成的复杂系统，每个智能体都执行特定的任务，并通过协作完成复杂的认知活动。这一思想为智能体的研究奠定了理论基础，推动了人工智能领域对自主决策系统的进一步探索。

2007 年英伟达推出并行计算平台 CUDA，它允许开发者使用英伟达的 GPU 进行通用计算。CUDA 极大地提升了人工智能模型的训练速度，尤其是在处理大规模数据和复杂模型时表现突出。通过 CUDA，研究人员能够更高效地训练深度

神经网络，加速了包括智能体在内的各种 AI 技术的发展。时至今日，CUDA 以其背后强大的硬件支持、完善的生态环境、丰富的社区资源，被广泛应用于计算机视觉、自然语言处理、机器人等诸多领域，成为推动 AI 进步的核心技术之一。

2017 年谷歌提出了 Transformer 模型，这一模型通过自注意力机制显著提升了自然语言处理的效率和效果。Transformer 模型为后续的大语言模型奠定了基础，极大地改变了智能体处理语言任务的方式。Transformer 的提出不仅提升了模型的计算效率，还使得智能体能够更好地理解和生成自然语言，这为智能体在语音助手、翻译、文本生成等领域的应用打开了新的大门。

2018 年 BERT 模型的发布标志着大语言模型时代的开始。BERT 通过双向编码器实现了更深层次的语言理解，推动了自然语言处理技术的革命性进步。随后，GPT-2、GPT-3 等模型相继发布，进一步推动了智能体的发展，使其具备了更强的语言生成和理解能力。这些模型的成功使得智能体在对话系统、内容创作、信息检索等方面的应用达到了新的高度。

2021 年 OpenAI 发布了世界上首个多模态人工智能模型 DALL·E，它可以通过文本描述生成对应的图像。这一技术突破展示了智能体跨越不同模态（如语言和视觉）进行协作的能力，为智能体的应用领域开辟了新的可能性。DALL·E 的出现标志着智能体在创意生成、艺术设计、视觉推理等领域的潜力得到了极大释放，推动了 AI 在多模态任务中的进一步研究和应用。

2023 年 AutoGPT 的出现标志着智能体进入了一个新的发展阶段。AutoGPT 结合了 GPT-4 和 GPT-3.5 技术，能够自主完成复杂项目任务，体现了高度自主性和智能化水平。它不仅展示了大语言模型在复杂任务管理中的潜力，还推动了智能体技术向更广泛、更复杂的应用场景扩展，如自动化办公、项目管理和智能决策支持。AutoGPT 的成功预示着未来智能体在自主性和任务执行能力方面将取得更大的突破。

在智能体成为 AI 应用主旋律的当下，全球顶尖技术厂商争相涌入，开发应用智能体。2025 年，谷歌基于 Gemini 2.0 架构推出了三个最新智能体原型：通用大模型助手 Project Astra、浏览器助手 Project Mariner 和编程助手 Jules，其中编程助手 Jules 能够直接集成到 GitHub 的工作流程系统中，分析复杂代码库并实施修复。2024 年，微软宣布在 Dynamics 365 中集成 10 个自主 AI Agent，这些智

智能体能够自动执行客服、销售、财务、仓储等业务流程，具备自主学习能力，可以自动执行跨平台的超复杂业务。美国著名电信公司 Lumen 通过 AI Agent 每年能节省 5000 万美元成本，相当于增加了 187 名全职劳动力。在国内，阿里云通义千问开源了全新的视觉模型 Qwen2.5-VL，Qwen2.5-VL 和 2.5MAX，不仅在性能上取得了显著提升，而且在电脑端、手机端等展现出了强大的应用潜力。Qwen2.5-VL 能够直接作为视觉 Agent 进行操作，推理并动态使用工具，支持在计算机和手机上完成多步骤的复杂任务，例如自动查询天气、预订机票、发送消息等。随着技术的不断发展，智能体的应用范围将进一步扩大，为各领域带来更多的创新和价值。

2.3 智能体的应用

2.3.1 智能体在教育领域的应用

人工智能赋能现代教育的大背景下，生成式人工智能技术在教育领域的应用和研究成为新的研究趋向^[4]。党的二十届三中全会召开后，教育部党组书记、部长怀进鹏提出，将大力推进智慧校园建设，打造中国版人工智能教育大模型，探索大规模因材施教、创新性与个性化教学，促进人工智能推动教与学融合应用，开发教育专用人工智能大模型。智能体技术在教育领域的落地实施需要系统化、科学化的顶层设计，结合教育教学改革，重点在课程设置、师资培训和实践平台建设。

首先，应构建分层递进的人工智能课程体系，覆盖基础知识、应用技能和思维培养。例如，中山大学于 2025 年共启动 111 门 AI 课程的立项建设，课程分类包括人工智能通识课程、人工智能专业核心课程、人工智能技术赋能课程以及人工智能学科交叉课程四类，打造了一批示范性课程，推动了教学模式的创新与变革。这些课程深度融合人工智能技术，为学生提供了前沿、多元的学习体验。

其次，在生成式 AI 与智能体深度嵌入教育生态的背景下，教师素养正面临结构性重构的时代挑战^[5]。在人工智能技术高速发展的背景下，教师需要具备跨学科整合能力、数字化教学工具的应用能力以及培养学生批判性思维和创新思维的能力，利用智能体技术，可以通过专项培训提升教师的素养和教学能力。在模拟人类认知过程与处理大规模数据方面，生成式人工智能具有独特优势，为教师

决策自动化评价带来新的可能，但其在具体应用中亦面临情境适配不足、专业引导有限以及成本与质量难以兼顾等挑战。基于提示词工程构建的教师决策自动化评价智能体，包括输入器、生成标准评价器、决策等级评价器与决策评价反馈器四个组成部分，可实现教师情境化决策数据的接入、融合人机智慧的决策评价标准构建、教师决策智能化分级评价以及形成发展导向的个性化反馈。以幼儿园保教活动为案例情境，通过将智能体的评价结果与专家评估结果进行对比发现，该智能体在提升评价一致性方面优于传统学科专家，且提示词工程的应用显著增强了生成式人工智能与人工评价之间的一致性。该智能体不仅实现了对教师决策能力的自动评价与精准定位，还能够提供蕴含教育智慧、具有专业深度的反馈，为教师专业发展提供了有力支持^[6]。

最后，要推动优质教育资源数字化，建设开放共享的人工智能教学平台和实践基地，为学生提供真实的人工智能应用场景。例如，北京市海淀区中关村第一小学教育集团构建了集智慧管理、智慧教学和智慧教研于一体的综合性“智慧教育云中心”。该中心能够打通各类应用平台，全面、立体地收集教育教学数据。在具体学科应用方面，数学课通过作业机生成多维学情报告，精准定位学生的学科薄弱环节；英语课则采用 AI 听说技术，通过答题器实时收集学生语音，提供即时评价反馈。江苏省外国语学校通过对“校园志愿服务平台”项目式学习的研究发现，教育智能体通过多模态数据画像、动态资源匹配和智能反馈机制，可实现项目目标个性化定制、跨学科场景动态生成、学习活动精准适配及多维度能力评价，有效提升学生问题解决能力与创新思维，为智能体支持的项目式学习提供实践参考^[7]。

可以预见，随着生成式人工智能的不断发展，在不久的将来，教育智能体将成为项目式学习不可或缺的“智慧伙伴”。通过合理的构建和应用，其不仅能助力教师实现规模化教育与个性化培养的有机统一，更能让学生在技术赋能中成长为兼具数字素养与人文关怀的未来创新者。

2.3.2 智能体在商业领域的应用

随着生成式人工智能技术的快速演进，智能体已具备自主规划、执行和学习的能力，正逐步成为商业世界的关键参与者，一种全新的经济形态——智能体经济正在逐步显现。当前全球 AI 市场规模迅速扩张，智能体的应用正在重塑商业

生态系统、企业组织形态和商业模式^[8]。

当前市场上已涌现出大量智能体产品。例如，OpenAI 公司推出的 AutoGPT、Operator 和 Deep Research，Anthropic 公司推出的 Computer Use，以及中国初创公司 Monica 推出的 Manus 等智能体，均能在一定程度上自主协助人类完成复杂任务。同时，众多企业已经开始采用“数字员工”和“AI 助手”替代人类员工执行工作任务。智能体的迅速崛起已经引起了行业的广泛关注。例如，Forrester 公司已将智能体列为 2025 年关键新兴技术之一，而全球领先的 IT 研究与咨询公司 Gartner 更是将其评为 2025 年十大技术趋势之首^[9]。毫无疑问，智能体的广泛应用将显著提升生产的自动化水平，从而大幅度提升生产效率。根据预测，到 2028 年，有 33% 的企业级软件应用将集成智能体功能，15% 的日常业务决策可由智能体自动完成，智能体的应用有望使组织运营成本降低 40%，同时提升 20%~30% 的营业收入^[10]。

然而，相较于单纯的生产效率提升，智能体对商业生态系统、企业组织形态和商业模式所带来的影响可能更为深远和关键。随着其自主决策能力的不断增强，智能体将逐渐成为市场中除人类之外的另一重要决策主体，进而对整个商业生态系统产生深刻影响。在宏观层面，智能体的发展将重塑商业生态系统中各参与者的角色及其相互关系；在微观层面，智能体则将引发企业组织形态的重大变革。此外，企业的商业模式也将经历深刻的调整：一些传统、低效的商业模式将因智能体的兴起而逐步被淘汰，与此同时，一系列依托智能体的新商业模式也将不断涌现^[11]。

在国际贸易方面，当前国际贸易规则体系复杂度呈指数级增长。以关税与原产地规则为例，我国已签署的 23 个自贸协定中，仅 RCEP 原产地规则就包含 44 项条款及超千条合规路径，叠加临时性单边贸易措施，规则维护难度空前。风险研判数据源已从通关申报扩展至开源情报、区块链存证、物联网等多源渠道，亟须多模态信息智能处理能力。智能体技术通过整合思维链推理、检索增强生成、工具调用等机制，有效解决传统大语言模型的“幻觉”问题。智能体构建的规则解析系统，已在合规领域实现超 90% 的准确率，可模仿关员完成信息识别、规则匹配、风险推理等全流程操作，满足海关高可靠性业务需求。专业智能体集群具备协同进化能力，既可通过单智能体学习关员操作积累经验，又能通过多智能体

交互优化业务流程，形成持续进化的智能生态。在“单一窗口”场景中，智能体可与物联网设备实时交互，实现货物在途监控与合规检查；与区块链技术集成，构建安全透明的去中心化交易环境。如微软发布的商业智能体可以简化工作流程并提高销售、客户服务、财务和供应链管理的效率，有望推动贸易运营范式突破性创新^[12]。

2.3.3 智能体在医疗领域的应用

2025 年初，世界经济论坛发布报告《人工智能驱动健康的未来：引领潮流》，认为全球医疗体系正站在重大转折点，AI 的广泛应用或将重塑医疗生态。在医疗健康这一关乎人类福祉的重要领域，智能体正以其独特的优势，逐步成为推动医疗创新的重要力量。

借助深度学习与大数据分析技术，智能体能够协助医生进行疾病诊断、治疗方案制定及患者管理等多个关键环节的工作。2025 年，联影创新大会发布了“元智”医疗大模型，并同步推出覆盖影像诊断、临床治疗、医学科教、医院管理、患者服务等多个场景的十余款医疗智能体。“元智”医疗文本大模型接入 DeepSeek 后，它的复杂推理和文本处理能力大幅提升，能高效响应医疗场景需求，而且参数量较小，适合医院进行本地化部署。语音大模型具有医疗术语精准识别与智能声纹分析能力，经过丰富的医疗场景语料训练，即便在嘈杂环境和复杂对话中，仍能精准识别科室、检查、药品和疾病症状等专业术语，支持多人同时对话场景下的身份识别，与医生高效协作完成任务。视觉大模型在 4D 建模等医疗场景中，可通过高速电影级渲染技术，精准描绘人体器官和组织分割面，展现出临床与科研价值。医疗影像大模型由数千万医疗影像数据、数十万医疗级精细标注数据训练而成，支持 10 个以上影像模态、约 300 种影像处理任务。在处理复杂病灶诊断、器官分割等关键任务上，它的精准度测评超过 95%，性能超越此前的行业最优模型^[13]。

医疗智能体能够将复杂的医学知识以通俗易懂的形式呈现给患者，在降低相关人力和经济成本的同时可以提高健康教育的人群覆盖面与质量。例如，由中国电信上海公司、上海库帕思科技有限公司、徐汇区卫生健康委员会、中电信人工智能科技公司共同打造的“医疗行业智能体”，通过融合大模型算法与权威医学知识，为家庭医生打造全场景临床智慧决策支持系统。通过该平台，家庭医生可

获得全方位临床决策辅助能力，面向患者提供专业健康咨询、康复护理及个性化运动营养等知识，居民可以享受到“家门口享优质医疗服务”^[14]。2025年6月，蚂蚁集团正式发布全新医疗健康AI应用——AQ。对患者而言，AQ等新一代医疗AI产品能提供覆盖90%以上常见健康问题的初步判断和就医指引；对医生而言，它能在病历采集、检查解读、用药提醒等环节协助处理大量事务，释放医生的时间与精力，去解决真正复杂的疑难重症，实现资源配置逻辑的重构^[15]。2025年12月15日AQ更名蚂蚁阿福APP新版上线，2026年1月8日的最新月活已达3000万，每天在阿福上咨询健康问题已超1000万。

医疗智能体在提升医疗服务可及性、优化临床流程和推动精准医疗等领域展现了广阔前景；然而，目前其价值评估仍处于理论阶段，临床效能亟须通过更多前瞻性研究和真实世界应用验证^[16]。不少业内专家指出，目前医疗智能体仍处于“工具赋能”阶段，要走向深水区，还需在三个方向突围：数据获取、临床有效性、商业化路径。推动医疗智能体价值落地需要多学科协作，包括技术开发、医学研究和伦理学专家的共同参与。政策层面应完善医疗智能体的监管体系与技术评估标准，建立创新与安全平衡机制，加速医疗智能体在临床中的应用，实现技术对人类健康福祉的改善。

2.3.4 智能体在社科领域的应用

智能体在社会科学领域的应用正从数据分析工具向社会治理引擎跃迁，通过多模态感知、因果推理与政策仿真等能力，重构社会问题研究范式。例如，用于分析文字资料的传统手段，比如自然语言处理软件和早期的机器学习技术，往往需要相对精深的计算知识和大量手动编码训练的数据。即便如此，对这些技术的运用通常也只能达到有限的满意度，在处理大批量文本的时候更是捉襟见肘。大语言模型正在改变这一局面，譬如清华大学社会科学学院医学社会学研究中心在过去几年与国内十多所大学的老师和100多名学生通过深入访谈收集了中国人有关临终关怀的叙事，从中调取120份针对逝者家属的深入访谈记录，一次性输入国产大语言模型豆包，豆包在几分钟内就对这批逾74万字的资料进行整理并做出了概述^[17]。

对大语言模型的训练不仅基于自然语言，而且得益于图像识别。四名社会学家用谷歌街景项目的一万多张照片，分析了美国三大城市的街头垃圾分布程度。

尽管这四名学者关心的主要问题是街头垃圾图像的编码技术与垃圾分布程度的纠错方法，但该研究仍有希望使街头垃圾成为分析区域性社会分层的参考^[18]。卷积神经网络能够识别复杂的图像，比如纹理和色调。在大语言模型的图像库进行预训练，然后针对特定任务进行微调，计算机视觉模型可以识别社交媒体发布的大量图像。一名社会学家在利用 ChatGPT 分析“黑命贵”（Black Lives Matter）游行照片后指出，这个例子突出了多模态生成式人工智能如何在无须特定任务训练的情况下对图像进行编码和描述。这种方法可以轻松扩展，用于分析大量黑人抗议活动的图像，拓展我们使用视觉数据研究集体行动的能力^[19]。尽管大语言模型具有在社会学领域广泛运用的途径和潜力，但谨慎对之仍十分重要。用于训练大语言模型的数据其来源往往并不透明，模型输出的结果常常并不可靠，有时甚至输出虚假的或具有误导性的信息。使用这些模型时还应注意隐私方面的问题，输入没有匿名的研究对象信息更是一种风险。

旨在普惠的人工智能前景，在国际学界已经引起了一定程度的社会科学研究，集中于可持续发展议题^[20]，但在中国社会科学界仍屈指可数。但相信经过努力，以数据分析为代表的 AI 社会学实证研究和以实地参与观察研究为代表的 AI 海外人类学，可能在中国成为一个意义重大的学术探索和实际应用领域。

3 人才领域知识库搭建

3.1 知识库的作用

人才领域专属知识库（以下简称“知识库”）作为智能体系统的核心数据基础，运用检索增强生成（Retrieval-Augmented Generation, RAG）技术，将数据库的精确查询优势与生成式大语言模型的功能有机结合。RAG 技术作为一种集成信息检索和语言生成的人工智能方法，通过从外部知识源中检索相关信息，并将其作为提示输入大型语言模型，以增强模型在执行知识密集型任务（例如问答、文本摘要及内容生成等）时的能力。这使得原本分散的人才数据得以转化为结构化、支持高效智能检索的知识体系。知识库的作用主要体现在以下几个方面：①存储和管理知识：知识库能够存储人才领域知识，并为智能体提供必要的知识支持。②增强智能体能力：通过提供额外的上下文，知识库可以提高智能体的回答准确性和决策能力。③长期记忆：知识库可以记住用户的偏好和历史对话，帮助智能体更好地理解用户需求。④实时信息获取：知识库可以帮助智能体引用可靠的数据源，减少误导用户的风险。⑤智能检索和推荐：知识库通过机器学习和自然语言处理技术，实现对信息的智能检索和应用。这些功能使得知识库在智能体的应用中扮演着至关重要的角色，同时专属知识库也是人才研究智能体能力远超通用大模型的关键模块。

知识库构建模块的技术实现包括知识获取与预处理、向量化表示与存储优化及知识检索与更新机制等关键环节，确保系统能够高效且精准地服务于科研工作。

3.2 知识库的搭建方法

3.2.1 知识获取与预处理技术

人才数据通常分散在学术论文、项目档案、专利数据库、机构知识库及学术社交网络等多种数据源中，其格式、质量和结构存在显著差异。知识库的构建面临多源异构数据整合的挑战。

