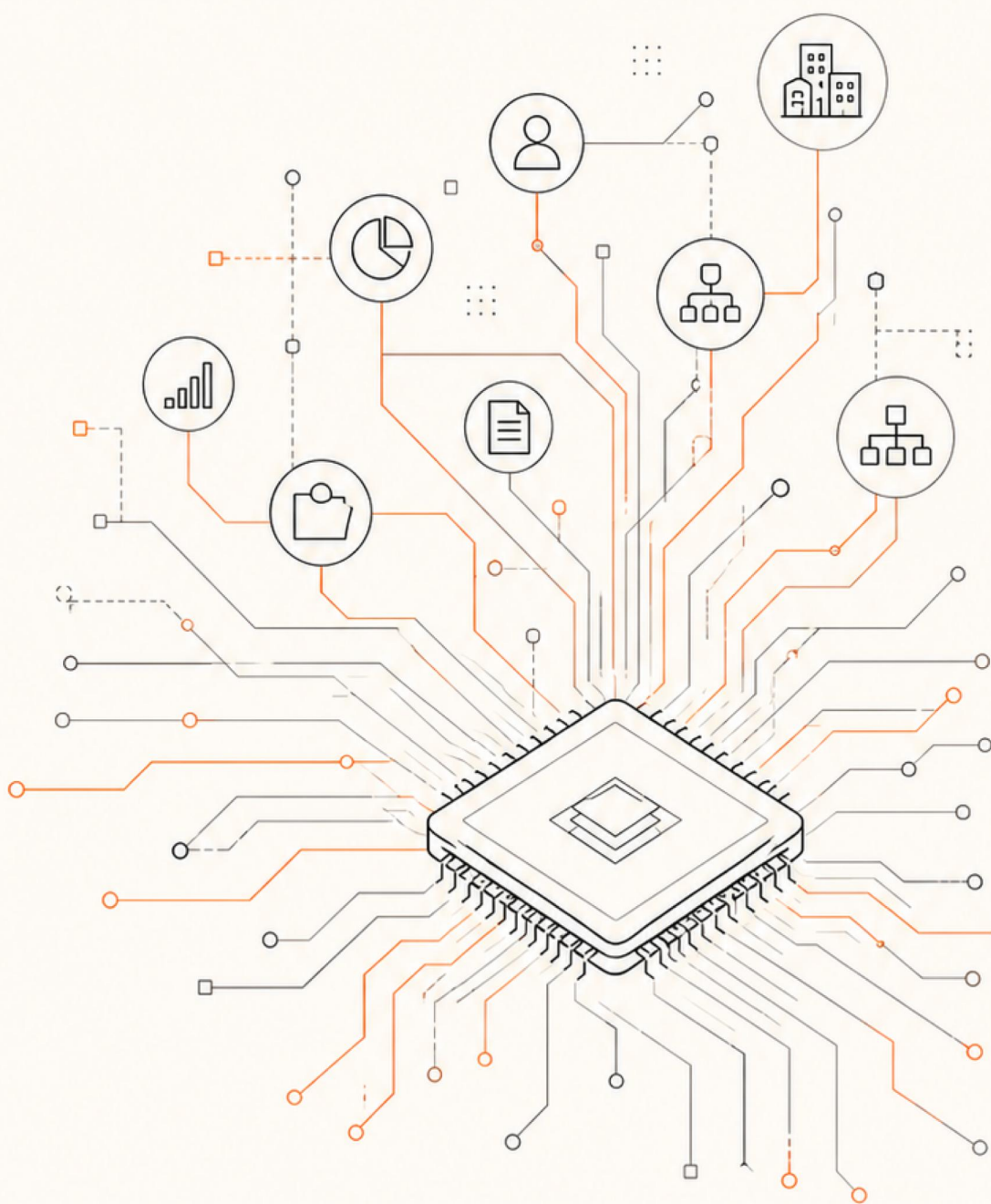


企业 AI 转型白皮书

给管理者的企业 AI 转型指南



版权声明

© 2026 曾俊 | 硅基行动

本白皮书由曾俊创作并发布，著作权归作者所有。

为促进企业 AI 转型知识传播，允许个人及企业在保留作品完整性、作者署名和来源信息的前提下，对本白皮书进行非商业性转发与分享。

未经作者书面授权，不得对本白皮书进行删改、改编、汇编或以他人名义重新发布；不得将其用于商业出版、付费课程、商业培训、咨询交付或其他直接营利用途。

引用本白皮书内容时，请注明作者、书名和来源。建议引用格式：

曾俊：《企业 AI 转型白皮书——给管理者的企业 AI 转型指南》，硅基行动，2026。

商业转载、培训使用及其他授权事宜，请联系作者或硅基行动。

欢迎将本白皮书的完整 PDF 转发给有需要的管理者和企业同仁。

目录

目录

前言

前言 这本书写给谁

第一篇 认知：AI 转型到底在转什么

第 1 章 AI 转型到底转的是什么

第 2 章 企业现在处在哪一层

第 3 章 AI 转型和数字化转型：不是一回事，但前后相连

第 4 章 为什么 AI 转型是一把手工程

第二篇 场景：从培训、知识库到 workflow

第 5 章 从 AI 培训到员工真正用起来

第 6 章 知识库为什么难点不是模型

第 7 章 从岗位经验到 AI workflow

第 8 章 AI 内容工厂如何改变生产方式

第 9 章 从单点工具到组织级协同

第三篇 经验：哪些坑最容易踩

第 10 章 企业 AI 落地最常见的 20 个坑

第 11 章 真实项目里总结出的 20 条经验

第 12 章 老板的 AI 项目风险识别清单

第四篇 路径：企业应该怎么开始

第 13 章 先培训、先诊断，还是先做智能体

第 14 章 如何拆解业务

