

主动性智能体与企业AI工作流入门课纲公开版

面向新员工、管理干部、业务骨干和数字化转型推进人员，讲清大模型、Copilot、工作流、Agent 与 Skill 的区别，以及企业智能体试点的适用场景和风险边界。

公开边界

- 本资料为公开版课纲，不包含客户名称、内部业务材料、定制提示词库、账号配置、项目交付细节或任何可识别客户信息。

课程定位

- 本课程用于帮助企业先建立对智能体的理性认知，避免把智能体简单理解为“更高级的聊天机器人”。
- 课程不做单一产品说明会，而是围绕任务规划、工具调用、流程协作、权限边界和人工复核讲清企业落地逻辑。
- 适合作为智能体试点前的认知培训，也适合与企业AI工作流、知识库和Skill化工作坊组合交付。
- 公开版不包含客户内部账号、自动化脚本、系统配置、业务数据或定制交付文件。

适合对象

- 管理干部：需要理解智能体对组织流程、岗位分工和管理方式的影响。
- 业务骨干：希望判断本岗位哪些任务可以进入AI工作流或智能体试点。
- 数字化推进人员：需要识别低风险试点任务，并建立治理、权限和评估意识。
- HR与培训负责人：希望把智能体热点转化为团队可理解、可讨论、可行动的培训内容。
- 企业内训师：希望把智能体概念讲成能落到岗位任务的课程。

核心问题

- 员工听到Agent、Copilot、工作流、Skill等概念，但分不清它们分别解决什么问题。
- 企业想追智能体热点，但不知道哪些场景适合先试点，哪些场景风险过高。
- 很多团队直接上工具，却忽视账号权限、数据输入、日志审计和人工复核。
- 业务部门提出很多需求，但没有统一的判断标准，导致试点方向分散。

课程模块

- 模块1：从工具到智能体：讲清AI从单轮问答、内容生成，逐步走向任务规划、工具调用、多步骤执行和流程协作的演进。
- 模块2：核心概念区分：区分大模型、Copilot、工作流、Agent、Skill和自动化脚本的不同作用，用企业任务案例说明何时该用哪类能力。
- 模块3：热点案例观察：以主动性智能体和工作流工具为例，分析会话隔离、任务执行、工具调用、技能封装和本地控制等能力。
- 模块4：企业应用场景：讨论资料搜集、文档整理、流程跟进、知识库问答、报告生成、任务提醒等低风险试点任务。
- 模块5：风险边界与试点设计：围绕权限、数据、账号、审计、人工复核和责任边界，设计一个可控的智能体试点任务。

课堂产出

- 一张Agent、工作流、Skill、Copilot概念对照表。
- 一份部门智能体试点场景初筛清单。
- 一个低风险智能体任务设计草稿。
- 一套权限、数据、账号、审计、人工复核检查点。
- 一份后续升级为知识库或智能体原型的建议路径。

课后落地建议

- 先从低风险、低权限、可人工复核、结果容易判断的任务开始。
- 不要一开始追求全自动；先把输入、输出、复核、日志和责任边界定义清楚。
- 当某个工作流稳定后，再评估是否接入知识库、系统接口或智能体平台。
- 培训后可以继续进入Skill化工作坊、企业知识库建设或智能体MVP共创。

常见问题

Agent和普通大模型有什么区别？

大模型主要负责理解和生成内容，Agent更强调围绕目标进行任务规划、工具调用和多步骤执行。企业培训中需要先讲清边界，避免把Agent神化。

企业适合先做哪些智能体？

建议先从资料整理、制度问答、会议跟进、报告初稿、任务提醒等低风险场景开始，暂缓涉及付款、审批、敏感数据和强责任判断的场景。

智能体课程会不会变成工具演示？

公开版课程强调概念、场景、边界和试点设计，工具只是案例载体，不把课程做成某一个产品的操作说明。

培训后可以继续做开发吗？

可以。培训适合发现需求和统一认知，后续可把高频、低风险、标准化任务升级为知识库、工作流或智能体MVP。

数据安全怎么控制？

需要从账号权限、输入数据、知识库范围、日志审计、人工复核和责任边界六个方面控制，不建议把敏感数据直接交给外部模型。