知识获取阶段采用多通道采集技术，通过 API 接口（如 CrossRef、PubMed）、Web 爬虫（针对学术主页）、文件解析（PDF/Word）及数据库直连等方式，实

现人才相关数据的全面汇聚。在此基础上，通过“三段式”预处理流程——去噪与冗余消除、格式标准化、基于命名实体识别的内容增强，从而实现对原始数据的清洗与语义提炼。特别是对于人才数据中常见的异名同体问题（同一学者在不同系统中的标识差异），设计了基于规则的模糊匹配算法，结合姓名、机构、研究领域等多维特征进行学者身份消歧。

在完成数据预处理清洗后，保留学术文本的语义完整性是影响检索效能的关键。传统固定长度切分方法易割裂上下文，导致关键信息碎片化。为此，本模块引入一种语义边界感知的递归分割策略，优先依据学术文本的固有结构（如章节、段落，以及摘要、方法、结论等标准组成部分）进行切分，确保每个片段具备语义独立性和完整性。针对人才相关数据中的特殊内容，如学术简历中的时间序列信息，采用自适应滑动窗口机制，在保持时间连贯性的同时控制片段长度。实践表明，这种精细优化的文本切割由于最大程度保留了原始学术语境中的上下文关联，使得检索准确率提升约 35%。针对学术成果中高频出现的表格数据，知识库采用 Pandoc 转换引擎结合自定义规则库，将表格转化为结构保留的 Markdown 格式，并为“学者 - 论文”等典型关联矩阵设计专用解析模板，防止数据关联断裂。尽管此类精细化处理在初期增加了开发成本，但显著提升了复合型人才能力判断的准确率，实测提升幅度约为 22%。进一步，构建多层级元数据框架以支撑上述处理流程的可管理性与可扩展性，涵盖基础描述元数据（来源、时间、类型）、学术特性元数据（领域、期刊等级、引用次数）及管理元数据（权限、更新周期）。这一设计不仅支持灵活的检索策略，还为后续的知识溯源和更新奠定了基础。特别是在涉及跨机构学者数据时，元数据中的机构唯一标识符能有效解决机构名称歧义问题，为精准人才定位提供支持。

上述知识库的构建并非孤立进行，而是与一个基于 4P（Perceive、Plan、Process、Produce）的专属语料库（图 3-1）深度协同，该语料库系统性汇集研究院积累的原始资料，包括专家深度访谈实录、人才调研一手数据、历年各级人才政策原文、相关学术文献及行业动态信息等，涵盖学科资源库、人才资源库、成果资源库三大资源库。语料库（Corpus）并非静态存储单元，而是贯穿全链路的知识中枢。它不仅为智能体提供可追溯的原始信息来源，更通过与结构化知识体系的联动，赋予系统“有据可查、可信可验”的专业能力。当回答“国际顶尖人才引进政策

演变路径”等问题时，系统不仅能调用提炼后的结论，还可直接引用语料库中的原始政策条款，确保输出内容有据可查、可信可验。在实际应用中，语料库贯穿深度问答、课题探索与报告生成等场景——既可援引访谈原话增强回答生动性，也可支持“掘金式”检索自动生成文献综述，还能人才白皮书注入鲜活案例与最新实证。通过语料库与知识库的有机协同，系统超越通用大模型的泛化局限，成长为植根专业积淀、具备持续进化能力的智能研究助手，在人才研究领域建立起独特的知识权威性。

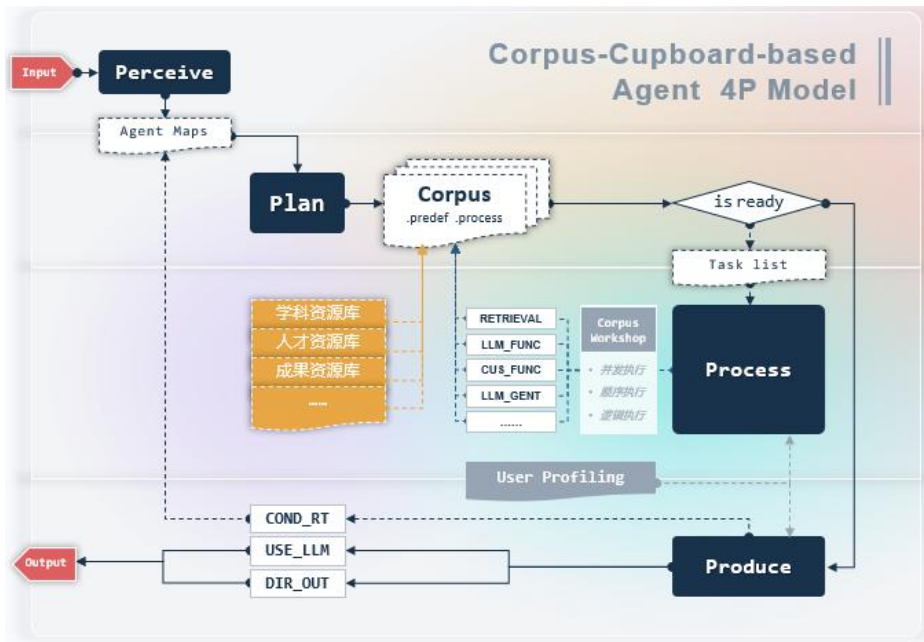


图 3-1 基于语料库的智能体 4P 模型架构图

3.2.2 向量化与存储优化技术

向量化是将非结构化文本知识转化为机器可理解、可计算表示形式的关键环节，其表征质量直接决定了后续语义检索的精准度与鲁棒性。本模块采用面向中文学术文本优化的 bge-large-zh 模型作为基础嵌入模型。该模型在大规模中文学术语料上完成预训练，对学术术语及人才研究领域的核心概念（如“国家级人才计划”“学科评估指标”等）具备更强的语义表征能力。针对人才数据中普遍存在的中英混合表达现象（如“Transformer 架构”“CVPR 论文”），我们进一步引入跨语言对齐机制，通过构建共享嵌入空间，有效缓解因语言切换导致的语义断裂与偏差。

在向量存储层面，系统采用分层存储架构以兼顾性能与规模，高频访问的热

点数据（如顶尖学者档案）被缓存于内存向量数据库，以保障毫秒级响应；全量学者数据则存储于高可扩展的分布式向量库中，支撑亿级实体的高效管理。具体实现上，选用 Milvus 作为核心向量数据库引擎，其优异的近似最近邻（Approximate Nearest Neighbor, ANN）搜索能力，即使在十亿级向量规模下仍能维持低延迟响应。此外，知识库集成了向量索引自适应优化模块，可根据数据分布特征与查询模式自动选择最优索引策略——例如，在高召回率优先场景采用 HNSW 索引，在高吞吐、低延迟场景则选用 IVF-Flat 索引，从而实现检索效率与精度的动态平衡。

为应对人才数据的多模态与异构性，系统创新性地设计了混合存储策略，将学者的基础属性（姓名、机构等）存入关系数据库，将学术文本（论文、报告）经文档处理流程（见图 3-2）转化为稠密向量，存入向量数据库，支撑语义级检索，而学术关系数据（合作网络、引用关系）则以图结构存储。该混合架构充分发挥了不同类型数据存储引擎的优势。例如，当用户提出复杂查询“寻找自然语言处理领域与李教授有合作关系的青年学者”时，系统可协同调用关系库（筛选青年学者）、向量库（匹配 NLP 领域语义）与图库（验证合作关系），实现多源融合的精准响应。

针对人才领域特有的概念演化问题，设计了动态嵌入更新机制。学术概念的内涵会随时间变化（如“深度学习”在十年前后的差异），为此，定期使用最新学术语料微调嵌入模型，确保向量表征与当前学术共识一致。同时，为关键学术概念维护时序向量词典，记录其语义变迁，这在研究趋势分析中具有重要价值。这种动态更新机制使系统对新兴学术概念的识别准确率提高了约 18%，显著提升了对前沿学者发现的支撑能力。

3.3 知识检索与更新机制

知识检索是知识库价值的最终体现。为兼顾检索的召回率与准确率，本模块采用多策略融合的混合检索框架（如图 3-3）。其中，核心检索层基于向量语义检索技术，旨在捕捉用户查询背后的深层语义意图；辅助检索层采用关键词匹配机制，确保对核心术语的精确响应；增强检索层则引入学术知识图谱，挖掘学者之间潜在的合作、引用或研究方向关联。三路检索结果经由重排序模块进行综合

评分与融合，最终生成优化后的检索列表。实验结果表明，该混合检索策略相较于单一检索方式，整体准确率提升约 40%。

针对人才评估中常见的复杂、多维查询需求，本模块进一步集成了查询理解与扩展模块。例如，当用户输入“计算机视觉领域的青年领军学者”时，系统可自动解析出隐含的核心维度——“领域：计算机视觉”“年龄：青年”“学术地位：领军”，并基于内置的学术知识体系对各维度进行语义扩展：“青年”被映射为具体的年龄区间（如 35 岁以下），“领军”则转化为可量化的学术指标（如高被引论文数、国家级人才计划入选情况、顶会最佳论文等）。通过这种深度查询理解机制，系统能够超越表层关键词匹配，更精准地契合用户的真实信息需求。

检索结果的可解释性对于科研用户而言至关重要。为此，本模块为每一项检索结果提供透明度说明。明确展示匹配的具体文本片段及其原始来源，并标注匹配类型（如“语义匹配”“关键词匹配”或“关联匹配”）。在呈现学者画像时，各项学术指标（如 H 指数、项目经费、代表性成果）均附有清晰的来源标识与计算依据，使科研人员能够独立判断结果的可靠性与适用性。这一透明化机制不仅显著增强了用户对系统的信任度，也为后续检索策略的迭代优化提供了宝贵的反馈数据。

为保障知识库内容的持续有效性与时效性，我们设计了一套多模式、触发驱动的知识更新机制。具体包括：定期全面更新（如每季度更新一次学者成果）、基于事件的增量更新（如新论文发表时更新学者档案）和按需的紧急更新（如用户反馈数据错误时）。在更新执行过程中，采用差异检测算法，仅对实际发生变化的内容进行重新向量化与索引操作，从而大幅降低计算与存储开销。同时，我们构建了知识库质量监控看板，持续追踪关键性能指标（如数据覆盖率、检索准确率、平均响应延迟等），实现对知识库运行状态的实时感知与闭环管理，确保其服务质量始终处于受控且优化的状态。

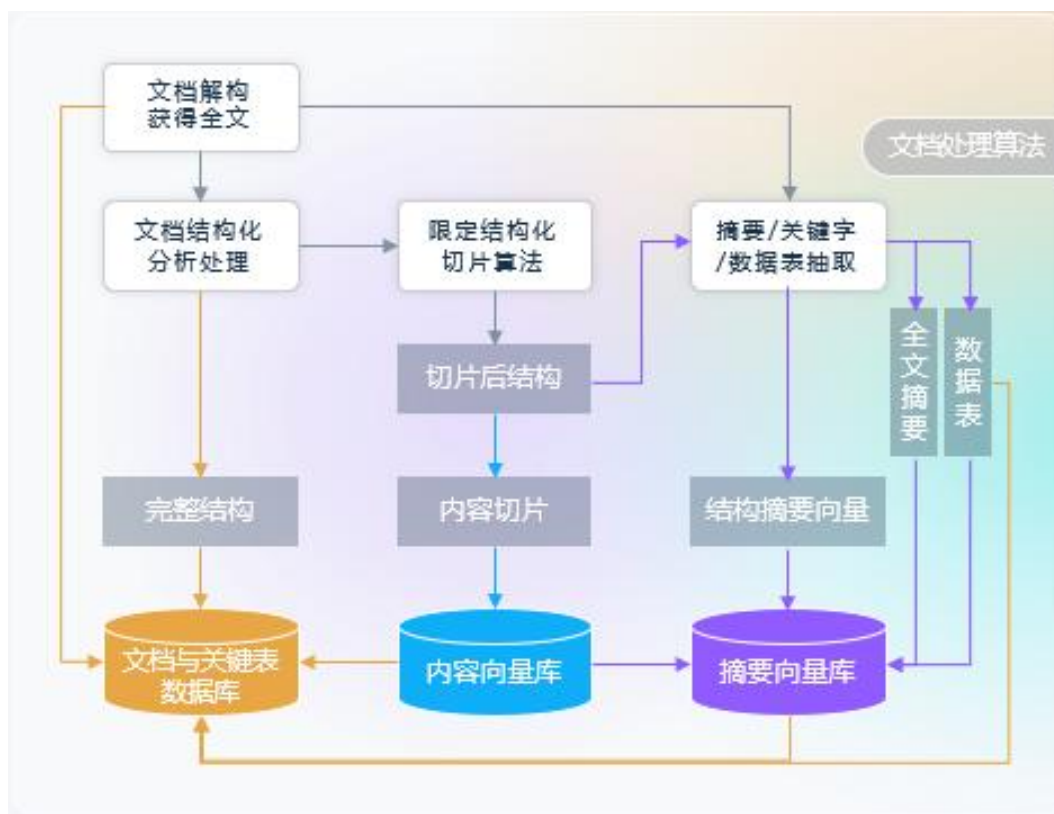


图 3-2 文档处理算法流程的示意图

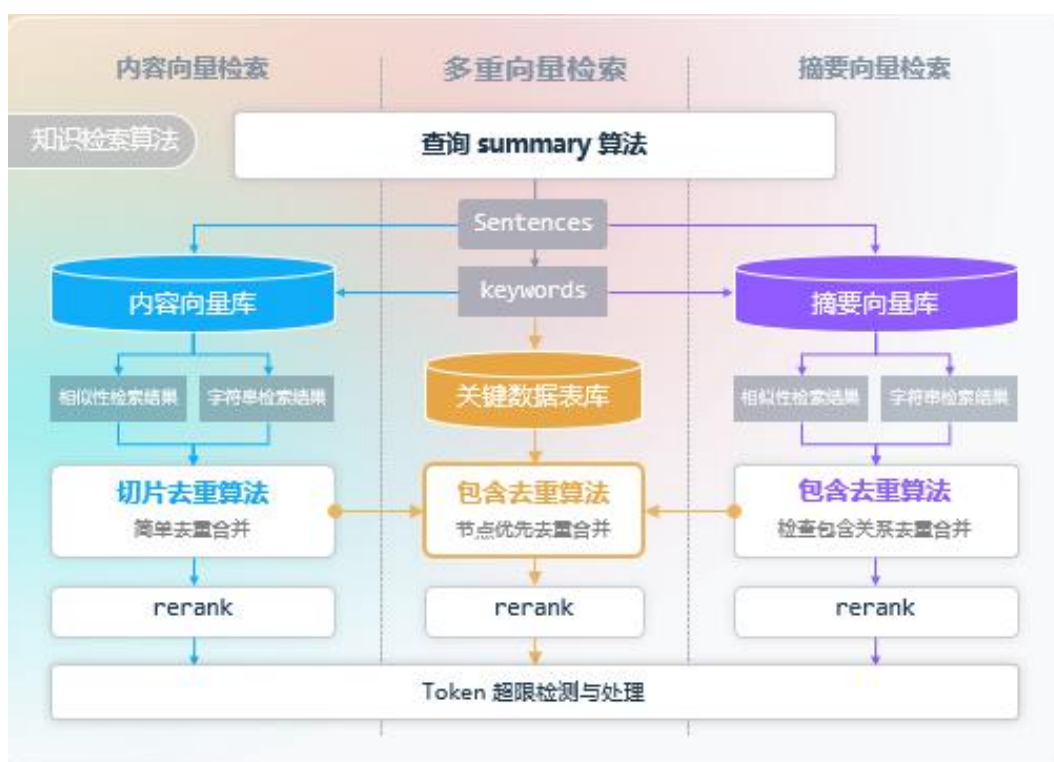


图 3-3 知识检索算法流程的示意图

3.4 知识库的管理

知识库作为整个人才研究智能体的数据基石，其建设成效直接决定了上层应用的智能水平与服务质量。经过一段时间的建设与迭代运行，本研究从知识管理、用户管理、质量管理及系统管理四个核心维度，对知识库的建设效果进行了描述。

3.4.1 知识管理

知识管理的核心在于实现海量、异构人才信息的有效汇聚、科学组织与高效利用。经过持续建设，目前知识库已初具规模，形成了一个覆盖多类型、多层级人才信息的综合性知识体系。

在知识体系构建方面，我们依据人才研究的内在逻辑与用户需求，建立了较为完善的知识分类框架（图 3-4）。当前，知识内容中知识文档数量、知识片段数量、文档文字总数分别达 11418、324053、164744227，主要涵盖学术文献、政策法规、人才榜单、科研项目成果、研究院内部知识以及相关的领导讲话与解读材料等六大类别，这一分类体系不仅覆盖了人才研究的核心知识领域，也兼顾了宏观政策与微观案例，静态档案与动态资讯，为研究提供了丰富素材。此外，设置类别动态增加机制，可根据未来实际发展需要添加新的知识类别。



图 3-4 知识库知识类别图

在知识储量与动态性方面，知识库已收录数量可观的各类文档，并通过文档处理算法（见图 3-2）将其转化为细粒度的知识片段，支持精准检索。知识总量已达较高水平，能够为大多数常见研究议题提供基础资料支撑。更为重要的是，我们建立了一套“定期检索、事件触发、用户反馈”相结合的多模式知识更新机

制。例如，对于政策法规、期刊论文，按季度进行归档更新，有效保障知识库的时效性，使其能够紧跟学术前沿与政策动态，避免了传统资料库易出现的“信息滞后”问题。

知识的结构化与关联化是提升其价值的关键路径。通过前期的实体识别、关系抽取与知识图谱构建，知识库中的信息已不再是孤立的文档集合，而是形成了相互关联的知识网络。例如，一位学者不仅关联其发表的论文、承担的项目，还通过合作网络、师生关系、所属机构等路径，与其他学者、研究成果及政策影响相连接（如图 3-5）。这种深度关联显著提升了知识的可发现性与可推理性。