第 15 章 哪些场景值得做，哪些不值得做

第 16 章 为什么企业应该先从人机协同开始

第 17 章 ROI 怎么算，才不容易被画饼

第五篇 判断：怎么看懂智能体、RAG 和技术方案

第 18 章 智能体到底是什么

第 19 章 如何判断一个智能体项目是真落地，还是 Demo

第 20 章 知识库 / RAG 项目怎么判断靠不靠谱

第 21 章 自建、采购、外包、平台化怎么选

第 22 章 安全、准确性和责任边界怎么把关

第六篇 组织：谁来牵头，怎么持续用起来

第 23 章 谁牵头，各角色怎么分工

第 24 章 怎么让员工持续用起来

第 25 章 不同部门和行业，应该从哪些 AI 场景切入

结语

结语 AI 转型的终点，是组织能力

附录

附录 企业各岗位 AI 应用场景速查表

参考

参考来源

作者

曾俊，北京硅基行动科技有限公司创始人 / CEO，前互联网大厂产品专家，拥有 8 年一线互联网大厂产品与复杂业务系统建设经验。现专注企业 AI 落地服务，主要提供企业 AI 实战培训、企业 AI 诊断与咨询、智能体定制、AI 落地推进。曾主导 100+ 个企业 AI 实战项目，覆盖制造业、餐饮、电商、内容生产等场景。擅长从企业真实业务出发，帮助团队建立 AI 认知、判断高价值场景，并形成可执行的落地方案。

联系作者

微信：ZF1235813266

公众号：曾俊 AI 实战笔记



微信二维码

前言 这本书写给谁

我见过很多老板，他们手里攥着一堆 AI 工具的采购单，他们的问题是“我要不要上 xxx”。这个问题其实问得太靠后了。更靠前一个问题应该是：你准备好为 AI 做什么了吗？

工具会越来越强，强到看着像一个能做事的人。但越像能做事的人，越需要你先搭好让它持续做事的条件。数据得是干净的、流程得是清楚的、出了错得有人管、责任得落到人头上。这些东西 AI 自己长不出来，得企业自己搭。

