

《AI驱动的审计与监察3.0》课程大纲

第1 | 课程名称、对象与课程核心

课程名称

《AI驱动的审计与监察3.0：智能体、DeepSeek、OpenClaw与审计工作升级》

一句话定位

这门课不是单纯教大家“怎么用聊天机器人”，而是帮助审计、监察、纪检、合规团队理解：

- 现在市面上常见的 AI 工具到底能做什么
- 哪些能力适合审计工作
- 哪些场景值得优先落地
- 哪些地方必须注意数据安全性与边界控制

课程适合对象

- 内部审计负责人及骨干
- 监察 / 纪检 / 合规团队
- 风险管理、内控管理人员
- 正在关注 DeepSeek、豆包、智能体、OpenClaw 等工具的企业管理者
- 需要和业务部门一起推进 AI 落地的 IT / 数据团队

课程关键词

- DeepSeek
- 豆包

- 智能体
- OpenClaw
- 审计自动化
- 数据安全
- 流程升级

第2 | 这门课具体讲什么

这门课不是泛泛谈 AI，而是具体讲五件事

1. 讲工具

用外部客户熟悉的语言，讲清楚这些东西和审计的关系：

- DeepSeek：为什么它会让审计分析、总结、解释更快
- 豆包：为什么它适合很多企业员工先入门、先体验
- 智能体：为什么未来不是只会聊天，而是开始能串任务、做步骤、调工具
- OpenClaw / 类 OpenClaw 工作台：为什么它让很多人兴奋，为什么它又不能乱上

2. 讲场景

这门课会讲几类最快上手场景：

- 报销 / 费用异常检测
- 招投标雷同与围标分析
- 举报 / 投诉 / 线索初筛
- 内控测试与流程检查
- 风险扫描与管理层简报
- 文档比对、制度比对、合同初步审阅

3. 讲流程

- 它怎么嵌进审计流程
- 哪一步可以自动
- 哪一步只能辅助
- 哪一步必须人工判断

4. 讲边界

- 数据能不能直接放进去
- 什么内容能上外部模型
- 什么内容必须本地化 / 私有化 / 脱敏处理
- 个人工作台和企业正式环境有什么区别

5. 讲落地顺序

帮助企业回答：

- 先做什么
- 谁来做
- 如何先做出小成果
- 如何避免一开始就做成 PPT 工程

第3 | 为什么企业现在需要这门课

过去很多 AI 课程的问题

本课程会避免：

- 太偏概念
- 太偏演示
- 太少谈场景优先级
- 太少谈数据安全
- 太少谈真实流程怎么接进去

不希望看到“同学听完热血，回去落不了地。”

现在企业真正面临的变化

- DeepSeek、豆包这类工具已经越来越多人接触过
- “智能体”这个词开始越来越常见
- OpenClaw 这类工作台 / 代理类工具已经让一些人开始尝试搭流程
- 但与此同时，数据安全、权限边界、模型幻觉、结果留痕这些问题也会一起冒出来