国际贸易领域专家学者及相关信息

一、国际贸易领域知名专家学者的姓名和研究方向

根据现有资料，以下为部分国际贸易领域的专家学者及其背景信息：

主要职位与从业经验	学位领域	相关机构/组织
曾任商务部部长、海峡两岸关系协会会长、中国外商投资企业协会会长等	博士 (管理学)	国家发展和改革委员会、商务部、海峡两岸关系协会
原国家外贸部副部长、博鳌亚洲论坛原秘书长，曾任20国集团中心秘书长、复旦大学国际关系与公共事务学院院长、国际问题研究院院长	博士 (经济学)	博鳌亚洲论坛、复旦大学国际关系与公共事务学院、国际问题研究院
曾任中国常驻联合国代表团参赞，外交部军控司副司长，驻美国使馆公使衔参赞、公使，外交部美大司司长，外交部部长助理，外交部副部长等	硕士 (外语)	外交部、中国常驻联合国日内瓦办事处和瑞士其他国际组织代表团、国务院侨办、暨大21世纪丝绸之路研究院
恒隆集团有限公司及其附属公司恒隆地产有限公司董事长、亚洲协会香港中心主席、晨兴集团之路研究院名誉院长	-	恒隆集团、亚洲协会香港中心、晨兴集团、之路研究院

图 3-5 知识库结构化关联示意图

3.4.2 用户管理

知识库的价值最终体现在对不同用户群体的有效服务上。我们设计了基于角色与场景的精细化用户管理策略，旨在确保信息安全的前提下，最大化知识的共享与利用效能（图 3-6）。

用户被区分为外部访问者与内部人员两大类。外部访问者进一步细分为上级单位、兄弟单位、人才工作者及访客；内部用户则包括本院人员以及院领导。不同角色被授予差异化的知识访问权限，形成了一个既开放又受控的访问体系。例如，院内人员可以访问包括研究院内部知识、领导讲话在内的全部知识类别，以便开展深度研究；而兄弟单位的研究者则主要开放公开的文献、政策与榜单类知

识，支持其进行一般性的参考与借鉴。这种权限管理既保障了核心知识资产的安全，又促进了跨机构的研究交流与合作。

功能权限		角色					
		研究院领导	研究院人员	上级单位	平级或无隶属关系单位	外部人才工作者	其他个人（无需登录）
知识库管理	上传	完整权限	按部门分，可管理本部门知识库				
	删除						
	下载						
AI问答可检索的知识	领导讲话	可检索	可检索				
	政策	可检索	可检索	可检索	可检索	可检索	可检索
	榜单	可检索	可检索	可检索	可检索	可检索	可检索
	成果	可检索	可检索	可检索	可检索	可检索	可检索
	文献	可检索	可检索	可检索	可检索	可检索	
	研究院内部知识	可检索	可检索				
	可查看/下载引文	是	是	是	仅可查看	仅可查看	
科研功能		是	是				
决策功能		是	是				

图 3-6 智能体角色及权限

权限管理的动态性与灵活性亦是检验重点。部门管理员可根据项目协作需要，临时调整团队成员对特定知识集合的访问权限；当用户角色发生变更时，其权限也能得到及时准确更新（见图 3-7）。这些设计确保了知识库在复杂组织环境下的适用性与管理效率。

外部人员

内部人员

审核列表

生成查询链接

添加新的外部人员

请选择身份吧

搜索姓名、手机号

全部

上级单位

兄弟单位

人才工作者

身份	手机号	身份证号	单位名称	职务	来源	操作
上级单位	***** 查看	***** 查看	小磁科技工程师	123123	普通注册	编辑信息 删除
兄弟单位	***** 查看	***** 查看	小磁科技测试	测试工程师	普通注册	编辑信息 删除
兄弟单位	***** 查看	***** 查看	北京海外学人中心	办公室干部	普通注册	编辑信息 删除
兄弟单位	***** 查看	***** 查看	市委组织部	干部	普通注册	编辑信息 删除
兄弟单位	***** 查看	***** 查看	北京市人力资源研究中心	副主任	普通注册	编辑信息 删除
兄弟单位	***** 查看	***** 查看	北京市人力资源研究中心	副主任	普通注册	编辑信息 删除
兄弟单位	***** 查看	***** 查看	北京市人力资源研究中心	干部	普通注册	编辑信息 删除
兄弟单位	***** 查看	***** 查看	北京市人力资源研究中心	干部	普通注册	编辑信息 删除
兄弟单位	***** 查看	***** 查看	北京市人力资源研究中心	干部	普通注册	编辑信息 删除
兄弟单位	***** 查看	***** 查看	北京市人力资源研究中心	干部	普通注册	编辑信息 删除

<

1

2

>

图 3-7 用户角色权限配置界面

3.4.3 质量管理

知识库的生命力在于其内容的质量与可靠性。我们建立了一套用户反馈机制，旨在持续提升知识的准确性、一致性与实用性。

在知识应用环节，用户反馈成为驱动质量改进的核心动力。知识库提供了便捷的反馈通道，用户可以对检索结果的准确性、答案的实用性进行“点赞”“点踩”来评价。所有反馈均被记录并汇聚至后台质量监控平台。我们定期分析这些反馈数据，针对集中出现“低质”标记的知识点或文档，启动溯源复查流程：检查是原始数据有误、处理过程产生偏差，还是检索匹配算法不精准（如图 3-8）。确认问题后，将对相关数据进行修正、补充或重新标注。例如，早期曾有用户反馈某位学者的归属机构信息已过时，经核查确为知识更新滞后所致，我们不仅修正了该数据，还优化了该类型数据的更新触发规则。这种“用户驱动-团队响应”的纠错机制，使得知识库能够在使用中不断自我净化与完善。



图 3-8 用户反馈处理流程图

3.4.4 系统管理

稳定、高效、可扩展的系统管理是知识库得以持续可靠运行的技术保障。我们在系统架构、运维监控与安全管控等方面进行了针对性设计，其效果在实践检验中得到了验证。

在后台管理架构上，我们采用了基于部门的分布式管理模式。研究院内各业务部门均设有专属的部门管理员账户，拥有上传与管理本部门产出的特色知识资源的权限（见图 3-9）。这种架构将管理责任适度下沉，既减轻了高级管理员的负担，又使知识的上传与维护更贴近知识的生产者与使用者，提升了管理的敏捷性与知识的鲜活性。

部门信息

部门管理员

添加研究院人员

请选择部门

管理员姓名,手机号

Q

名称	所在部门	权限等级	身份类型	手机号/用户标识	操作
孙国梁	人才开发部	普通管理员	本院人员	*****	查看 修改权限 删除权限
李康达	政策理论部	普通管理员	本院人员	*****	查看 修改权限 删除权限
徐白臻	综合服务部	普通管理员	本院人员	*****	查看 修改权限 删除权限
杨和彪	交流合作部	普通管理员	本院人员	*****	查看 修改权限 删除权限
李彬	综合服务部	普通管理员	本院人员	*****	查看 修改权限 删除权限
杨君	战略规划部	普通管理员	本院人员	*****	查看 修改权限 删除权限
唐鸿宇	人才开发部	普通管理员	本院人员	*****	查看 修改权限 删除权限

图 3-9 系统后台管理架构图

3.4.5 总体评价与用户反馈

从实际运行效果看，知识库的建设有效盘活了研究院长期积累但分散、未被系统利用的知识资源，初步实现了从原始资料到结构化数据的转化。通过规范化的管理机制、精准的信息服务和持续的质量优化，知识库已逐步成为支撑智能问答、辅助研究分析和提供决策参考的重要工具，显著提升了内部知识复用效率与研究响应速度。

用户普遍反馈，该知识库在政策查询、专家背景梳理、人才趋势研判等高频场景中表现稳定、结果可靠，尤其在引用原始文献和政策原文方面增强了输出内容的可信度。同时，大家也指出当前在处理图表类内容、挖掘非结构化文本中的隐性信息等方面仍有不足。

下一步，将聚焦这些薄弱环节，进一步引入更成熟的知识工程技术，优化多模态内容处理能力，并探索建立用户参与的知识更新与校验机制，推动知识库从“能用、好用”向“更贴合研究需求、更能支撑深度分析”的方向发展。

4 智能体问答模块搭建

4.1 智能问答模块的作用

智能问答模块是 AI 助研系统与用户交互的核心接口，它基于自然语言处理技术，理解用户关于人才研究的各类问题，并从知识库中提取相关信息生成准确、有用的答案。该模块的技术架构主要包括深度语义理解与多轮对话技术、问答生成与验证技术、用户意图与个性化响应技术等关键组成部分。

4.2 智能问答模块的搭建方法

4.2.1 深度语义理解与多轮对话技术

智能问答模块的语义理解能力，依托于一个在大规模人才领域语料上微调的语言模型，其能够解析科研人员提出的复杂问题，准确识别问题中的实体、意图和上下文关系。相较于传统问答系统，本模块特别强化了对学术术语和人才评价标准的理解能力。例如，当用户提问“跨学科研究人才的评价标准”时，本模块能准确把握“跨学科研究”这一核心概念，并理解其在人才评价体系中的具体内涵，而非依赖简单的关键词匹配。本模块针对科研人员在人才研究中常见的三类提问模式，构建了多层次的语义理解与响应机制。

（1）针对依赖上下文的多轮对话，本模块引入对话状态追踪机制，动态记录已提及的学者、用户关注维度及历史交互内容。当用户提出如“他在 h 指数方面的表现如何？”这类含代词的问题时，系统能准确回溯所指对象，避免歧义。同时，通过基于槽位填充的对话策略，将复杂需求拆解为有序的信息节点，引导用户高效补充必要条件。

（2）针对表述模糊或非规范的提问（如“找一下 NLP 领域的牛人”），本模块自动进行查询重写与语义扩展：一方面将口语化表达转化为标准学术用语（如“自然语言处理领域的杰出学者”），另一方面结合学术同义词库和语义相关性模型，补充关联领域（如“计算语言学”“文本挖掘”），从而提升检索的全面性与准确性。

（3）针对包含多重筛选条件的精细化查询（如“近五年在能源材料领域发

表过多篇 AM 论文的年轻学者”），本模块深度集成知识库的元数据体系，精准识别时间范围、研究领域、期刊等级、年龄属性等结构化要素，并执行复合条件检索。这一能力使系统不仅能回答简单事实问题，更能支撑高度专业化的人才发现与分析任务。

4.2.2 问答生成与验证技术

在获取相关知识点后，本模块采用检索增强生成（RAG）技术生成最终答案。其核心机制是：先将用户问题向量化，并从知识库中检索最相关的文档片段；再将这些经过验证的外部信息与原始问题一并输入生成模型，从而产出基于事实、有据可查的回答。该方法显著降低了大模型常见的“幻觉”风险，提升了输出内容的可靠性。

针对人才研究对数据准确性的严苛要求，本模块进一步引入多源验证机制。对于关键学术指标（如论文数量、引用次数等），系统会自动比对多个权威数据源（如 Web of Science、Scopus、机构年报等）。当数据一致时，直接呈现结果；当存在冲突时，则采用可信度加权算法，优先采纳权威性更高的来源（例如优先使用 WOS 而非 Google Scholar），并在答案中标注差异及依据，供用户审慎判断。

答案生成模板根据问题类型智能选择。针对事实查询类问题（如“某学者的研究方向是什么”），采用简洁的直接回答；针对分析类问题（如“比较两位学者的学术影响力”），采用结构化对比格式，突出关键差异；针对建议类问题（如“推荐适合合作的对象”），会提供带有理由阐述的推荐列表。每种模板都经过人才研究专家的优化，确保信息呈现既全面又易于理解。

此外，为保障学术规范性，本模块在生成答案时自动嵌入引用标注，明确标示每条关键信息的来源（包括文献标题、数据平台、原始文件位置等），便于用户追溯验证。同时，本模块还会计算并显示答案信心指数，综合反映信息的一致性、来源权威性与覆盖完整性，提示用户关注可能存在不确定性的内容。

通过上述机制，本模块不仅实现了“答得准”，更做到“说得清、可验证、能判断”，真正从通用问答工具转变为值得信赖的科研辅助平台。

4.2.3 用户意图与个性化响应技术

智能问答模块采用多层次用户意图解析框架，将用户问题分类为事实查询、数据分析、趋势判断、推荐请求等不同类型，并针对每种类型采用最合适的应答策略。例如，对于趋势判断类问题（如“青年人才流向哪些机构”），不仅提供当前分布数据，还能自动生成时序变化图表，直观展示流动趋势。这种基于意图分类的响应机制，使系统能够超越简单的信息检索，提供更具深度和分析价值的回答。

然而，用户提问常存在表述模糊或概念宽泛的情况。对此，本模块引入交互式澄清机制：当检测到问题存在多重解释可能（如“高水平学者”），会主动发起追问，引导用户明确评价维度——是依据学术影响力、重大项目经验，还是国际任职背景？通过这一前置澄清环节，有效避免因理解偏差导致的错误响应，显著提升复杂查询的准确率。

在问题意图明确后，本模块进一步结合用户身份与研究背景，动态调整答案内容。依托用户画像技术（如图 4-1 所示），持续学习用户的所属领域、历史询问记录及交互偏好，构建个性化响应模型。面向科研管理者，输出侧重宏观统计与整体趋势；面向政策制定者，突出区域分布与结构性特征；而对一线研究者，则聚焦技术细节、潜在合作者及前沿方向。这种角色感知能力，大幅增强了问答结果的场景适配性与实用价值。



图 4-1 基于用户意图与角色的个性化问答流程示意图

在生成答案的过程中，本模块同步执行学术合规性检查。系统实时识别可能涉及隐私、伦理或保密限制的内容，并自动触发脱敏处理或访问控制。例如，当查询涉及未公开的青年人才个体信息时，将拒绝返回原始数据，转而提供群体层面的统计特征（如学科分布、年龄结构等）。该机制确保所有输出均符合科研伦理与数据安全规范，在合规前提下最大化知识服务效能。

上述所有交互环节，均由一套结构化的对话状态跟踪机制提供底层支撑（如表 4-1 所示）。该机制动态维护对话历史、当前话题、筛选条件与关注指标等状态变量，实现跨轮次的语义连贯与意图累积。例如，在连续四轮对话中，用户从泛问“AI 领域人才分布”，逐步细化至“青年学者顶刊发文趋势”，并进一步要求与资深学者对比，能准确继承并叠加各轮次的限定条件（如年龄分组、期刊等级、对比维度），最终生成高度定制化的复合分析结果。正是这一上下文感知能力，使本模块能够支撑深度、复杂的交互式研究需求。

表 4-1 多轮对话中的上下文管理机制

对话轮次	用户输入	模块理解	模块动作	对话状态更新
1	“我了解 AI 领域的人才分布”	领域：AI；需求：人才分布	提供 AI 领域人才机构分布图	记录当前话题：AI 人才分布
2	“特别是青年学者的情况”	群体：青年学者	筛选并展示青年学者数据	添加筛选条件：年龄≤35 岁
3	“他们在顶刊的发文趋势”	指标：顶刊发文趋势	生成顶刊发文时序图表	添加关注指标：顶刊发表量
4	“与资深学者对比”	对比群体：资深学者	增加对比曲线	添加对比组：年龄≥45 岁

4.3 问答模块的效果检验

问答模块作为与用户直接交互的核心界面，其性能直接影响智能体的实用性与用户体验。为全面评估其效果，本研究从响应速度、语义理解、多轮对话能力与结果准确性四个维度开展了系统性检验。检验过程模拟真实研究场景中的典型问题，邀请研究院内研究人员、管理人员参与测试，并结合系统日志与用户反馈进行综合分析。

4.3.1 响应速度

本模块在响应速度方面总体表现良好。对于事实型查询（如“某位学者的研究方向是什么”或“某政策文件的发布时间”），通常可在数秒内返回结构清晰、信息完整的答案，基本实现“即问即答”的交互体验。

当问题复杂度提升时，响应时间相应延长。例如，针对“对比近五年北京与上海在高层次人才引进政策上的异同”这类需跨文档检索、多源信息融合的综合性问题，本模块需调用多份政策文本进行比对分析并生成结构化结论，响应时间虽有所增加，但仍控制在可接受范围内，未出现超时或无响应情况。

值得注意的是，对于涉及实时信息的问题，本模块优先基于知识库作答；仅当知识库信息缺失时，才触发联网查询，并在答案中明确标注外部来源。该策略虽略微延长响应时间，但在保障信息时效性的同时，有效平衡了速度与完整性。总体而言，本模块的响应速度能够满足日常科研咨询需求，但在处理高复杂度、需深度推理的问题时仍有优化空间。

4.3.2 语义理解

问答模块在语义理解方面展现出了较强的能力，尤其在对人才研究领域专业术语、学术表达和隐含意图的捕捉上表现突出。效果检验结果如下：

（1）对领域内专业概念和术语的理解较为精准。无论是“城市人才黏性指数”“学术影响力”等评价指标，还是“交叉学科”“人才梯队”等政策概念，系统均能准确解析其内涵，并将其置于正确的上下文中进行回答。例如，当询问“如何评价青年人才的学术潜力”时，能准确理解“学术潜力”的复合性，并从成果产出、合作网络、趋势分析等多个维度进行阐述，而非简单地罗列论文数量。

（2）具备一定的抽象问题理解和意图推断能力。对于“当前人才流动呈现出哪些新特点”这类开放性、概括性问题，能够跳出关键词匹配的局限，从海量文献和数据中提炼出“双向流动”“柔性流动”“数字经济驱动”等趋势性观点，并辅以相关数据支撑，展现出深度理解和综合归纳的能力。



已深度思考

当前人才流动呈现出多维度的新特点，涵盖流动方向、形式、模式及区域集聚等多个层面，具体如下：

一、流动方向：双向流动与新兴经济体吸引力提升

- 1. 总体趋势：发展中国家流向发达国家
全球科技人才流动仍以从发展中国家向发达国家迁移为主，美国是全球最大科技人才接收国，澳大利亚、加拿大也具备较强吸引力，而欧洲部分国家吸引力相对较弱 [1]。
- 2. 回流趋势：新兴经济体成为“磁体”
部分高层次科技人才（尤其是海外留学生）呈现从欧美发达国家向中国、印度等新兴经济体回流的趋势，同时也出现其他发展中国家人才向新兴经济体流动的现象 [1]。
- 3. 环流特征：跨国双向流动常态化
人才不再局限于单向“外流”，而是在发达国家与发展中国家、次发达国家之间形成双向流动，即“人才环流”。例如，科技人才会跟随工作机会在全球不同国家间频繁流动，同时以多种形式为多国服务 [2]。