我把这种能力，叫“**承接能力**”。AI 的价值从来不会自动发生，它只发生在一个有承接能力的组织里。一家没有承接能力的企业，买再贵的模型、请再牛的供应商，最后大概率也只是多了一堆没人用的工具、几个进不了真实生产的 demo。

2026 年 4 月，斯坦福数字经济实验室出了一份报告，跟踪了 51 个真正落地成功的企业 AI 项目，横跨 9 个行业、7 个国家。他们最后总结一句话：**技术是能跑的，难的是技术之外的一切**。他这个观点和我做企业 AI 服务，接触了上百家企业的体感完全一致。

所以这本书从头到尾都在讲一件事：**企业怎么把 AI 接进自己的业务里，让它持续、安全、可控地产生价值。**

这本书写给企业管理者

这本书写给**真正掏钱、拍板、要为结果负责的人**，技术负责人也可以看，但不能把它当成技术手册，这本书不讲技术细节。

这本书只讲：**作为老板，你需要懂哪些判断，这件事到底靠不靠谱、要花多少钱、能挣回来多少、出了事谁来负责**。你脑子里盘旋的这几个问题，这本书就回答这几个问题。至于模型底层怎么跑、Agent 怎么搭，这是落地的人该懂的事。你要懂的是：怎么判断你请的人靠不靠谱，他们的规划能信几分。

这本书讲真实落地

此前我在某互联网大厂做风控系统的 AI，一个客诉问题原来要走完客服、工单、分析师一整条人墙，动辄一整套流程，我们用 AI 把它压到分钟级；一套账号风险体系，拦截效果明显提升，风险画像覆盖更完整，也帮助团队更早识别和处置高风险账号；一个多智能体的风控大脑，多个子 Agent 协作，让客诉量明显下降，也显著降低资损。

现在我做企业 AI 服务，接触了上百家企业，服务了几十家企业。包括餐饮、制造、跨境电

商、银行、政府，有做成的，也有做砸的、谈崩的、中途放弃的。

这本书讲的每一个项目、每一个客户、每一个数字，都是在一线项目里验证出来的。哪些项目成了、哪些项目没成、哪些项目半途放弃、哪些项目遇到了哪些坑等，这些问题你大概率也都会遇到。

第 1 章 AI 转型到底转的是什么

同一个词，五个老板说的是五件事

做企业 AI 服务后，接触了上百家企业，聊得最多的一个词就是“AI 转型”。但聊到后面我发现一个怪现象：十个老板嘴里的“AI 转型”，可能是十件完全不同的事。来看看几个场景，你可以对照一下自己是哪一种。

第一个，做电商的，跟我说：我们早就在做了，上个月请了个讲师，给全公司连续做了几场培训，教大家用 DeepSeek、写提示词、做 PPT，现场气氛很好。我问他，培训完两周，还有多少人在用？他愣了一下说，好像没怎么统计。

第二个，制造业的，说：我们买了一套 AI 系统，接了大模型，搭了个知识库，员工有问题可以问它。我问，问出来的答案准吗？他说，刚上线那阵大家挺新鲜，后来发现经常答非所问，慢慢就没人问了。

第三个，做连锁餐饮的，说：我们找外包做了两个智能体，演示那天我看了挺满意，效果不错。我问，现在真正用在生意上了吗？他说，demo 是挺好，但一放到门店真实场景里，问题就冒出来了，还在改。

第四个，做跨境的，干脆什么都没动，只是焦虑：同行在做、政府在推、抖音上天天刷到 AI 案例，他知道这事不能不看，但一放到自己公司，全是问号。先从哪个部门动手？先培训还是先上系统？怎么判断哪个方向值钱？怎么不花一笔冤枉钱、最后只做出一个没人用的 demo？