所以企业真正需要的不是再听一遍热词

而是建立一种更成熟的判断：

- 什么值得先试
- 什么不能乱试
- 什么适合面向个人效率
- 什么适合进入部门流程
- 什么必须纳入治理和数据边界设计

第4 | 课程模块与典型内容（上）

模块一：从大模型到智能体，审计工作到底变了什么

这一部分会讲清楚：

- 从“聊天式 AI”到“任务型 AI”的变化
- 为什么智能体会影响审计和监察工作
- 为什么很多人提效率，但真正的变化是流程结构被改写

模块二：主流工具怎么看——DeepSeek、豆包、OpenClaw

不是工具评测，而是从审计场景出发去看：

DeepSeek 能做什么

- 报告初稿
- 异常解释
- 文本总结
- 风险归纳
- 问题描述优化

豆包能做什么

- 员工入门体验
- 轻量知识问答
- 日常总结与整理
- 普通办公协助

OpenClaw / 类 OpenClaw 工作台该怎么看

- 为什么它适合做实验田、个人工作台、原型演示
- 为什么它不能直接等同于企业正式生产方案
- 为什么越强的能力，越要配套边界与隔离

模块三：异常场景

典型内容包括：

- 报销 / 费用异常检测
- 招投标文本雷同分析
- 关联关系排查
- 图片巡检 / 现场异常识别

这一模块会重点讲：

AI 能帮你找线索，但不能替你做最终认定。

第5 | 课程模块与典型内容（下）

模块四：任务场景

典型内容包括：

- 举报 / 投诉 / 线索分流
- 初步调查计划生成
- 法规条款与事实映射
- 合规问答与制度查找

这一部分解决的是：

很多高频、重复、文本密集型工作，能否先由 AI 完成“初筛、归类、建议”。

模块五：流程场景

典型内容包括：

- 内控测试
- 内控设计缺陷分析
- 费用审计流程嵌入
- 文档比对与制度比对
- 审计工作台 / 小型工作流设计

这一部分会讲清楚：

- AI 在流程中该放在哪里
- 哪些步骤适合自动化
- 哪些步骤必须人工复核
- 如何形成可回看、可解释、可留痕的结果

模块六：分析场景

典型内容包括：

- 管理层风险扫描
- 舆情与外部风险监测
- 行业研究与专项研究
- 风险周报 / 月报辅助生成

核心目标是让审计 / 合规团队不只是“发现问题”，而是能成为管理层的持续风险输入源。

第6 | 数据的安全与治理边界

这门课会专门讲一块：数据安全

因为一旦进入 AI 审计，最容易出问题的，不是模型回答慢一点，而是：

- 敏感数据外泄
- 不该上传的材料被上传
- 个人试用和企业正式使用边界不清
- 模型输出进入正式报告却没人复核
- 工具权限过大却没有控制

课程中会讲到的数据安全核心问题

1. 什么数据可以上外部模型

例如：

- 公开信息
- 脱敏后的样例
- 通用模板类内容

2. 什么数据要谨慎处理

例如：

- 涉及个人隐私的材料
- 商业秘密
- 审计底稿原件
- 纪检监察线索
- 未公开经营数据
- 合同、报价、客户信息等敏感材料

3. 企业常见误区

- 觉得“只是试一下”就没风险
- 觉得“内部人自己用”就没问题
- 觉得“结果挺准”就可以直接引用
- 觉得“本地部署”就自动安全

课程会给出的基本原则

- 效率提升不能凌驾于数据边界之上
- AI 可以辅助，但正式结论必须人工负责
- 工具能力越强，越要重视权限与留痕
- 个人实验环境与企业正式环境要分开看

第7 | 课程形式、交付结果与预期收益

可选授课形式

半天版

适合：

- 高管
- 审计 / 合规 / 监察负责人
- 需要快速形成认知的管理层

一天版

适合：

- 部门骨干
- 项目经理
- 审计、合规、内控、监察团队集中培训

两天版 / 工作坊版

适合：

- 希望明确筛出试点场景
- 希望形成原型方向或内部行动清单的企业

课程结束后，企业能带走什么

围绕本单位情况，至少形成：

- 1份 AI 审计场景优先级清单
- 2—3个适合先试的落地方向
- 1套“工具—流程—数据安全”基本判断框架
- 1份内部推进时的注意事项清单

企业参加后的实际收益

- 不再停留在“只会问模型几个问题”
- 能区分哪些工具适合展示，哪些工具适合试点
- 能识别哪些场景最容易先做出结果
- 能更早意识到数据安全和治理问题
- 能用更稳妥的方式推进 AI 审计和监察工作

最后一段总结

这门课真正要解决的，不是“大家知不知道几个 AI 名字”，而是让企业真正搞清楚：

在 DeepSeek、豆包、智能体、OpenClaw 这些新工具不断冒出来的情况下，审计与监察工作到底该怎么升级，哪些能做，哪些不能乱做，哪些值得先做。