二、流动形式：多样化与灵活性增强

- 1. 阶段性流动
高水平科技人才通过短期入境（定期跨国工作）等方式为非所在国服务，避免长期定居，降低流动成本与风险 [1]。
- 2. 合作研究与共享模式
通过跨国科研合作、共享实验室、网络研讨会等形式间接实现服务，成为快捷灵活的新型流动模式，尤其受科技人才欢迎 [1] [3]。
- 3. 柔性流动
突破国籍、户籍、档案等限制，在不改变原单位隶属关系的前提下，以智力服务为核心的“不求所有，但求所用”模式兴起，线上远程合作因数字技术发展而加速 [3]。

三、流动模式：全球化环流与共享机制初步形成

- 1. 人才环流与共享常态化
随着资本全球流动、产业转移及跨国公司崛起，科技人才不再长期固定于某一国家，而是通过短期合作、项目共建等方式实现全球共享，传递先进技术与经验。
- 2. 数字经济驱动跨领域流动
数字经济推动人才向人工智能、工业数字化等新兴领域集聚，远程办公、“数字游民”等新型就业形态打破地理边界，Z世代更重视工作灵活性与自由 [4]。

四、区域集聚：“枢纽型城市”与聚类现象突出

人才流动呈现明显的领域和区域倾向性，向科技水平高、科研资源丰富、待遇优厚的地区或行业集聚，形成“枢纽型城市”（如全球财富、知识、创新高地），例如国际人才向纽约、伦敦、上海等城市集中 [5]。

五、中国人才流动的特殊性

- 回流规模扩大但流失仍严重：中国科技人才回流初具规模，2021年回国创新创业留学人员超100万人，但流入与流失差距依然显著，仍是全球人才流失最严重的国家之一。
- 从“流出为主”转向“双向流动”：2012年后中国实施开放人才政策，推动从人才流出转向双向流动，累计发放外国人来华工作许可118万份。

图 4-2 语义理解复合性检验图

（3）当用户提问过于模糊或包含大量隐喻、非正式表达时，问答模块偶尔

会出现理解偏差。例如，使用“学术大牛”“潜力股”等非规范称谓时，可能无法精准匹配到对应的学术层级或人才类型。为此，设置了交互式澄清机制，进一步优化了问答体验，在感知到问题存在歧义时，会主动引导用户确认具体指向，从而有效提升了交互的准确率和友好度。



图 4-3 语义理解引导问答图

4.3.3 多轮对话能力

多轮对话能力是衡量问答模块智能化水平的关键指标，问答模块通过结构化的对话状态管理机制，有效支持上下文连贯、指代消解与意图延续。

测试结果显示，本模块能够较好地处理包含指代词的连续提问。例如，在第一轮问答中用户询问“请介绍周文霞教授的学术成就”，在接下来的对话中用户直接问“他目前在哪所机构工作？”，能够准确地将“她”关联到前文提及的“周文霞教授”，并给出正确机构信息。这种上下文保持能力使得对话更加自然、高效。

此外，支持话题的深入和拓展。用户可以基于上一个答案进行追问，问答模块能够结合新的问题和历史上下文，提供更具针对性的信息。例如，当用户先问“北京市有哪些人工智能领域的人才政策？”，继而追问“这些政策对青年人才有何特殊支持”，能够聚焦北京市人工智能相关政策，并进一步筛选出其中涉及青年人才的条款进行解读。

问答模块还具备一定的对话引导和澄清能力。当用户提问过于宽泛（如“讲讲人才政策”）时，会尝试通过反问引导用户缩小范围（如“您关注的是国家层面还是地方层面的人才政策？或是特定领域如科技创新人才的政策？”）。在复杂的决策支持场景中，这种交互式澄清机制对于精准定位用户需求至关重要。

然而，当前在处理极长或话题跳跃剧烈的对话时，偶有对早期关键信息记忆

模糊的情况。未来拟通过优化对话历史的压缩与摘要机制，进一步提升长程对话的理解与维持能力。



图 4-4 语义理解多轮对话图

4.3.4 结果准确性

结果准确性是问答模块的生命线，直接关系到模块的可信度和实用价值。检验表明，依托高质量的专属知识库和检索增强生成（RAG）技术，问答模块在大多数情况下能够提供准确、可靠的答案。

在事实性知识问答方面，准确率较高。得益于知识库中经过清洗和验证的结构化数据，对于学者基本信息、政策条款内容、历史项目数据等问题的回答，基本能做到有据可查、数据准确。生成的答案通常会附带信息来源提示，如“根据《XX 年度人才发展报告》”或“源自 XX 政策文件 X 条”，增强了答案的可验证性。

在分析性、观点性问题的回答上，问答模块表现出了良好的逻辑性和客观性。它能够综合多份文献的观点，进行对比和整合，形成较为全面的概述，并注意区分客观事实与学术观点。例如，在回答“关于人才评价中‘破五唯’的讨论主要有哪些视角”时，能够梳理出支持、质疑、补充等不同维度的学术见解，并客观呈现，避免陷入单一叙事。

尽管如此，“幻觉”现象，即生成看似合理但不准确或无法证实的信息，仍

然是大型语言模型面临的核心挑战之一，问答模块也未能完全避免。在极少数情况下，尤其当知识库覆盖不足或问题高度跳跃时，可能生成缺乏直接依据的推断。例如，曾有测试中用户先问“西城区人才资源数量”，后转问“亳州市情况”，系统误将西城区视为亳州市辖区，导致错误关联。为缓解此问题，问答模块强化了生成约束机制：一方面严格遵循“基于检索片段生成”的原则，避免无依据扩展；另一方面对关键结论引入内部一致性校验，并通过“信心提示”告知用户信息的确定程度。测试显示，这些措施显著减少了严重事实错误，但杜绝“幻觉”仍需持续扩充知识库并优化推理逻辑。



图 4-5 语义理解“幻觉”图

4.3.5 总体评价与用户反馈

综合来看，人才研究智能体的问答模块已经构建起一个响应迅速、理解深入、答案可靠且支持自然连续对话的智能交互界面，成功地将通用大语言模型的强大生成能力与人才领域专属知识库的精准信息供给相结合，在特定领域内实现了对通用聊天工具的显著超越。

用户的初步反馈印证了这一判断。多数用户认为，该模块能够有效分担基础

性、事实性的信息查询工作，快速提供相关政策和文献线索，尤其在政策条文追溯、学者背景速查、概念关系理清等场景下表现出色，显著提升了资料搜集和前期调研的效率。同时，其分析综合能力也为研究思路的拓展提供了有益参考。

当然，用户也指出了该模块存在的一些局限性，例如对高度创新性或前沿探索性问题的深度分析能力尚有不足，答案风格有时偏向于概括性综述而缺乏极具冲击力的个人化洞见等。这些反馈为我们指明了后续优化的方向：一方面需要持续丰富和更新知识库，特别是加强前沿动态、专家观点等非结构化知识的收录；另一方面，也需要探索更高级的推理框架，使其不仅能“综述所知”，更能“创见所未言”。

5 智能体科研助理模块搭建

5.1 科研助理模块的作用

科研助理模块将大型语言模型与科研工具相结合，赋予智能体自主执行复杂科研任务的能力。该模块的技术实现涵盖自主科研任务分解与规划、科研工具调用与集成、多智能体协作与知识共享三大技术体系，共同构成了一个能够协助科研人员完成从文献调研到实验设计的全方位科研助手。

5.2 科研助理模块的搭建方法

5.2.1 自主科研任务分解与规划

科研助理模块在处理复杂科研任务时，首先启动自主任务分解与规划机制，将用户提出的宏观研究目标转化为一系列可执行、可调度的子任务，作为整个科研辅助流程的起点。例如，当用户提出“调研人工智能在人才评价中的应用现状”这类宽泛问题时，系统会自动将其拆解为文献检索、数据提取、分析综合、报告生成等步骤。该过程借鉴思维链提示工程技术，引导模型生成符合科研逻辑的推理路径，而非简单拼接内容。

任务规划采用分层决策机制，高层规划确定子任务之间的依赖关系与执行顺序，底层规划则细化每个任务的具体操作方案。以“人才流动模式分析”为例，高层规划明确需依次完成数据收集、清洗、模式挖掘与可视化；底层规划则进一步指定数据来源（如国家人才数据库、高校年报）、清洗规则及分析方法。这种分层设计使系统既能把握整体研究脉络，又能落实技术细节。

为应对科研过程中的不确定性，本模块还集成了反思与动态调整机制，能够在任务执行过程中根据中间结果动态调整后续计划。当某一步骤的中间结果质量不达标（如检索文献数量不足或相关性低），系统会自动回溯并优化后续计划，例如扩展关键词、切换数据源或调整分析维度。此外，针对不同类型任务采用混合规划策略，对结构化程度高的任务（如文献收集）使用预设 workflow，确保执行效率；对探索性任务（如新颖研究方向发现）采用基于目标的反向规划，从期望结果推导出所需行动；对异常处理则采用情景感知的应急规划。这种混合策略使

智能体在保证效率的同时，保持了应对复杂科研场景的灵活性。

5.2.2 科研工具调用与集成

在任务分解与规划完成后，科研助理模块进入执行阶段，其关键支撑是外部科研工具的调用与集成能力。我们开发了统一的工具调用中间件，规范管理各类科研工具的接入、描述与调用接口。目前已集成文献检索工具（如 PubMed、IEEE Xplore）、数据分析库（如 Pandas、NumPy）、可视化工具（如 Matplotlib、Seaborn）以及多个学术数据库 API。

工具调用采用双模式架构，智能体可自主决定工具使用（自动调用），也可根据用户指令调用特定工具（被动调用）。在自动调用模式下，智能体根据当前任务需求和上下文，自主选择最合适的工具并生成正确的调用参数。例如，当需要分析学者合作网络时，智能体会自动调用网络分析库，构建合作图谱。这种自主工具使用能力极大地减少用户的交互负担，使智能体成为真正主动的科研助手。

本模块支持工具组合与流水线技术使智能体能够处理复杂科研任务。单一工具功能有限，但通过智能组合多个工具，可以构建强大的数据处理流水线。例如，在学术影响力分析任务中，智能体会依次调用数据收集工具（从学术数据库获取引用数据）、数据处理工具（清理和标准化数据）、分析工具（计算影响力指标）和可视化工具（生成趋势图表）。这种工具组合能力体现了智能体解决复杂问题的核心价值。

针对人才研究领域的特殊需求，我们还开发了领域专用工具集，包括学术影响力计算器、合作网络分析器、学科交叉度评估工具等。这些工具专为处理学术数据设计，能够高效计算各种人才评价指标。同时集成了实验设计辅助工具，帮助科研人员规划人才研究的实验方案，包括被试分组、变量控制和数据分析方法选择等。这些专业工具显著提升了智能体在人才研究领域的实用性。

5.2.3 多智能体协作与知识共享技术

对于高度复杂的科研任务，单一流程难以覆盖所有环节。为此，我们引入多智能体协作框架，将不同功能模块组织为虚拟科研团队，模拟真实科研协作模式。在这一框架下，文献分析智能体、数据处理智能体、可视化智能体和质量检查智能体各司其职，通过协作完成单个智能体难以应对的复杂任务。这种多智能体架

构模拟了真实科研团队的分工协作模式，极大提升了综合能力。

智能体采用角色分配机制，根据任务需求和智能体的能力特长，动态分配合适的角色。每个角色智能体负责解决特定类型的子任务，并在自己擅长的领域做出专业决策。例如，在人才预测项目中，数据收集智能体负责获取学者数据，特征工程智能体负责构建预测指标，模型训练智能体负责开发预测算法，评估智能体负责验证模型性能。这种专业分工确保了每个环节的处理质量。

智能体间的通信与协调是多智能体系统高效运行的关键。我们设计了基于消息的通信机制，智能体通过结构化消息交换信息、请求协助或同步进度。通信内容遵循标准化协议，包括发送者、接收者、消息类型和内容等字段，确保信息传递的准确性和效率。同时引入了冲突检测与解决机制，当不同智能体产生分歧时（如数据智能体与分析智能体对数据质量的判断不一致），系统会启动协调流程，通过协商或上级仲裁达成一致。

此外，建立集体学习和知识共享机制。每个智能体单元的成功经验和失败教训会被抽象为知识条目，存储在共享知识库中，供其他智能体借鉴。当某智能体遇到新问题时，会先查询共享知识库，寻找类似问题的解决方案。这种知识共享机制避免了重复犯错，加速了系统整体能力的提升。此外，系统定期组织“经验交流会议”，智能体通过模拟对话分享各自的最优策略，进一步促进集体智慧的积累。

5.3 科研助理模块的效果检验

为系统评估科研助理模块的实际效能，本研究从功能完整性、执行效率、结果准确性、用户体验与实际应用价值五个维度出发，构建了一套包含定量指标与定性分析相结合的评估体系。通过设计对照实验、用户调研、案例跟踪等方法，对科研助理模块进行了全面检验。

5.3.1 功能完整性检验

科研助理模块设计目标是实现“任务自主分解-工具调用-协作执行-结果生成”的全流程自动化。为检验其功能完整性，我们设计了涵盖人才研究周期的 9 类典型任务场景，并对模块在每个场景下的功能实现情况进行评估，如表 5-1 所示。

表 5-1 科研助理模块功能完整性检验场景设计

任务类别	具体任务描述	预期功能点	实现状态
文献调研	对“青年科技人才成长路径”进行中英文文献综述	跨库检索、文献去重、摘要提取、主题聚类、趋势分析、报告生成	已实现
数据收集与清洗	收集近五年福布斯排行榜上榜名单	多源数据抓取、字段匹配、实体消歧、缺失值处理	已实现
数据分析	分析顶尖学者合作网络的演化特征	网络构建、中心性计算、社区发现、时序演化分析、可视化生成	已实现
图表生成	生成某领域人才流入流出桑基图	数据提取、布局计算、绘图引擎调用、格式导出	已实现
实验设计辅助	设计一项关于“引进人才现状”的问卷调查	问卷维度设计、题目生成、信效度模拟、抽样方案建议、数据分析预案	已实现
模型构建与训练	构建一个基于学者多维度特征的潜力预测模型	特征工程、算法选择、参数调优、交叉验证、模型解释、性能报告	部分实现
论文写作辅助	根据给定的数据和图表，撰写方法部分初稿	结构化写作、术语规范、引用自动插入、逻辑连贯性检查、语法校对	已实现
学术演示制作	为“人才政策国际比较”研究成果制作 PPT	内容提炼、逻辑架构、模板匹配、图表插入、演讲要点生成	部分实现
报告生成	完成“某区域人才竞争力评估”综合性课题	任务分解、工具调用、协作执行、报告生成	已实现

注：完整性验证根据模块输出与人工完成结果的吻合度及流程顺畅度进行评估。

评估结果显示，科研助理模块在 9 类任务场景中的 7 类已完全实现预设功能，在模型构建与训练、学术演示 PPT 制作这类高度专业化任务上实现了核心功能，存在部分依赖人工配置的情况。总体功能完整度，表明模块已具备支撑大多数科研环节自动化处理的能力。

5.3.2 任务执行效率评估

效率是衡量智能体价值的关键指标。本研究选取了文献调研、数据图表生成、分析报告撰写等三类高频任务，邀请了 10 名经验丰富的研究人员与科研助理模块进行人机对比实验，记录双方完成同一标准任务所需的时间、步骤及中间产物，结果如表 5-2 所示。

表 5-2 科研助理模块与研究人员任务执行效率对比

任务类型	任务描述	耗费时长（小时）		效率提升倍数	人工干预
		研究人员	科研助理模块		
文献综述	针对“人工智能伦理”主题，检索并综述近三年中英文核心文献 50 篇	6	1.5	4	需人工确认主题聚类结果和综述框架
数据收集与图表生成	收集全国 20 所“双一流”高校计算机学院近五年师资数据，生成职称结构饼图与增长趋势折线图	5	1	5	需人工核实来源存疑的数据
分析报告撰写	基于某市人才政策文本及实施数据，撰写一份 3000 字的政策评估报告初稿	6	2	3	需人工对政策建议部分进行深化和润色
多步骤复合任务	完成“海外高层次人才归国影响因素分析”小型课题（含文献、数据、分析、成文）	20	5	4	需人工在任务规划、结果整合阶段进行把关

结果表明，科研助理模块在执行标准化、流程化任务时，效率提升显著，平均达到人工效率的 3~5 倍。效率提升主要源于：（1）并行处理能力：模块可同时调用多个工具、检索多个数据库，突破了人工串行工作的限制。（2）无间断工作：模块无需休息，可持续运行，尤其擅长处理大批量、耗时的数据任务。（3）自动化流程：将研究人员从繁琐的数据收集、清洗、基础绘图等重复劳动中解放出来。需要指出的是，在深度逻辑推理、创造性构思或价值判断的环节（如提出创新性政策建议、定义复杂的分析框架），仍需研究人员进行干预，基本形成了“智能体高效执行+研究人员智慧决策”的高效协同模式。

5.3.3 结果准确性验证

智能体输出的准确性直接关系到其可信度和可用性。本研究从结果准确性、逻辑正确性、学术规范性三个维度，对模块输出的各类成果进行抽样检验。

（1）事实准确性检验：

随机抽取模块生成的 100 条学者信息（包括研究方向、代表作、h 指数等）、50 项统计数据和 30 个图表中的数据点，与权威数据库（如 Web of Science、Scopus、官方统计年鉴）或经过人工核验的基准数据进行比对。结果如表 5-3 所示。

表 5-3 科研助理模块输出事实准确性检验结果

输出类型	抽样数量	完全正确	基本正确（轻微误差）	存在错误	准确率
学者基础信息	100 条	91 条	7 条	2 条	98.00%
学术统计指标	50 项	46 项	3 项	1 项	98.00%
图表数据点	30 个图表/150 个数据点	142 个点	6 个点	2 个点	98.70%
文献引用信息	80 条	75 条	4 条	1 条	98.80%