第五个，做女装品牌的，看了很多短视频，他张口就要一个“全自动、像人一样思考”的 AI。我们给了 demo，质量很好，她回了一句“不够 AI native”，就没了下文。她心里那个标准，是 AI 博主造出来的幻象，现实里没有任何一家公司能交付。

实际上：上面这些，除了第五个被带跑偏的，前四个每一个都有它的价值。培训对，搭知识库对，做智能体也对，焦虑着不敢乱动也比瞎冲强。

但它们都只是一部分。把任何一件单独拎出来当成“AI 转型的全部”，就错了。这是第一章想讲清楚的事，因为你对“转什么”的理解有多偏，后面踩的坑就有多深。这五家企业，后面都会在某一章里再出现一次，变成一个具体的坑。

AI 转型真正转的，是五样东西的连续推进

我把做培训、做企业知识库、做风控 AI、做内容 workflow、做智能体项目的经验攒在一起，得出一个结论：**真正的企业 AI 转型，是一套分层推进的过程。**

把它拆成五层，从浅到深。这五层在后面会一层一层拆开，这里先建立个整体印象。

企业AI转型五层架构



第一层，人的转型：员工敢不敢用、会不会用 AI，这是地基中的地基。企业做 AI 实战培训，核心干的就是这一层：把人从“只会打开对话框聊两句”的水平，往“能把 AI 嵌进自己日常工作”那一层推。这一层没过，后面全是空谈。

第二层，数据与知识的转型：企业的数据和知识，能不能被 AI 调用。我们给一家精密制造企业做 ISO 审核智能体，大部分功夫都花在把它那堆又脏又乱的历史文档，扫描件、复杂表格、流程图、带密码的压缩包。清洗成 AI 能读、能查的干净知识上。数据这一层不夯实，模型再强也是对着锅粥干瞪眼。

第三层，岗位与流程的 AI 化：一个个岗位、一条条流程被 AI 改造。我在某互联网大厂做的支付客诉智能体就是：把一个客诉问题从原来走完“客服→工单→分析师”的完整流程，压到分钟级。但前提是，得先把分析师脑子里的判断逻辑，一条条拆成了能写下来的流程。

第四层，组织协同方式的变化：AI 从个人工具变成组织能力。我们正在做的一个区域特色连锁餐饮项目，设计成分层智能体体系：从面向顾客的点餐，到面向店长的经营，再到面向总部的集团管理。AI 一路从单点长成了贯穿整个组织的协同方式。

第五层，业务模式与经营方式的变化：最后一层，是业务模式与经营方式变化。这是最深的一层。它解决的问题是：企业能不能从“用 AI 优化原业务”，升级为“用 AI 重新设计业务”。用 AI 写内容，用 AI 做客服，用 AI 做销售辅助，用 AI 做数据分析，用 AI 做合同审核。这些都重要，但本质上还是在优化原有业务。

更深一层，是 AI 改变企业的产品形态、服务方式、获客路径、销售方式、客服体系、供应链、内容生产和决策模式。比如一个跨境电商团队，原来做内容生产，要靠人工写脚本、拍摄、剪辑、发布。后来通过 AI 脚本、AI 配音、AI 剪辑、人工审核，把内容生产变成一条流水线。

这当然很有价值。但它还只是“用 AI 重做了一条流水线”。更深一步，是商业模式也发生变化。

比如从按视频数量收费，变成按转化效果收费；从卖人力服务，变成卖 AI 驱动的增长系统；从单次交付，变成持续运营，这才接近业务模式变化。第五层不是每个企业一开始都能做到，大部分企业目前还停留在第一层到第三层，少数企业开始进入第四层。