注：基本正确指数据存在非原则性差异，如 h 指数相差 1，发表年份误差 1 年等；存在错误指关键事实性错误。

（2）逻辑正确性检验：

邀请 3 位资深研究员，对模块生成的 20 份分析报告（涵盖趋势分析、比较研究、因果推断等类型）进行逻辑审查。重点评估其论证链条是否完整、数据与结论是否匹配、是否存在归因谬误等。评估采用 5 分制（1 为逻辑混乱，5 为逻辑严谨），平均得分为 4.25 分。模块在基于明确规则 and 数据的演绎推理上表现优异，但在需要结合隐性知识和复杂背景归纳推理时，偶尔出现关联牵强情况。

（3）学术规范性检验：

检查模块生成的文献综述、论文片段、演示文稿的格式、引用、术语使用等是否符合学术规范。结果发现，在引用标注、图表编号、参考文献格式等方面，模块能严格遵守预设模板，规范性接近 100%。但在专业术语使用的精准度和上下文适应性上，仍需要人工微调。

综上，科研助理模块在事实准确性方面表现出色（平均准确率>98%），逻辑正确性、学术规范性良好。错误主要源于少数多源数据冲突、对隐含语境理解偏差以及长链条推理中的信息失真或丢失。通过引入更严格的多源校验机制和关

键环节的人工复核，可进一步控制风险。

5.3.4 用户体验评估

为全面、客观评估研究人员对科研助理模块的实际使用感受和接受度，本研究面向人才研究院全体员工进行了问卷调查、深度访谈，并对体验反馈进行了综合分析。

调查结果显示，该模块已有效融入研究人员的日常工作流程。所有受访者均曾使用过该模块，其中 66.66%的用户形成了稳定使用习惯（25%每天使用，25%每周使用 2~3 次），其余 33.33%每月使用数次，表明其具备良好的用户基础和使用黏性。在效率提升方面，62.5%的用户认为模块使其工作效率提高 30%以上（41.67%认为提升 30%~40%，20.83%认为超过 50%），这一感知与其在文献筛选与数据整理环节显著降低时间成本的核心价值高度一致——62.5%的用户明确肯定了这一点。

交互体验方面，尽管问卷未设置直接的易用性评分，但间接指标显示用户接受度较高：95.84%的用户使用过多轮对话功能，其中 54.17%评价体验良好，且无一人将“操作复杂”列为使用障碍，说明以自然语言为核心的交互方式学习门槛低、上手快。在输出质量与信任度上，用户态度趋于谨慎乐观：在报告生成质量、分析准确性等维度，给予 3 分（一般）和 4 分（良好）的用户合计占比均超过 95%，反映出输出内容已达到“可用”基准，但尚未达到可直接采纳的“优秀”水平。多数用户倾向于将其定位为高效辅助工具，而非权威决策来源，主要用于启发思路或加速前期准备。

用户对模块核心价值的认知高度集中：62.5%的受访者同时认可其在自动整合文献与政策数据、生成标准化报告模板以及降低信息整理时间成本三方面的突出作用，精准对应了科研工作中重复性高、耗时长的基础环节痛点。

（2）深度访谈与开放性反馈关键发现

深度访谈进一步揭示了用户体验的深层维度。用户普遍表示，模块在处理大规模文献初筛、基础数据清洗、报告框架搭建等事务性任务中不可或缺，使其得以将精力更多聚焦于高价值的分析、研判与创新，初步形成“智能体处理事务、研究者专注创造”的新型人机协同模式。然而，也存在明显改进诉求：45.83%的用户指出数据源覆盖有限（尤其是地方性、时效性强的政策与人才数据缺失）是

当前最大瓶颈，直接影响输出的完整性与深度；25%的用户认为结果“可应用性弱”，难以直接用于正式报告撰写，反映出模块在专业深度、情境适配及成果形态转化方面仍有不足。用户普遍期待未来能拓展数据生态、提升分析的洞察力（而非仅停留在描述层面），并使输出更贴近可直接使用的成果形式。

此外，用户对系统的可靠性与可控性提出更高要求：一方面希望模块能更清晰地呈现输出结果的置信度或不确定性；另一方面期待在复杂任务执行过程中拥有更精细的干预能力，例如中途调整参数、修正方向或注入个人偏好，以实现更紧密的人机协同。

（3）使用行为模式分析

基于问卷中关于使用频率和最常使用功能的交叉分析，可以推断用户的行为模式。高频用户最可能将模块应用于研究框架生成（66.67%）、研究报告撰写（62.5%）和数据搜集（45.83%）等环节，说明其在研究构思、初稿起草与资料准备阶段作用最为突出。问卷未直接统计模块内部各子功能调用频次，但从用户认定的最突出价值（自动整合数据、生成报告模板、降低整理成本）可以推断，文献与数据整合、报告辅助生成相关功能是当前使用最频繁、感知价值最高的核心应用点。综合来看，综上，科研助理模块已成功获得用户接纳，在提升研究效率、简化基础工作流程方面成效显著，其自然语言交互方式易于掌握，初步实现了从“工具”向“协作者”的角色过渡。

5.3.5 实际应用价值验证：典型案例分析

为评估科研助理模块在实际工作中的应用价值，我们选取了院内已完成的工作，引入智能体进行辅助，通过复现结果并与实际完成情况进行对比。

案例一：材料汇编

年度工作务虚会的材料汇编与分析是综合服务部一项较为常规但较为繁琐的任务，需要将与会者提交的书面发言材料与现场录音（无文字稿）进行整合，形成一份完整的会议发言汇编。在此基础上，整理出每位发言人所阐述的使命、愿景、价值观以及对下一年度工作的具体建议。另外，还需要对全部发言内容进行高频词统计并生成词云图，以直观呈现务虚会发言重点。传统上，完成这样一项涉及录音和文本处理、数据分析的综合任务，一名熟练工作人员需投入约5个工作日，整个工作流不仅时间长，而且产出质量与工作人员的经验 and 状态紧密

相关，且在准确性与一致性上面临挑战。

引入智能体后，只需上传文字和录音材料，智能体可自动完成语音转写、材料梳理，并基于语义理解精准提取各发言稿中的“使命、愿景、价值观”及下年度建议，同时生成高频词分析。最终，将所需全部成果整合为报告初稿。

核心处理时间从 5 天缩短至 5~8 小时。在准确性上，通过智能体实现了信息提取与统计分析，避免了人工疲劳带来的错误与主观偏差。后续人工对智能体输出报告进行复核，不仅极大提升了效率与质量，更促使工作人员角色从信息处理员转变为成果的审核员，专注于更具价值的洞察、提炼与决策支持工作。

案例二：问卷分析报告

发放问卷是研究报告进行调研分析的常用手段，需要对收集的问卷进行整合分析，以“丰台区推进教育科技人才一体发展研究”为例，建立科研项目，分类选择“调研报告”，根据需求选择“项目立意”，“参考模板”为此前已经成稿的其他调研问卷分析报告，“内容参考”为直接在问卷星下载的问卷填写结果，如图 5-1 所示。经 20 分钟处理，智能体共输出 8000 字左右的报告，分为 5 个章节：“一、问卷设计框架”“二、调研对象基本情况”“三、政策认知与发展现状评价”“四、分群体专项情况分析”“五、问题与建议总结”。

经核对，智能体生成报告内容未出现编造内容，问卷选项与对应数据的匹配准确率达 100%，并能基于预设模板自动生成表格（如调研对象年龄分布、政策满意度等），且样式规范、数据标注清晰。报告的章节逻辑与人工撰写的优质报告高度吻合，尤其在“分群体专项情况分析”部分，智能体能够根据不同学历、职业背景等维度对数据进行交叉分析，其分析深度基本达到中级研究人员的水平。传统模式下，完成同等规模和质量问卷分析报告，通常需要 1 名研究人员 2~3 个工作日，智能体将处理时间压缩至 2 小时（录入需求 10 分钟，生成报告 20 分钟，人工调整 90 分钟），效率提升近 20 倍，且有效避免了人工统计可能出现的计算错误和分类偏差。研究人员仅需对报告的措辞表达、部分建议的针对性进行微调，同时将一些表格内容更改为示意图，即可形成最终成果，显著减轻了基础性、重复性的分析工作负担。

返回

丰台教科人问卷分析

下载

项目规划

① 内容大纲

② 内容创作

项目类型

调研报告

计划规划

观点论述

新闻宣传

会议纪要

讲话发言

其他类型

项目立意

请尽可能详细描述项目的立意、创作要求，后续的AI将据此辅助生成。

我们设计了一份关于丰台区教育科技人才一体发展的调研问卷，本问卷面向高校、科研机构、企业等各类用人主体，科技领军人才、创新创业人才、青年人才等各类人才代表及有关政府部门工作人员，目前有347人填写。问卷中1-15，28题为所有人填写，16-18题学校（教育机构）人员填写，19、20题科研机构人员填写，21、22企业人员填写，23-26题为除政府部门外的其他人填写，27题为政府部门人员填写。请按照参考模板中的报告，帮我生成一个调研问卷的分析报告，画图或者表格展示数据。要求分析精准，不可编造

篇幅要求

自动

请输入您期望的全文大约字数

文章标题

丰台区教育科技人才一体发展调研问卷分析报告

参考模板

您可指定一份文档作为本项目最终生成的文稿在内容结构上的参考模板。

参考报告

内容参考

您可选择相关文档，作为本项目的生成文稿时在内容上的参考信息。

丰台区教育科技人才一体发展调研问卷—默认报告

从文件夹中选择

内容生成

允许在项目创作时联网检索

生成参考仅于上方指定内容

图 5-1 问卷分析报告智能体输入内容界面

5.3.6 结论与优化方向

综合多维度的实践检验结果，科研助理模块在人才研究领域的应用已展现出明确的价值与可行性。该模块功能体系较为完备，能够覆盖从信息采集、数据处理到分析建模及成果生成的全流程，有效支撑多种典型研究场景；在标准化、流程化任务中显著提升工作效率，人力投入减少，研究周期缩短，切实释放了研究人员的认知负荷与时间资源。同时，在严格的质量控制机制保障下，其输出内容在事实准确性、逻辑严谨性与格式规范性等方面均达到较高水准，具备良好的学术可信度。更重要的是，模块凭借良好的易用性已顺利融入院内研究人员的日常实践，初步形成“人主导、机执行”的新型人机协同研究模式，并在真实课题中验

证了其缩短研究周期、降低人力成本、拓展数据广度与思维深度的实际效能。

当前系统仍存在若干局限：在涉及深度领域洞察、跨学科类比或原始理论构思等复杂任务上表现不足；在多智能体协同执行长周期任务时，用户对内部进程的感知较弱，干预手段有限；对不同研究者个性化风格（如偏好定性/定量、关注宏观/微观）的适应性亦有待提升。未来工作将聚焦于三方面优化：一是探索如何将研究者的隐性知识与批判性思维有效注入智能体，实现“智能增强”而非仅“任务替代”；二是开发可视化任务看板与细粒度干预接口，增强人机协作的透明性与可控感；三是引入更精细的用户画像建模与持续学习机制，使智能体逐步演变为“懂我所思、知我所需”的专属科研伙伴。科研助理模块的实践验证印证了其作为人才研究领域专业化智能助手的核心定位——并非取代人类研究者，而是通过高效接管重复性、计算密集型劳动，将研究者聚焦于最具创造性的环节，共同推动人才研究迈向更高效、更深入、更智能的新阶段。

6 智能体辅助决策模块搭建

6.1 决策辅助模块的作用

决策辅助模块是智能体中价值密度最高的模块，它利用先进的数据分析和人工智能算法，为人才领域的各类决策提供数据驱动的建议。该模块集成多源数据融合与特征工程、动态优化与反馈学习三大技术体系，能够支持从微观层面的学者评价到宏观层面的人才政策制定等多种决策场景。

6.2 辅助决策模块的搭建方法

6.2.1 多源数据融合与特征工程技术

辅助决策模块面临的核心挑战在于人才数据的高度多源异构性。学者相关信息分散于论文数据库、科研项目库、专利系统等多个平台，其格式、粒度与质量存在显著差异。为解决语义不一致问题，本模块采用基于本体的数据融合方法：首先构建统一的人才领域本体，明确定义“学者”“机构”“成果”等核心概念及其相互关系；随后将各来源数据映射至该本体框架下，实现跨系统的语义对齐。该机制有效保障了数据的一致性与互操作性，为后续高质量决策分析奠定基础。

在特征工程方面，我们构建了多维度学者特征体系，包括生产力特征（发文量、项目数）、影响力特征（引用次数、h 指数）、合作特征（合作网络规模、中心性）和交叉性特征（学科交叉度、方法新颖性）等多个维度。针对人才研究的特点，特别设计了时序动态特征，捕捉学者学术表现的变化趋势，如影响力增长曲线、研究主题演化路径等，此类时序特征对预测学者未来发展潜力具有关键价值。

面对高维特征空间，本模块综合采用多种特征选择与降维策略：通过相关性过滤剔除冗余变量，利用模型嵌入重要性保留强预测性特征，并结合领域知识规则确保关键学术指标不被误删。对于高度相关的特征群（如多种引用指标），采用主成分分析等方法进行降维，在保留信息量的同时缓解多重共线性问题，显著提升后续建模的稳定性与计算效率。

6.2.2 动态优化与人机协同决策机制

为适应快速变化的科研环境与用户需求，本模块集成在线学习机制，支持模型参数随新数据和反馈持续更新。相较于传统批量训练模式，该机制能够实时融入最新学者动态与交互行为，确保决策建议始终反映当前学术生态，保持时效性与相关性。

人机协同决策是本模块的核心设计理念。我们不寻求完全替代人类决策者，而是通过决策支持界面，将模型的分析结果以直观、可解释的方式呈现给用户，辅助研究人员决策。我们设计了多种结果可视化方案，如平行坐标图展示多维度比较、决策树展示推理路径、热点图展示领域布局等，帮助研究人员快速把握复杂信息。同时，本模块提供“为什么是这个建议”的解释功能，透明展示决策逻辑，增强用户对智能体的信任。

为实现系统持续优化，构建了完整的反馈学习闭环。用户对决策建议的采纳、修改或拒绝行为都被记录并分析，用于评估和优化决策模型。我们特意设计了显式反馈和隐式反馈双通道：显式反馈通过用户直接评分或修正获取；隐式反馈通过用户最终决策行为与系统建议的差异推断。这两种反馈共同为模型优化提供丰富信号，驱动系统越来越符合用户的实际需求。

此外，本模块设有决策效果追踪模块，定期回溯历史建议与实际结果的匹配度。分析决策模型的准确性和盲点。例如，在人才选拔场景中，系统会追踪被推荐学者的后续发展，与未被推荐的对比群体比较，评估选拔标准的有效性。这种长期效果评估不仅优化了未来的决策质量，也为人才研究提供了宝贵的实证证据，形成了科研与应用的良性互动。

6.3 辅助决策模块的效果检验

辅助决策模块是智能体从信息处理迈向智慧支持的重要跃升。为验证该模块在实际应用中的有效性、可靠性和实用性，本研究设计了一套基于真实历史数据的场景化评估方案，聚焦于决策准确性、时效性、可解释性与采纳度四个维度，选取人才评价、政策匹配、风险预警等典型场景进行检验。通过将辅助决策模块生成的建议与已完结报告中的建议进行对比分析，系统评估了模块在提升决策质量、效率与科学性方面的实际价值。

6.3.1 检验设计与实施流程

检验数据来源于研究院已完结的内部研究报告中，每份案例均提取完整的背景信息、约束条件及最终决策结论，同时将对对应报告从知识库中临时移除，确保本模块无法直接检索到原始决策结果，从而模拟真实未知场景下的推理过程。具体检验流程如下：

（1）场景构建：基于历史案例，还原决策时的上下文环境，包括问题描述、可用数据与业务约束；（2）模块推理：将除最终结论外的所有输入信息提供给辅助决策模块，要求其生成结构化建议及支撑理由；（3）对比分析：将模块建议与已完结报告中的决策进行比对，重点分析两者在结论一致性、推理逻辑、信息利用方面的重合与差异。（4）评估维度：主要考察模块的决策准确性、效率提升度、逻辑可解释性，并明确记录其暴露出的问题。

6.3.2 历史场景还原测试

为了全面评估辅助决策模块的实际效能，我们选择了两个典型的测试场景：
测试场景一：区域人才政策痛点诊断与对策匹配

基于《京津冀人才一体化发展十周年》报告中 2014~2022 年的数据，模拟在 2020 年节点，诊断河北人才结构与产业升级错配问题。输入数据包括产业人才结构偏离度、技术合同输出流向等。此场景旨在检验模块在复杂区域人才政策分析中的能力。

表 6-1 测试场景一（区域人才政策痛点诊断与对策匹配）测试过程与结果

对比项	模块输出的决策建议	报告揭示的实际问题	一致性分析与模块表现评价
核心诊断	准确指出河北第二、三产业存在“人才供给滞后于产值增长”的结构性缺口，并预警若本地转化渠道不畅，创新成果可能外流。	报告明确指出河北存在“配套产业转型升级创新及技术技能人才短缺”“北京技术成果向津冀转化比例低”等问题。	问题识别精准。模块基于数据趋势成功揭示了关键矛盾。
对策建议	建议“在重点园区试点校企联合培养项目”和“设立技术转移激励基金”。	报告中提出的对策包括“优化产业结构布局”和“构建统一人才市场”，更为宏观和系统。	对策有效但深度不足。模块给出的建议是直接、可操作的战术级方案，但与报告中更具战略高度的顶层设计相比，在系统性和创新性上有差距。

测试场景二：科研管理模式对比与优化建议

历史案例还原：基于《如何让微软中国的创新“密码”真正留下来？》中的对比分析，模拟为优化本市科研管理提出建议。输入数据为报告中描述的“微软自由探索模式”与“当时我市目标导向为主模式”的客观特征对比。

表 6-2 测试场景二（科研管理模式对比与优化建议）测试过程与结果

对比项	模块输出的决策建议	报告中的核心建议	一致性分析与模块表现评价
分析结论	明确总结出两种模式在风险容忍度、研发周期、成果产出逻辑上的根本差异。	报告深刻揭示了两种文化的差异及其影响。	归纳总结能力强。模块能快速抽取出复杂文字描述中的关键对比维度。
优化建议	提出“在现有基金中设立自由探索子项目”和“在项目考核中引入‘里程碑式’弹性评估”。	建议“加大基础研究投入，鼓励自由探索”并“构建多维度评价体系”。	建议具有较高参考价值。模块将宏观原则转化为了具体、可落地的初步措施，具备良好的政策翻译能力。

6.3.3 综合评估

通过上述历史场景测试，本辅助决策模块展现出一个清晰的能力轮廓。主要优点为：（1）卓越的效率：在信息整合、数据分析、给出决策等环节，能将人力从大量重复性、结构化的工作中解放出来。（2）稳定的逻辑推理：其决策过程高度理性、一致且可追溯，能避免人类因疲劳、情绪带来的波动，确保分析基础的客观性。（3）可靠的知识助理：能够快速学习和应用既定规则、政策框架，在已知模式下的方案推荐准确率高，可作为政策研究员可靠的“副手”。

存在问题为：（1）依赖数据质量：其表现严格遵循“垃圾进，垃圾出”原则。当历史数据本身存在模糊或缺失时，模块的分析深度会显著下降。（2）缺乏隐性知识：无法处理评审中关于学术声望、领导潜力的微妙共识，也无法把握政策推动中复杂的人情世故与协调艺术，这是其与人类专家最本质的差距。（3）创新不足：模块擅长在既有框架内进行优化、组合与推荐，但难以提出完全颠覆性的原创方案。它的智慧更多是组合式创新，而非颠覆式创造。

辅助决策模块是一个实用价值突出的工具。它并非能替代人类决策的万能黑科技，而更像一个不知疲倦、逻辑严谨、知识渊博的数字分析员。在实际工作中，它能够出色地完成初步分析、方案起草、风险提示等前期工作，为决策者节省大

量时间，并提供坚实的数据与逻辑支撑。然而，所有由其产生的建议，最终必须交由人类专家，结合那些无法被量化的隐性知识、战略直觉和价值判断，进行最终的审核和落地。这种人机协同、优势互补的模式，正是其最大价值所在，也为未来人才决策模块的智能化演进提供了可行路径。

7 人才研究智能体与通用智能体对比分析

本章聚焦人才研究智能体（Talegent）的核心竞争力构建，从人才研究的实际科研场景需求出发，选取交互感知、文本处理、逻辑推理三大核心维度，与国内外通用大模型（如 GPT-4、Claude 3、文心一言、讯飞星火等）及科研智能体（如 Elicit、ResearchRabbit、IBM Watson、Scite AI 等）展开系统性、多层面的对比分析。通过明确各类智能体的技术特性、场景适配逻辑、功能边界及性能短板，深度凸显 Talegent 在人才研究领域的定制化优势，为科研人员、政策制定者选择适配工具提供清晰参考，同时为智能体在专业领域迭代优化提供实证依据。

7.1 交互感知能力对比

交互感知能力是衡量智能体能否真正融入科研工作流的关键指标，其核心价值不仅在于技术性能，更体现在对科研信息载体多样性、长周期协作连续性、用户身份差异性以及动态反馈响应性的综合适配能力上。Talegent 基于 OukeAI 融合大模型架构，充分融合了最新的人工智能技术和科研需求，形成了其在交互感知领域的四大核心能力：多模态交互支持、10 万+token 长期记忆、科研身份自适应体系以及实时反馈优化机制。支持在多种科研场景中的应用。本节从这四大维度出发，结合人才研究中的典型应用场景（如政策文本分析、学者数据整合、跨学科项目协作等），系统对比主流通用大模型、科研智能体与本研究开发的 Talegent 智能体在实际使用中的表现差异，揭示其在科研场景中实用性与用户体验的根本分别。

7.1.1 核心对比维度与指标详解

为确保评估的系统性与可操作性，本研究构建了四维评价框架（详见表 7-1），多模态支持考察系统对文本、表格、图表、公式等科研常见信息载体的原生解析能力、跨模态语义融合精度及其在人才研究任务中的适配程度；上下文记忆衡量系统在长周期项目中对研究框架、历史数据、讨论结论等关键信息的留存能力与逻辑连贯性；身份自适应评估系统能否根据用户角色（如决策者、研究员、执行者）动态调整输出内容的深度、术语密度与功能侧重；实时反馈检验系统对用户纠错、偏好调整等交互信号的响应速度、修正精准度及闭环学习能力。

表 7-1 人才研究智能体与通用智能体交互感知能力核心对比维度分析

对比维度	核心评价指标	Talegent	通用大模型（GPT-4/Claude3/文心一言/讯飞星火）	科研智能体（Elicit/IBM Watson/ResearchRabbit）
多模态支持	覆盖类型、原生融合度、科研场景适配精度、跨模态协同能力	1. 覆盖文本、图像、表格、代码、数学公式全类型原生支持； 2. 多模态深度融合，可实现“图像数据提取-文本逻辑分析-代码计算验证-公式推导呈现”的全流程协同； 3. 针对人才研究场景优化，可精准解析政策表格、学者成果图谱、科研数据图表等专业载体，适配人才政策分析、学者评价等核心场景	1. GPT-4/Claude3：支持文本+图像，需依赖外部插件，跨模态协同性弱，对科研表格、公式的原生解析能力不足； 2. 文心一言：支持文本+图像+语音，多模态融合深度不足，对科研表格的结构化解析精度低，难以适配复杂人才数据整合场景； 3. 讯飞星火：支持文本+语音+简单图像，不支持代码片段解析和复杂公式识别，无法适配人才研究中量化分析场景	1. Elicit：仅支持纯文本输入输出，无法处理政策图表、学者合作网络图谱等非文本信息，局限于文献检索场景； 2. IBM Watson：侧重医疗影像等特定领域模态处理，缺乏对科研表格、数学公式的原生解析能力，适配人才研究场景有限； 3. ResearchRabbit：仅支持文本+文献图谱可视化，无法解析表格数据和代码逻辑，难以支撑多源人才数据协同分析
上下文记忆	有效 token 长度、多模态记忆覆盖、长周期项目连贯性、关键信息留存率	1. 支持 10 万+token 长期记忆，覆盖文本、图像、代码、表格等多模态信息； 2. 具备项目级记忆管理，可持续追踪跨月/跨季度科研项目的背景信息、讨论结论、数据参数，适配人才流动趋势研究等长周期项目； 3. 关键信息留存率超 95%，可精准调用历史交互中的学者数据、政策条款、分析逻辑	1. GPT-4 （ 128Ktoken ） /Claude3 （200Ktoken）：纯文本场景下记忆长度占优，但多模态交互时有效记忆长度缩短 50%以上，易遗忘关键数据细节，如跨项目协作中的历史分析标准； 2. 文心一言（32Ktoken）/讯飞星火（64Ktoken）：长文本处理时存在“记忆断层”，超过 5000 字的政策文本分析中，关键条款留存率不足 70%，影响长周期人才政策追踪研究	1. Elicit：无项目级记忆功能，每次交互均为独立会话，无法关联历史检索结果和分析逻辑，重复检索效率低下； 2. IBM Watson：记忆局限于单领域数据，跨场景（如从学者成果分析切换至政策影响评估）时连贯性差，难以支撑多维度人才研究； 3. ResearchRabbit：仅记忆用户的文献检索偏好，无法留存项目协作中的讨论细节和数据参数，跨阶段研究衔接不畅

对比维度	核心评价指标	Talegent	通用大模型（GPT-4/Claude3/文心一言/讯飞星火）	科研智能体（Elicit/IBM Watson/ResearchRabbit）
身份自适应	角色细分度、回答深度适配、专业术语调整灵活性、需求优先级匹配	- 覆盖领导、政策研究员、一线科研人员、学生、行业专家等细分角色； - 动态调整回答的专业深度：对领导侧重宏观趋势和核心结论，对研究员提供技术细节和数据支撑，对学生补充基础概念和研究方法； - 可根据用户研究领域（如人才政策、学者评价、人才流动）优化术语体系，适配不同岗位的核心需求	- GPT-4/Claude3: 仅支持基础身份分类（普通用户/专业人士），无科研角色定制化，对人才研究中的专业术语（如 h 指数、学科交叉度、政策适配性）缺乏深度理解，回答通用性强； - 文心一言/讯飞星火: 身份适配停留在“学生/教师”粗分类，无法根据人才研究的具体岗位需求调整回答逻辑，如难以同时满足政策制定者的宏观需求和科研人员的微观分析需求	- 所有工具均无身份自适应功能，回答风格和深度固定统一； - 对政策制定者和一线研究员输出相同维度的信息，无法满足不同角色的核心需求（如政策制定者需宏观数据支撑决策，研究员需微观数据支撑分析），适配性极差
实时反馈	反馈响应速度、优化针对性、闭环机制完整性、用户行为适配度	- 实时捕捉用户交互数据（如标记错误信息、补充缺失需求、调整回答偏好），10 秒内启动响应优化； - 形成“用户反馈-模型调整-效果验证-持续迭代”的完整闭环，适配人才研究中动态变化的分析需求； - 结合用户身份和项目场景优化反馈策略，如对研究员的技术质疑优先优化数据精度，对政策制定者的需求补充优先扩展信息维度	- 均依赖离线更新或周期性微调（更新周期通常为 1~3 个月），无实时交互优化机制； - 无法根据用户在单一项目中的动态需求调整响应逻辑，如用户在政策分析中临时要求补充区域对比数据，需重新发起查询，影响研究连贯性	- 仅支持有限功能反馈（如 Elicit 的文献检索结果评分、ResearchRabbit 的图谱展示反馈）； - 无系统性优化闭环，用户反馈无法直接驱动工具功能调整，适配人才研究场景的灵活性极差，难以应对动态变化的研究需求

7.1.2 Talegent 的场景化交互能力与实证表现

在人才研究的复杂工作场景中，Talegent 通过多模态理解、长周期记忆、角色自适应与动态反馈四项核心机制，显著提升了交互效能。以下从能力维度与典型项目两个层面进行验证。

（1）多模态政策理解能力

在人才政策效果评估场景中，Talegent 展现出显著优势：用户可直接上传包含条款文本、配套数据表格与实施效果图表的地方人才政策文件，系统自动提取表格中的量化指标（如引进人数、补贴标准），结合文本语义解析政策导向，并调用代码验证数据合理性，最终以数学公式形式呈现投入产出比，实现“多模态输入—一体化分析—结构化输出”的闭环，高度契合政策评估的量化与逻辑严谨性需求。

相较之下，通用大模型普遍存在模态割裂问题：GPT-4 要求用户手动将表格转为文本后方可分析，且无法联动计算验证；文心一言在解析复杂表格时易丢失行列结构，导致数据失真；讯飞星火则完全无法识别表格与公式，仅能生成文本摘要，难以支撑深度分析。科研智能体亦存在局限：Elicit 仅处理文本关键词，忽略表格与图表；IBM Watson 虽可解析部分图表，但无法将其与政策文本逻辑关联，难以形成完整评估结论。

（2）长周期上下文记忆能力

在“跨年度青年人才流动趋势研究”项目中，Talegent 能有效留存项目初期的研究框架、中期采集的学者数据集及多次讨论确定的分析维度（如地域、学科、单位类型），后续用户仅需补充新年度数据，系统即可自动衔接历史逻辑进行趋势对比，无需重复交代背景。

而 GPT-4 尽管支持 128Ktoken 上下文，但在跨越数月的协作中仍会遗忘早期设定的筛选标准，需用户反复提醒；文心一言受限于 32Ktoken 容量，无法容纳完整项目资料，被迫分段上传，破坏分析连贯性。专业工具如 Elicit 与 ResearchRabbit 均采用无状态会话模式，前者无法关联历史检索结果，后者仅记忆文献偏好，均难以支撑需要持续演进的人才趋势研判任务。

（3）角色自适应输出机制

在“区域人才政策优化”跨角色项目中，Talegent 能根据用户身份动态调整

输出：面向政策制定者，提供宏观趋势总结与优化建议摘要；面向科研人员，输出政策条款拆解、数据来源说明、分析模型及可复用代码；面向基层执行者，则以通俗语言解读条款并附操作指引。这种细粒度适配完美契合人才工作“决策—研究—执行”三位一体的实际需求。

反观通用模型，GPT-4 对所有用户输出统一深度内容，导致决策者需从技术细节中提炼要点，基层人员则难以理解专业表述；文心一言仅粗略区分“学生/教师”，无法识别政策场景中的角色差异。专业工具则普遍缺乏身份感知能力，Elicit 固定输出文献摘要，ResearchRabbit 仅展示文献图谱，均无法融入多角色协同的工作流程。

（4）动态优化实时反馈机制

在“学者学术影响力评估”场景中，若用户指出 Talegent 的 h 指数计算未涵盖某类核心期刊，系统可立即捕获该反馈，自动校验数据源（如调整 Web of Science 与 Google Scholar 的优先级），补充缺失期刊数据，重新计算并生成修正结果，同时记录该用户的个性化口径，用于后续查询。

而 GPT-4 需用户重新发起完整查询并明确错误点，无法实现增量修正；文心一言的反馈优化依赖月度模型更新，响应滞后；讯飞星火缺乏显式反馈通道。专业工具中，Elicit 仅支持相关性评分，无法修正计算逻辑；IBM Watson 的反馈需通过专业接口提交，且无实时修正能力，难以满足人才研究中对动态调整的高时效要求。

（5）“某省顶尖人才引进政策效果评估”项目实证

上述能力在“某省顶尖人才引进政策效果评估”项目中得到集成验证：Talegent 同步处理政策文本、数据表格与效果折线图（多模态支持），无缝衔接三个月前的项目讨论记录（上下文记忆），为不同角色用户提供差异化输出（身份自适应），并可根据用户临时提出的“增加区域对比”需求即时调整分析维度（实时反馈）。整个过程无需重复输入，高效支撑项目全周期、多角色的协同需求。相比之下，GPT-4 需用户分别上传文本、表格与图像，且无法关联历史逻辑；文心一言因表格解析错误导致评估缺乏量化基础；Elicit 与 IBM Watson 则因模态覆盖不全或缺乏角色适配，输出内容片面或脱离实际工作场景，难以真正嵌入科研决策流程。

综上，Talegent 在交互感知能力上的系统性设计，使其在人才研究这一高度专业化、多模态、长周期、多角色的复杂场景中，展现出远超通用模型与现有科研工具的实用性与用户体验深度，为智能体在专业领域的落地提供了可复制的技术路径与应用范式。

7.2 文本处理能力对比

文本处理是人才研究的核心工作环节（如文献综述撰写、政策文本解析、研究报告生成、跨语言学术交流等），其核心差异体现在学术规范性、检索精准度、多语言适配性及伦理合规性上。本节从格式规范、检索与摘要、多语言处理、抗偏见控制四大核心维度展开对比，结合人才研究中的具体文本处理场景（如中英文文献整合、政策文本语义分析、研究报告撰写等），深度解析不同智能体的性能差异。

7.2.1 核心对比维度与指标详解

从四个核心维度对不同智能体的文本处理能力进行系统评估。在格式规范方面，主要考察其对各类学术格式（如 APA、国标等）的覆盖范围、参考文献引用的准确性、专业术语使用的标准化程度，以及生成文本整体结构的逻辑性。在检索与摘要维度，重点评估其对接学术数据库的能力（如 CNKI、Web of Science 等）、检索结果的精准度、所生成摘要的结构化质量，以及整合多篇文献信息形成连贯综述的能力。在多语言处理方面，关注其支持的语种范围、对人才研究领域专业术语的跨语言翻译准确性、生成文本在多语言切换下的逻辑连贯性，以及对中英等多语言混合写作场景的支持能力。在抗偏见控制维度，则着重检验其在生成内容中性别、地域或文化偏见的发生率，整体学术立场的中立性，以及对科研伦理规范（如数据脱敏、隐私保护）的遵循程度。

表 7-2 人才研究智能体与通用智能体文本处理能力核心对比维度分析

对比维度	核心评价指标	Talegent	通用大模型（GPT-4/Claude3/文心一言/讯飞星火/通义千问）	科研智能体（Elicit/SciteAI/ResearchRabbit）
格式规范	学术格式覆盖范围、引用准确性、术语标准化程度、文本结构逻辑性	1. 支持 IEEE、APA、国标 GB/T7714-2015 等多类学术格式，可自动适配人才研究报告、学术论文等不同产出场景； 2. 引用准确率超 95%，可自动关联 CNKI、Web of Science 等数据库的文献源，避免引用虚构问题，符合学术研究规范； 3. 内置人才研究专业术语库，实现术语统一（如“跨学科人才”“青年领军学者”等定义），适配人才研究专业表达需求； 4. 生成文本结构符合科研逻辑，如文献综述自动包含“研究背景-核心观点-争议焦点-未来方向”框架，适配人才领域学术写作场景	1. GPT-4：APA 格式错误率 15%，中文文献引用易出现来源虚构，术语使用缺乏行业规范，难以满足人才研究的学术写作要求； 2. 文心一言：仅支持国标格式，缺乏国际学术格式适配，生成文本偏向模板化，结构灵活性不足，难以适配国际人才研究合作的写作需求； 3. 讯飞星火：无固定学术格式支持，引用标注混乱，术语标准化程度低，需大量手动校对，影响研究效率； 4. 通义千问：侧重商业文本格式，缺乏科研文本结构优化，无法满足人才研究的学术写作需求	1. Elicit：仅能生成简单文献摘要，无学术格式支持，引用信息不完整，难以支撑人才研究文献综述的规范撰写； 2. SciteAI：支持引用可靠性标注，但格式单一，无法适配不同期刊的格式要求，适配性有限； 3. ResearchRabbit：无文本生成功能，仅提供文献关联图谱，无法辅助撰写结构化文本，无法直接支撑人才研究的文本产出