这五层最重要的一个性质：它是递进的，不能跳。人不会用，谈数据没意义；数据没打通，谈流程改造是空中楼阁；流程没改造，组织协同就无从谈起。

我经常给客户举一个最简单的例子：如果你的员工连豆包都没用过，你跟他聊智能体、聊业务 AI 化，没有任何意义。那是把第五层的话，说给还站在第零层的人听。

回头看本章开头那五个老板，你就全明白了。请讲师做培训的老板，在动第一层（人）；买系统、搭知识库的老板，在动第二层（数据与知识）；做智能体 demo 的老板，想够第三层（岗位流程），但往往因为前两层没夯实，demo 一进真实业务就垮；那个还在焦虑、什么都没动的老板，其实是最清醒的。他至少知道自己不知道；而那个被博主带跑、张口要“全自动 AI”的老板，是把第五层的幻象，当成了第一步的起点。

每个人都对，每个人又都只摸到了大象的一条腿。AI 转型是这五件事连续、有序地往下走。

AI 是放大器，不是发动机

讲了五层，一个贯穿全书的逻辑：**AI 是放大器，不是发动机**。它放大的是企业本来就有的东西，本来的能力，和本来的混乱。

我在某大厂做风控时，做过一个复杂项目，原始数据有几百个字段。一开始团队想把字段全喂给模型，让它自己学，结果效果很差，字段太多，噪音太大，模型抓不住重点。后来我们花了大量时间，和业务专家一起做字段筛选、口径统一、特征梳理，把真正关键的数据沉淀出来，效果才上来。

这件事教给我的不是“模型不行”，恰恰相反：模型很行，行到能把你的数据里的混乱，原原本本地放大给你看。你数据干净、口径统一，它就放大你的秩序；你数据散乱、规则打架，它就放大你的混乱。

前面提到的那家做 ISO 审核的制造企业，是另一个角度的同一个道理。他们有很多知识，但那些知识全是“写给人看的”，散在大量乱七八糟的文件里。你直接把这堆东西丢给再聪明的大模型，它照样答得驴唇不对马嘴。问题出在喂给它的材料本来就是一锅它消化不了的粥。AI 这面镜子，照出来的是这家企业多年来数据治理上欠的债。

理想汽车的李想讲过一个观点，他说 **AI 是照妖镜，也是放大镜**。强的人有了 AI 会更强，因为他能更快验证想法、更快迭代，但 AI 也会暴露一个人、一个组织的能力不足。一个人用 AI 提效十倍，上下游却跟不上，说明组织流程有问题；数据拿不到，说明数据治理有问题；成果没法评价，说明绩效体系有问题。（来源：李想×罗永浩公开访谈）

所以，AI 这面镜子，照的从来都是同一件事：你这家企业到底准备好了没有。

难是企业接不接得住

国内现在大量企业在用 AI，但 95% 的试点没回报。中间卡在哪？很多老板的第一反应是：是不是模型不够强？是不是该换个更好的工具、找个更牛的供应商？

我的观点是：几乎都不是模型的问题，这也不只是我一个人的观点，几乎所有认真做过调研的机构，结论都指向同一个方向。RAND 公司 2024 年访谈了几十位一线数据科学家，总结出 AI 项目失败的五大根因。排在第一的，是业务领导层对 AI 的误解，五条里几乎没有一条是纯技术问题。BCG 提了个 10-20-70 原则：一个 AI 项目能不能成，算法只占 10%、技术和数据占 20%、人和流程占 70%。斯坦福报告：技术是能跑的，难的是技术之外的一切。

我做一线项目的一个观点，和这些顶级机构完全同频：企业 AI 做不起来，70% 的功夫卡在业务和数据上，不在模型。那这“接不接得住”，具体接的是什么？我把它拆成五种能力，姑且叫**承接能力**。

第一，任务定义能力：你说不说得清，哪个岗位、哪条流程、哪个决策点适合 AI 介入、要它减什么成本、提什么质量。说不清，AI 项目就变成“大型许愿现场”。前面那个做女装选品的客户就栽在这。她要一个“像人一样思考的选品大脑”，但说不清到底要它在哪个具体环节、解决哪个具体问题。需求是一团模糊的幻想，再强的 AI 也只能跟着瞎猜。

第二，上下文整理能力：你的数据不是少，是乱；知识锁在人脑里、群聊里、旧文档里。那家制造企业的大量脏文档，就是最典型的反面教材：不先把这堆上下文整理干净，接 AI 进去，只会让它更快地、更理直气壮地胡说。

第三，工具连接能力：真实工作发生在 CRM、ERP、OA、邮件、各种业务系统里。AI 得能进去读信息、触发动作，才产生效率。我在某互联网大厂做的那套客诉智能体，客服原来要在 ERP 里提工单、等分析师；智能体的价值，正是替他打通了“遇到问题→自助拿到答案”这条原本要跨好几个系统、好几个人的链路。

第四，反馈迭代能力：AI 早期一定会出错，关键是有没有人“养”它，把错误变成下一轮改进。没人养，它就永远停在玩具阶段。这就是为什么那么多企业的试点，跑两周发现“不太稳定”就放弃了。问题在于没人负责把它的错误一轮轮喂回去、调出来。

第五，治理与责任能力：AI 生成的回复谁负责、自动流程出错谁兜底、敏感数据怎么审计。这决定了你能不能从“试试看”进到“长期用”。个人玩 AI，顶多自己信息乱一点；企业一旦让 AI 进真实工作流，这些责任问题答不上来，项目就永远卡在“试点”这一步，不敢往生产走。

你回头看这五种承接能力，没有一种是“买个更聪明的模型”能解决的。它们全是企业自己的事，是组织、是数据、是流程、是责任。

老板们最常问我的是“我们要不要上 AI、上哪个工具”。这个问题问得太靠后了，更靠前、也更值钱的问题是：你准备好为 AI 做什么了吗？

这一章，你该带走的判断

如果你是企业老板，这一章我希望你带走三条判断。

第一，别把任何单点动作当成 AI 转型的全部。培训、买系统、做智能体，都对，但都只是其中一层。真正的转型是五层连续推进，人、数据、流程、组织、业务。你现在卡在哪一层，下一章会帮你量出来。

第二，AI 是放大器，不是发动机。它放大你本来的秩序，也放大你本来的混乱。所以上 AI 之前，先问问自己：我这家企业，是有底子等着被放大，还是有一堆问题会被照得更清楚？两种都不是坏事，但你得心里有数，别指望 AI 替你把烂摊子自动收拾好，它只会把烂摊子照得更清楚。

第三，难的不是 AI，是承接能力。从麦肯锡到 BCG 到斯坦福，所有顶级机构和我的一线经验，都指向同一件事：70% 的功夫在人和流程，不在模型。你缺的不是更强的工具，是把工具接进真实业务、并让它持续产生价值的那套能力。

一句话

这章收束成一句话：**AI 的价值从来不会自动发生，工具越像“能做事的人”，越需要你这家企业先具备让它持续产生价值的承接能力。**

所以接下来你真正该问的：我现在在哪一层？下一步该去哪一层？哪件事最值得先做？把这三个问题想清楚，才是你 AI 转型真正的起点。后面每一章，都在帮你回答它们。

第 2 章 企业现在处在哪一层

先给你一把尺子

上一章我们提到，AI 转型是五层连续推进的过程。这一章我们把这把尺子展开，让你能拿它量自己。

我做企业调研时有个习惯，进门聊不了多久，心里就会给这家公司标一个位置：这家企业现在大概在第几层，这个判断靠的就是这五层模型。它不是一套理论，是我做培训、做知识库、做风控 AI、做内容 workflow、做智能体这一路攒出来的一张地图。