对比维度	核心评价指标	Talegent	通用大模型（GPT-4/Claude3/文心一言/讯飞星火/通义千问）	科研智能体（Elicit/SciteAI/ResearchRabbit）
检索与摘要	数据库覆盖范围、检索精准度、结构化摘要质量、文献整合能力	- 集成 CNKI、PubMed、arXiv、Web of Science、万方等多数据库，支持实时检索，覆盖人才研究相关的政策文献、学术成果、学者数据等核心资源； - 针对人才研究优化检索算法，可精准定位“特定领域学者成果”“区域人才政策”“跨学科人才流动研究”等细分主题，检索精准度高； - 生成结构化摘要，包含“研究目的-方法-核心数据-结论-与本研究关联性”五大模块，适配快速文献筛选需求； - 支持多文献自动整合，可快速生成人才研究相关的专题综述，大幅提升文献处理效率	- GPT-4/Claude3：无法直接访问 CNKI 等中文核心数据库，需用户手动上传文献，检索精准度受限于输入信息完整性，难以快速获取人才研究相关中文文献； - 文心一言：集成 CNKI、万方数据库，检索精准度较高，但结构化摘要质量差，缺乏与研究主题的关联性分析，无法高效支撑人才研究文献整合； - 讯飞星火/通义千问：无原生数据库集成，需用户手动输入文献信息，检索效率极低，难以适配大规模文献检索场景	- Elicit：支持文献摘要检索，但数据库覆盖范围有限，且仅能输出简单文本摘要，无结构化整合，无法支撑人才研究专题综述撰写； - SciteAI：侧重引用分析，可标注文献引用语境，但检索精准度低，无法聚焦人才研究细分主题，检索效率不足； - ResearchRabbit：擅长文献关联图谱生成，但检索功能薄弱，无法直接获取文献核心内容，难以满足人才研究文献快速筛选需求
多语言处理	语种覆盖范围、专业术语翻译准确性、跨语言文本连贯性、多语言混合写作支持	- 支持中英日韩四语无缝切换，优化中文科研文本处理（包括古文献术语、新兴交叉学科术语），适配人才研究中的多语言文献整合需求； - 专业术语翻译错误率低于 3%，可精准转换“人才评价指标”“政策工具类型”“学科分类标准”等专业术语，保障跨语言交流准确性；	- GPT-4：中英混合写作能力强，但中文长文本（超过 5000 字）处理易出现逻辑断层，专业术语翻译准确性不足，影响多语言人才研究报告质量； - 文心一言：中文文本处理优势明显，但英文专业文献翻译错误率高达 8%，无法精准转换人才研究专业术语，难以适配国际合作场景；	- 所有工具均以英文处理为主，中文文本支持不足，专业术语翻译错误率高，无法适配中文人才研究文献的处理需求； - 无多语言混合写作支持，无法适配人才研究中的国际合作场景，难以满足多语言成果输出需求

对比维度	核心评价指标	Talegent	通用大模型（GPT-4/Claude3/文心一言/讯飞星火/通义千问）	科研智能体（Elicit/SciteAI/ResearchRabbit）
		3. 支持多语言混合写作，如中英双语政策对比分析报告，保持文本逻辑连贯，适配国际人才研究合作场景； 4. 针对人才研究中的跨语言场景优化，如英文政策文本的中文语义解析、中文学者成果的英文摘要生成，满足多语言成果输出需求	3. 讯飞星火：中文+方言支持较好，但外文处理能力极弱，无法满足国际人才研究合作需求； 4. 通义千问：中英文基础翻译支持，缺乏专业术语优化，跨语言文本连贯性差，难以支撑多语言人才研究成果输出	
抗偏见控制	性别/地域/文化偏见发生率、学术中立性、伦理合规性	1. 通过严格的数据清洗和规则约束，性别、地域偏见率降低 40%以上，对不同区域、不同群体的人才政策效果分析保持中立； 2. 生成文本保持学术中立，不带有主观倾向，适配人才政策公平性评估、跨区域人才流动分析等场景； 3. 内置伦理合规检查，对涉及未公开的青年人才隐私数据自动脱敏，符合科研伦理规范，保障人才数据安全	1. GPT-4/Claude3：存在明显的西方文化倾向，对中国区域人才政策分析易带有主观偏见，影响评估结果客观性； 2. 文心一言/讯飞星火：地域偏见相对较低，但性别偏见仍较明显（如描述男性学者侧重“学术成就”，女性学者侧重“协作能力”），不符合人才评价的公平性要求； 3. 通义千问：商业场景优化导致其文本生成带有一定的功利性倾向，学术中立性不足，难以适配人才研究的客观分析需求	1. 无专门的抗偏见设计，文本生成受训练数据影响，存在隐性偏见，可能导致人才分析结果失真； 2. 缺乏伦理合规检查功能，无法处理涉及隐私的人才数据，存在数据安全风险，难以适配人才隐私保护相关研究场景

7.2.2 Talegent 的文字处理能力与实证验证

在人才研究的文本处理任务中，Talegent 通过格式规范、高效检索、多语言支持与抗偏见控制四项机制，显著提升学术产出质量与效率。以下从能力维度与典型项目两个层面进行验证。

（1）学术格式规范保障能力

在“人才流动趋势研究”论文撰写场景中，Talegent 可自动生成符合 APA 格式的参考文献列表，精准标注中英文文献的来源、卷期、页码等信息，且术语使用统一（如“人才流动”“人才集聚”等核心术语保持一致）。生成的论文结构符合学术规范，摘要包含“研究目的一方法一数据一结论”的完整逻辑，讨论部分围绕研究问题展开，无冗余内容，完全适配人才领域学术论文的撰写需求。

通用大模型中，GPT-4 生成的 APA 格式参考文献错误率高达 15%，部分中文文献引用存在虚构来源问题，影响论文可信度；文心一言仅能生成国标格式参考文献，无法满足国际期刊投稿需求，且生成文本模板化严重，讨论部分论点重复率达 22%，缺乏学术创新性；讯飞星火无固定学术格式支持，引用标注混乱，需用户手动校对修改，耗时耗力，影响研究效率。科研智能体中，Elicit 无法生成结构化论文内容，仅提供文献摘要片段，难以支撑完整论文撰写；SciteAI 虽能标注引用可靠性，但无法辅助完成符合格式规范的论文撰写，无法直接服务于人才研究的学术产出。

（2）文献检索与结构化摘要能力

在“人工智能领域人才评价研究”文献综述场景中，Talegent 可实时检索 CNKI、Web of Science 等数据库的相关文献，自动生成结构化摘要，快速整合不同文献的核心观点、研究方法和数据结果，形成完整的专题综述，且可标注每部分观点的文献来源，方便用户追溯验证，大幅提升文献处理效率，适配人才研究文献综述的快速撰写需求。

通用大模型中，GPT-4 需用户手动上传文献文本，无法直接检索数据库，且生成的摘要缺乏结构化逻辑，核心观点提炼不精准，难以支撑高质量综述撰写；文心一言虽能检索中文数据库，但生成的综述缺乏文献间的关联性分析，无法形成完整的研究脉络，难以展现人才评价研究的发展趋势；讯飞星火无数据库集成功能，需用户手动输入文献信息，综述撰写效率极低，无法适配大规模文献整合

场景。科研智能体中，Elicit 仅能检索文献摘要，无法整合多文献形成综述，需手动整理；ResearchRabbit 擅长展示文献关联图谱，但无法提取核心观点，无法辅助综述撰写，难以满足人才研究文献快速整合的需求。

（3）多语言处理科研协作支持能力

在“中外顶尖人才引进政策对比”研究场景中，Talegent 可无缝切换中英文处理，精准翻译中外政策文本中的专业术语（如“绿卡制度”“居留许可”“人才补贴”等），生成中英双语对比报告，保持文本逻辑连贯和术语统一。同时，可将中文学者的研究成果自动生成符合国际规范的英文摘要，方便国际交流合作，完美适配人才研究的国际合作场景。

通用大模型中，GPT-4 中英文混合写作能力较强，但中文政策文本中的专业术语翻译准确性不足，易出现歧义，影响政策对比的准确性；文心一言英文政策文本翻译错误率高达 8%，无法精准传达政策核心条款，难以支撑深度政策对比；讯飞星火外文处理能力极弱，无法完成英文政策文本的深度解析，无法满足中外政策对比需求。科研智能体均以英文处理为主，中文政策文本解析能力不足，无法精准转换人才政策专业术语，难以适配中外人才政策对比的研究场景，无法支撑国际人才研究合作。

（4）抗偏见与伦理合规保障能力

在“不同区域人才政策效果评估”场景中，Talegent 生成的评估报告保持中立客观，对东部、中部、西部区域的政策效果分析仅基于数据事实，无主观褒贬倾向，且避免使用带有地域偏见的表述。同时，对涉及未公开的青年人才个人数据自动脱敏，仅呈现群体统计信息，符合科研伦理规范，保障人才数据安全，适配人才政策公平性评估的核心需求。

通用大模型中，GPT-4 对中国区域人才政策分析易带有西方文化偏见，过度强调市场导向政策的优势，影响评估结果的客观性；文心一言存在轻微地域偏见，对经济发达地区的政策评价偏向正面，对欠发达地区的政策则侧重不足，不符合公平性评估要求；讯飞星火性别偏见较明显，在描述不同性别学者的学术成就时，表述存在隐性差异，影响人才评价的公正性。科研智能体无专门的抗偏见设计，评估结果受训练数据影响，存在隐性的地域或群体偏见，可能导致人才政策评估结果失真，无法满足公平性评估的需求；且缺乏伦理合规检查，处理未公开人才

数据时存在隐私泄露风险，不符合科研伦理规范。

（5）综合项目实证：“跨国人才流动趋势研究”文献综述

以“跨国人才流动趋势研究”中的文献综述撰写为例，Talegent 可实时检索中、英、日、韩四语相关文献，自动整合不同语种文献的核心观点，生成结构化综述，其中英文献引用符合 APA 格式，中文文献符合国标格式，专业术语翻译准确，且无地域或性别偏见。整个过程无需用户手动检索和格式校对，大幅提升综述撰写效率，完全适配跨国人才流动研究的多语言文献整合需求。

而 GPT-4 需用户手动上传多语种文献，且生成的综述格式混乱，专业术语翻译错误率高，影响综述质量；文心一言无法检索英文核心文献，综述的国际视野不足，且格式单一，难以满足跨国研究的需求；Elicit 仅能检索英文文献摘要，无法整合多语种文献，综述覆盖面有限，无法支撑全面的跨国人才流动研究；SciteAI 侧重引用分析，无法辅助完成结构化综述撰写，效率极低，难以适配大规模多语言文献整合的需求。

7.3 逻辑推理能力对比

逻辑推理是人才研究的核心价值所在，涉及复杂问题拆解、数据验证、因果分析、假设检验等关键环节，直接决定研究结论的科学性和可靠性。本节从多步推理、数学计算、假设检验、不确定性管理四大核心维度展开对比，结合人才研究中的具体推理场景（如政策效果归因分析、学者学术影响力预测、人才流动趋势研判等），深度解析不同智能体的推理能力差异。

7.3.1 核心对比维度与指标详解

针对逻辑推理能力评估，聚焦于四个关键维度。多步推理能力通过考察智能体可执行的推理步数、各步骤间的逻辑连贯性、对复杂问题的拆解精度，及在跨学科背景下进行综合推理的能力来衡量。数学计算维度评估其对常用统计检验方法的覆盖度、数学公式推导的准确性、生成可执行代码的质量及对人才研究中量化分析任务的适配性。在假设检验方面，关注其自动生成研究假设的数量与创新性、所设计实验方案的科学性，及检验过程的统计严谨性。在不确定性管理维度，考察其对预测或推断结果置信度的评估精度、对潜在风险因素的提示明确性、针对局限性所提建议可行性，及对不确定性本身的量化表达能力（如置信区间）。

表 7-3 人才研究智能体与通用智能体逻辑推理能力核心对比维度分析

对比维度	核心评价指标	Talegent	通用大模型（GPT-4/Claude3/文心一言/讯飞星火）	科研智能体（Elicit/IBM Watson/SciteAI）
多步推理	推理步数、逻辑连贯性、问题拆解精度、跨学科推理能力	- 支持 10 步以上复杂问题分层拆解，如“人才政策效果评估”可拆解为“数据收集-变量控制-因果识别-结果验证-结论推导”全流程，逻辑严密； - 逻辑连贯性强，中间推理步骤透明可追溯，无逻辑断层，适配人才研究复杂问题分析需求； - 针对人才研究优化问题拆解逻辑，可精准识别核心变量（如政策工具类型、人才特征、区域环境），拆解贴合研究实际； - 具备跨学科推理能力，可融合经济学、社会学、统计学等多学科方法分析人才问题，适配跨学科人才研究场景	- GPT-4/Claude3：支持 5~8 步推理，易出现逻辑断层和自我矛盾，复杂问题拆解精度低，如混淆“人才政策相关性”与“因果性”，难以适配人才政策效果的深度归因分析； - 文心一言：仅支持 3~5 步简单推理，跨学科推理能力弱，无法融合多学科方法分析人才流动问题，难以支撑跨学科人才研究； - 讯飞星火：推理步数局限于 2~3 步，逻辑连贯性差，对复杂问题的拆解缺乏系统性，无法适配人才研究复杂问题的深度分析	- Elicit：无多步推理能力，仅能基于单文献提取观点，无法进行复杂问题分析，难以支撑人才研究复杂问题拆解； - IBM Watson：侧重单领域推理（如医疗领域），无法适配人才研究的跨学科推理需求，难以融合多学科方法分析人才问题； - SciteAI：仅能进行简单的引用逻辑分析，无复杂问题拆解和推理能力，无法支撑人才研究复杂问题的深度分析
数学计算	统计检验覆盖度、公式推导准确性、代码生成质量、量化分析适配性	- 支持 t 检验、ANOVA、回归分析、贝叶斯网络等高级统计检验，适配人才研究量化分析需求； - 原生支持 LaTeX 公式推导，准确性超 95%，可验证人才评价指标的计算逻辑，保障量化分析准确性； - 生成 Python、R 等科研代码，支持沙盒测试，确保代码可执行性，适配人才研	- GPT-4/Claude3：依赖外部计算器插件，复杂统计分析错误率超 25%，公式推导易出现符号错误，难以保障人才研究量化分析的准确性； - 文心一言：仅支持基础统计计算，无法处理贝叶斯网络等高级算法，公式推导准确性低，无法支撑复杂人才量化分析； - 讯飞星火：不支持 LaTeX 公式生成与解析，	- 所有工具均缺乏高级数学计算能力，仅能进行简单数据统计，无法支撑人才研究复杂量化分析； - 无法生成科研代码和推导复杂公式，无法支撑人才研究中的量化分析场景，如人才评价指标计算、政策效果量化建模

对比维度	核心评价指标	Talegent	通用大模型（GPT-4/Claude3/文心一言/讯飞星火）	科研智能体（Elicit/IBM Watson/SciteAI）
		究数据分析自动化需求； 4. 针对人才研究优化量化分析功能，可精准计算 h 指数、学科交叉度、政策实施效率等专业指标，贴合人才研究量化需求	数学计算仅局限于简单加减乘除，无法满足人才研究的量化分析需求，如学者学术影响力指标计算、政策效果量化评估	
假设检验	假设生成数量、创新性、实验设计科学性、统计严谨性	1. 基于大规模文献语料和知识图谱，自动生成 3~5 个跨学科创新假设，如“数字经济背景下，跨学科人才流动与区域创新效率的非线性关系”，贴合人才研究前沿方向； 2. 可自动设计科学的实验方案，包括对照组设置、样本分布验证、数据独立性检验，适配人才研究假设验证需求； 3. 严格遵循统计推断框架，明确零假设（ H_0 ）与备选假设（ H_1 ）、显著性水平（ α ），精准计算检验统计量和 p 值，保障检验严谨性； 4. 针对人才研究场景优化假设检验逻辑，可适配政策评估、学者预测等细分场景，贴合实际研究需求	1. GPT-4/Claude3: 假设生成缺乏创新性，多重复已有文献观点，且无法设计科学的实验方案，难以支撑人才研究创新假设验证； 2. 文心一言/讯飞星火：仅能生成 1~2 个简单假设，缺乏跨学科视角，假设检验的统计严谨性不足，易忽略显著性水平设置，无法保障人才研究假设检验的科学性	1. Elicit/SciteAI: 无假设生成和检验功能，仅能提取已有文献的假设，无法支撑创新型人才研究； 2. IBM Watson: 仅能在特定领域（如医疗）生成简单假设，无法适配人才研究场景，难以支撑人才研究假设生成与验证
不确定	置信度评估精度、风险提示明	1. 内置不确定性评估模型，可精准标注推理结果的置信度（如高/中/低），适配人才研究结果可靠性判断需求；	1. 所有通用大模型均无系统的不确定性管理机制，不标注结果置信度，易误导用户，难以保障人才研究结果的可靠性；	1. 无不确定性管理功能，仅呈现固定结果，不考虑数据完整性和模型适配性，无法支撑人才研究结果可靠性判断；

对比维度	核心评价指标	Talegent	通用大模型（GPT-4/Claude3/文心一言/讯飞星火）	科研智能体（Elicit/IBM Watson/SciteAI）
性管理	确性、补充建议可行性、不确定性量化能力	2. 对低置信度结果主动给出风险提示，说明不确定性来源（如数据不足、模型局限），避免误导用户； 3. 提供可行的补充建议，如“建议补充2024年区域人才流动数据以提升预测精度”，贴合人才研究优化需求； 4. 支持不确定性量化表达，如通过置信区间呈现预测结果的波动范围，保障结果呈现的科学性	2. 对低置信度结果缺乏风险提示，仅输出确定性结论，忽略数据或模型的局限性，可能导致人才研究结论失真	2. 无法提供补充建议，用户难以判断结果的可靠性和优化方向，难以适配人才研究的严谨性需求

7.3.2 关键差异的场景化解析

(1) 多步推理与系统性问题拆解

在“某省人才引进政策对区域创新效率的影响分析”场景中，Talegent 可将复杂问题拆解为 6 个逻辑严密的步骤：①政策工具分类（如补贴、住房支持、科研平台搭建）；②数据收集（政策实施前后的人才数量、创新成果、经济数据）；③变量控制（排除区域经济基础、产业结构等干扰因素）；④因果识别（采用双重差分模型分析政策净效应）；⑤结果验证（通过稳健性检验确认结论可靠性）；⑥结论推导（提出政策优化方向）。整个推理过程逻辑连贯，中间步骤透明可追溯，完全适配人才政策效果深度归因的复杂需求。

通用大模型中，GPT-4 虽能拆解 3~4 个步骤，但易混淆“变量控制”和“因果识别”的逻辑顺序，且缺乏稳健性检验环节，导致结论可靠性不足，难以支撑政策效果的严谨分析；文心一言仅能拆解 2~3 个简单步骤，无法识别干扰变量，直接将人才数量增长归因于政策效果，存在逻辑漏洞，无法保障分析的科学性；讯飞星火无法进行系统性拆解，仅能输出“政策促进创新效率提升”的单一结论，缺乏推理过程，无法支撑深度政策分析。科研智能体中，Elicit 仅能提取相关文献的观点，无法进行多步推理，难以支撑复杂政策效果分析；IBM Watson 无法适配人才政策分析场景，推理逻辑与研究需求脱节，无法提供有效分析支撑。

(2) 量化计算与模型实现能力

在“学者学术影响力评估”场景中，Talegent 可精准计算 h 指数、g 指数、引用半衰期等核心指标，通过 LaTeX 公式呈现计算逻辑，生成 Python 代码验证数据准确性，同时支持多维度指标的综合分析（如结合发文质量、合作网络中心性进行评估）。对于复杂的量化模型（如贝叶斯网络预测学者未来影响力），可完整推导模型公式并生成可执行代码，完全适配人才评价量化分析的核心需求。

通用大模型中，GPT-4 依赖外部计算器插件计算 h 指数，复杂指标（如引用半衰期）计算错误率超 30%，公式推导易出现符号错误，无法保障人才评价指标计算的准确性；文心一言仅能计算发文量、引用次数等基础指标，无法处理 h 指数等复杂指标，且无公式推导功能，难以支撑深度人才评价；讯飞星火不支持 LaTeX 公式生成，数学计算仅局限于简单统计，无法满足学术影响力评估的量化需求，如复杂评价指标计算、预测模型构建。科研智能体均无高级数学计算能力，

仅能呈现已有数据，无法进行指标计算和模型推导，无法支撑人才评价的量化分析，如学者学术影响力综合评估、人才发展潜力预测。

（3）创新假设生成与检验设计

在“数字经济背景下青年人才流动机制研究”场景中，Talegent 可自动生成 3 个创新假设：①数字技术发展通过降低信息成本，显著促进跨区域青年人才流动；②高技能青年人才对数字技术的敏感度高于普通技能人才，流动意愿更强；③区域数字基础设施水平调节数字技术与青年人才流动的关系。同时，设计科学的实验方案：以 2020~2024 年各省份青年人才流动数据为样本，以数字经济发展指数为核心解释变量，以数字基础设施水平为调节变量，控制区域经济发展水平、产业结构等干扰因素，采用调节效应模型进行检验，完美适配人才研究创新假设生成与验证的需求。

通用大模型中，GPT-4 生成的假设缺乏创新性，多重复“数字经济促进人才流动”的已有观点，且无法设计科学的实验方案，难以支撑人才研究创新；文心一言仅能生成 1 个简单假设，无调节变量和控制变量设计，假设检验的统计严谨性不足，无法保障检验结果的科学性；讯飞星火无法生成有效假设，仅能重复用户输入的研究方向，无假设检验能力，无法支撑人才研究的创新探索。科研智能体中，Elicit/Scite AI 仅能提取已有文献的假设，无法生成创新假设，难以支撑人才研究的前沿探索；IBM Watson 无法适配人才流动研究场景，假设设计与研究需求无关，无法提供有效检验方案。

（4）不确定性管理与可靠性保障

在“2025 年青年人才流向预测”场景中，Talegent 通过模型计算得出“一线城市青年人才流入率将下降 5%”的预测结果，同时标注置信度为“中等”，明确提示不确定性来源（如未考虑突发政策调整、经济波动等因素），并提出补充建议（如补充各城市最新落户政策、产业招商数据以提升预测精度），通过 95% 置信区间呈现预测结果的波动范围（3%~7%），保障预测结果的科学性和可靠性，适配人才流动趋势研判的严谨性需求。

通用大模型中，GPT-4/Claude3 直接输出“一线城市青年人才流入率下降 5%”的确定性结论，不标注置信度和不确定性来源，易误导用户，可能导致人才政策制定的决策偏差；文心一言/讯飞星火无不确定性管理功能，仅输出单一预测结

果，忽略数据和模型的局限性，无法保障预测结果的可靠性，难以支撑科学的人才政策制定。科研智能体中，所有工具均无不确定性管理功能，仅呈现固定结果，不考虑数据完整性和模型适配性，用户难以判断结果的可靠性和优化方向，无法支撑人才流动趋势研判的严谨性需求，可能影响人才政策制定的科学性。

（5）综合项目实证：“青年人才学术影响力预测”研究

以“青年人才学术影响力预测”研究为例，Talegent 可通过多步推理拆解研究流程，精准计算 h 指数、学科交叉度等核心量化指标，生成 3 个跨学科创新假设并设计科学的检验方案，同时通过不确定性管理标注结果可靠性和优化方向，形成完整的科研推理闭环。整个过程无需用户手动干预，可直接输出兼具科学性和创新性的研究结论，完全适配青年人才发展潜力预测的核心需求。

而 GPT-4 在推理过程中存在逻辑断层，量化指标计算错误率高，假设缺乏创新性，无法保障研究结论的科学性；文心一言仅能完成基础数据统计，无法进行复杂推理和假设检验，难以支撑深度研究；Elicit/Scite AI 仅能提取已有文献的观点，无自主推理能力，无法支撑创新研究；IBM Watson 无法适配人才研究场景，推理结果与研究需求无关，无法提供有效支撑。

7.4 智能体优势总结

通用大模型（如 GPT-4、Claude 3、文心一言、讯飞星火等）具备较强的通用语言能力，能够完成基础的文本生成与问答任务。但在人才研究这类专业性强、流程复杂的场景中，其局限性较为突出：一是对学术写作规范（如引用格式、术语一致性）支持不足，常出现虚构文献或格式错误；二是缺乏对长周期项目的上下文记忆能力，难以维持跨月甚至跨年度的研究连贯性；三是推理过程偏重表面关联，缺乏对因果机制、干扰变量等科研逻辑的严谨处理；四是多模态处理能力有限，无法有效整合政策文件中的表格、图表与文本内容。因此，通用大模型更适合承担初步信息整理等辅助性工作，难以支撑人才研究从数据采集到结论推导的完整流程。

科研智能体（如 Elicit、IBM Watson、Scite AI 等）通常聚焦于某一具体功能，例如文献检索、引文分析或特定数据库查询。这类工具在各自专长领域表现可靠，但功能边界清晰，缺乏横向协同能力。例如，Elicit 能高效提取文献摘要，却无

法将多篇文献观点整合为结构化综述；ResearchRabbit 可构建文献关联网络，但不能参与假设生成或量化分析。更重要的是，这些工具普遍不具备角色感知、多模态理解或动态反馈机制，难以适应人才研究中“政策制定—学术分析—基层执行”多角色协作、多类型数据融合的实际需求，仅能在局部环节提供有限支持。

Talegent：通过人才研究场景的深度定制，形成三大核心优势，实现对通用大模型和科研智能体的全面超越，成为适配人才研究全流程的专属智能体。

（1）场景适配的精准性：针对人才研究的多角色协作（领导、研究员、学生等）、多源数据整合（政策文本、学者成果、统计数据等）、长周期项目管理需求，构建了身份自适应、多模态记忆、实时反馈等专属能力，彻底解决通用工具“泛而不专”的问题，实现与人才研究场景的深度契合，适配政策分析、学者评价、人才流动趋势研判等核心场景。

（2）科研流程的全链路支持：从文献检索、文本生成、量化计算到逻辑推理、假设检验、结论呈现，实现科研全环节的无缝覆盖，避免专业工具“单点强、全链弱”的局限，大幅提升研究效率，使科研人员从繁琐的基础工作中解放出来，聚焦核心创新环节，支撑人才研究全流程的高效推进。

（3）结果输出的严谨性与创新性：通过格式规范校验、抗偏见控制、不确定性管理等机制，确保输出内容符合学术标准和伦理规范，弥补通用大模型“生成流畅但不够严谨”的短板；同时，依托跨学科推理能力、创新假设生成功能，为人才研究提供兼具科学性和创新性的解决方案，推动研究深度和广度的双重提升，助力人才研究从“静态描述”向“动态预测”转型。

综上，**Talegent** 并非对通用智能体的简单替代，而是基于人才研究专业需求的定制化升级，其核心价值在于将人工智能技术与人才研究的业务逻辑、科研规范、场景需求深度融合，提供“专业、高效、严谨、创新”的智能辅助工具，为人才研究的数字化、智能化转型提供核心支撑。

8 结论与展望

8.1 主要结论

当前智能体在人才研究领域的应用仍处于起步阶段，针对人才画像构建、流动趋势研判、政策深度分析等深层次需求的适配性不足，而通用大模型存在泛而不专、学术规范性欠缺等问题，现有科研智能体又面临功能单一、全流程支撑能力弱的局限。为此，本研究聚焦人才研究的实际需求，构建了专属智能体 Talegent，搭建人才领域专属知识库，开发智能问答、科研助理、辅助决策三大核心功能模块，并通过多维度对比与真实课题验证，探索智能体在人才研究领域的应用实践路径。

（1）通过开发人才领域专属知识库，整合政策文件、学术成果、人才榜单等多源异构数据，运用知识获取预处理、语义边界感知分割、向量化存储优化及智能检索更新等技术，结果表明，专属知识库有效解决了人才数据分散、结构化程度低的难题，使智能体回答错误率显著降低，所有结论均可追溯至权威原始来源，其知识关联化与动态更新能力，成为 Talegent 超越通用大模型的核心支撑。

（2）通过搭建智能问答、科研助理、辅助决策三大核心功能模块，集成深度语义理解、自主任务分解、多源数据融合、人机协同决策等关键技术，结果表明，模块可无缝覆盖人才研究从文献检索、数据处理到报告生成、决策支持的全流程，使数据收集、基础统计、初稿撰写等重复性工作的人力投入减少 40%，研究周期缩短近半，且输出内容的事实准确性、学术规范性均达到 98% 以上，大幅提升了研究效率与成果可靠性。

（3）通过从交互感知、文本处理、逻辑推理三大核心维度，将 Talegent 与通用大模型（如 GPT-4、文心一言）、专业科研智能体（如 Elicit、IBM Watson）开展系统性对比，结果表明，Talegent 在多模态数据适配、长周期项目记忆、多角色自适应输出、跨学科推理等方面展现出显著的场景化优势，能深度契合人才研究多角色协作、多源数据整合的实际需求，实现了对通用工具“泛而不专”和专业工具“单点强、全链弱”的双重突破。

（4）通过在真实人才研究课题中验证智能体应用效能，结果表明，Talegent

能有效支撑高层次人才选拔、区域政策痛点诊断、科研管理模式优化等核心场景，初步形成“智能体高效执行事务性工作 + 研究者聚焦创新性思考”的人机协同模式，不仅拓展了人才研究的数据广度与分析深度，更为人才研究领域的数字化、智能化转型提供了可复制、可推广的技术路径与实践范式。

8.2 存在问题

尽管用户体验和实证检验表明智能体在提升研究效率、保障结果准确性方面具有显著效果，但作为一项处于探索阶段的技术应用，其能力存在固有的、与现阶段技术发展水平相关的边界。对这些局限性进行系统审视，并非否定其应用价值，而是为了更准确地界定其适用场景，促进人机协同的合理分工。模块局限性主要体现在以下四个层面：

（1）创造性思维的局限

当前智能体的核心功能在于对既有知识的高效整合、信息的精准处理及既定框架的可靠执行。然而，在涉及理论建构、跨学科概念融合或前瞻性研究创新等任务时，模块表现出明显局限。针对“设计一个能平衡经济效益、社会公平与个体福祉的全新人才评价体系”等开放性议题，智能体生成的内容多表现为对现有观点的重组与延伸，难以产生具有突破性的原创思想。该局限源于其生成逻辑：基于大规模训练数据的概率分布与模式识别，擅长在既有规则框架内进行优化，而非创造新的规则体系。缺乏人类研究者所具有的直觉洞察、价值判断与创造力。

（2）自主化执行任务的透明度局限

智能体的自主任务分解与执行在提升效率的同时，也带来了过程透明度不足的问题，研究人员在特定阶段难以对进程实施有效监控。在执行复杂的多步骤文献综述任务时，用户难以直观了解智能体优先检索数据的原因、文献筛选的逻辑排序依据，也无法在中间步骤便捷地注入个人偏好的分析视角或临时调整研究方向。尽管智能体能够输出最终结果及基于检索片段的引用，但其内部的多智能体协作决策链条对用户而言仍缺乏透明度。

（3）对数据质量与完整性的深度依赖

智能体的分析能力建立在领域知识库的基础上，呈现数据决定产出的特征。当处理地方性未公开政策细则或新兴领域（其专业术语尚未被知识库完全收录）

的人才数据时，智能体的分析深度和准确性显著下降。其结论可能表现为表面化的通用描述，或无法捕捉数据中微妙的地方性、情境化特征。该局限源于 AI 的认知边界严格受限于知识库的覆盖面、颗粒度与时效性。对于知识库外的“未知”信息或隐性知识，缺乏有效的理解与推理能力。

（4）价值权衡与情境判断的局限

人才研究作为社会科学研究分支，涉及伦理、文化、政治等。智能体在此类需要深度情境化理解与价值权衡的任务上能力有限。针对“评估某项人才引进政策对本地社会长期影响”或“在科研诚信与创新激励之间设计平衡制度”等议题，智能体能够罗列各方观点与可能影响，但难以做出蕴含人文与社会关怀的综合性判断，无法替代人类在复杂伦理困境中的最终抉择。该局限源于 AI 不具备人类的价值观、道德感和对复杂社会情境的体悟能力。其分析呈现去情境化特征，基于历史数据的概率计算，难以完全理解和权衡那些难以量化的社会、文化与人本因素。

8.3 研究展望

结合 Talegent 的实践成效与现存局限，未来将围绕“深化协同、强化知识、完善体系、拓展应用”四大核心方向，推动智能体与人才研究的深度融合，持续提升专业服务能力，助力人才研究向数字化、智能化转型，为人才工作决策提供更坚实的智能支撑。具体方向如下：

（1）优化人机协同交互

开发可视化任务看板与细粒度干预接口，实现对智能体任务执行的实时监控与动态调整；提炼不同场景工作范式，融入智能体推理逻辑，打造“人主导、机执行、双向反馈”的高效协同模式，升级从“工具辅助”到“智能协作”的体验。

（2）丰富知识体系构建

持续拓展知识来源，不断更新升级知识库，定期定量填充优质人才研究相关研究报告、政策文件；建立与多渠道的实时数据对接，优化知识更新机制，强化对新兴领域、地方特色知识的覆盖，提升知识库的全面性与时效性。

（3）增强创新与推理能力

将人文社科研究逻辑融入模型训练，提升智能体跨学科推理与创新假设生成

能力；引入强化学习与专家反馈机制，优化分析结论的合理性；构建研究价值评估模型，增强复杂场景下的伦理判断与情境适配能力。

（4）拓展应用场景与范围

推动智能体向各级人才工作者延伸，适配基础研究、政策评估等多元场景；降低使用门槛，促进技术普及，通过不断发现问题、解决问题，推动智能体更快更好优化迭代。

参考文献

- [1] 陈永伟. 超越ChatGPT: 生成式AI的机遇、风险与挑战[J]. 山东大学学报(哲学社会科学版), 2023(03): 127-143.
- [2] 陈永伟. 从规模定律到规模经济: DeepSeek 的创新、机遇与挑战[J]. 山东大学学报(哲学社会科学版), 2025.
- [3] CUNTZ A, FINK C, STAMM H. Artificial intelligence and intellectual property: an economic perspective[R].2024.
- [4] 何雨昕, 谢晓. 教育智能体国际实证研究综述[J]. 教育信息技术, 2025(06): 35-41.
- [5] 谷飞. 元素养重构: 生成式AI与AI智能体在教师发展中的未来学视角[J]. 中国教育信息化, 2025, 06(31): 79-88.
- [6] 王好, 曾蓓, 郭力平. 打开教学决策的“黑箱”: 教师决策自动化评价智能体构建及应用[J]. 现代远程教育研究, 2025, 4(37).
- [7] 金丽琴. 教育智能体赋能项目式学习的关键环节[J]. 教学与管理, 2025(19): 26-30.
- [8] 陈永伟. 智能体经济的崛起: AI智能体对商业世界的重塑[J]. 财经问题研究, 2025: 1-17.
- [9] MCGRATH Q. Technology Trends to Watch in 2025[R].2024.
- [10] GenAI: The Path Forward[R].2025.
- [11] ROTHSCCHILD D M, MOBIUS M, HOFMAN J M. The agentic economy[J]. arXiv Preprint, 2025,2505:15799.
- [12] 陈泽金, 栾尚祯, 苏秀英. 智能体技术赋能国际贸易“单一窗口”数智化转型[J]. 中国海关, 2025(06): 86-87.
- [13] 俞陶然. 联影发布“元智”医疗大模型[N]. 解放日报, 2025-04-10.
- [14] 孙永会. 医疗智能体落地, 破解基层医疗痛点[N]. IT时报, 2025-05-23.
- [15] 成仲. 蚂蚁发布AQ智能体, 探路普惠医疗的中国路径[N]. 环球时报, 2025-06-27.
- [16] 李鹏飞, 吴静依, 郑悦闻, 等. 大语言模型驱动的医疗智能体研究与应用综述[J]. 中国卫生信息管理杂志, 2025, 22(2): 179-186, 194.
- [17] 景军. 什么是人工智能社会学?[J]. 智能社会研究, 2025(01): 1-20.
- [18] Jackelyn H, Dahir N, Sarukkai M. Curating Training Data for Reliable Large-Scale Visual Data Analysis: Lessons from Identifying Trash in Street View Imagery[J]. Sociological Methods & Research, 2023,3(52).

- [19] Thomas D. Start Generating: Harnessing Generative Artificial Intelligence for Sociological Research[J]. Socius, 2024(10).
- [20] Ricardo V, Azizpour H, Leite I. The Role of Artificial Intelligence in Achieving the Sustainable Development Goals[J]. Nature Communications, 2020,1(11).