第一层，人的转型：员工敢不敢用、会不会用 AI，解决的是个人效率。

第二层，数据与知识的转型：企业的数据和知识，能不能被 AI 调用，解决的是底层能力。

第三层，岗位与流程的 AI 化：把一个岗位、一条流程，改造成“AI + 人”的新模式。

第四层，组织协同方式的变化：让 AI 从个人工具，升级成跨部门的组织能力。

第五层，业务模式与经营方式的变化：从“用 AI 优化原业务”，升级到“用 AI 重新设计业务”。

这五层有两个规律：

第一，它是递进的，越往下越深、越难、价值也越大。第一层动一动，员工效率提一点；第五层动起来，可能是整个生意的玩法变了。

第二，跳层往往出问题，这是我见过最普遍的坑，老板看了几个炫酷的智能体案例，热血上头，想直接从第一层跳到第五层“用 AI 重做业务”，结果发现底下两层（数据、流程）压根没准备好，项目一上真实业务就垮。底层没夯，上层的楼盖得越高，塌得越快。

第一层：人的转型，员工到底用没用起来

这是大多数企业最先看到、也最先动手的一层。就一句话：你的员工，敢不敢用、会不会用、用得好不好。能不能用 AI 写文案、做总结、查资料、做分析、生成 PPT、写代码、整理会议纪要、回答业务问题。这些都属于第一层，它解决的是个人能力问题。以前很多活儿要人从零干：写材料靠自己憋，查资料靠自己搜，做 PPT 靠自己排。现在 AI 先干第一轮，人再判断、修改、决策。

这一层的表现：公司买了 AI 工具、开了账号、甚至做了培训，但真正在用的，只有那么几个人，而且是偷偷用。大部分人听完培训该干嘛干嘛，工具账号躺在那儿吃灰。

做了培训，不等于这一层就达到了，仅约三分之一的员工，过去一年受过任何形式的 AI 培训，而受过培训的人里，有 65% 没法把培训学的东西，关联到自己的实际岗位上（来源：Docebo 调研）。

什么意思？就算你做了培训，大概率也卡在“听的时候热闹、回到岗位用不上”。这就是为什么很多公司 AI 预算打了水漂。不是工具不行，是人没真正进到使用状态。

我给一家新消费公司的财务部做 AI 培训，我把培训设计成三层，每一层都为了解决“用不起来”：

第一层，认知拉齐，先让大家搞清楚 AI 是什么、能干什么、不能干什么、怎么和它协作。认知不齐，后面全白搭。

第二层，工具学习，教真正好用的工具，并且讲清楚不同场景该选哪个，不是堆一堆工具名词。

第三层，业务实操，直接结合他们的应收、应付、报表、税务岗位，做场景实操。不是讲通用案例，是动他们手上每天的真活。

这套设计原因：AI 培训不是教工具，是解决“人能不能真正用起来”。你只教工具，员工回去还是不用；你得把他拉到自己的业务场景里，让他用 AI 解决一个他今天就头疼的问题，他才会真的进去。

判断你在不在第一层、过没过第一层，问自己一个问题就够了：抛开那几个 AI 爱好者，你公司里普普通通的中坚员工，有没有把 AI 用进了日常工作？如果没有，你还在第一层门口。

但就算过了第一层，AI 也只是员工的个人工具，还不是企业能力。个人效率提升了，组织整体可能纹丝不动。要往组织能力走，得继续往下。

第二层：数据与知识的转型，AI 项目做到一半，变成了治理项目

很多企业的 AI 项目都死在这一层，一开始 demo 很好看，老板看了也满意。可一进真实业务，问题全冒出来了：数据不完整、系统不打通、知识不结构化、权限不清晰、业务规则散落在不同人的脑子里。

这一层的表现是：“我们的 AI 项目，做着做着变成了数字化治理项目”，因为 AI 要真正进业务，前提是得有可调用的数据、可复用的知识。这两样没有，模型再强也跑不出价值。

很多企业明明有 ERP、CRM、OA、财务系统、客服系统，但系统之间没打通，数据口径不统一，关键业务信息还散在 Excel、聊天记录、邮件、个人电脑和员工脑子里。这种状态下，你接个再强的大模型进来，它也抓瞎。

前面提过那个几百字段的风控项目，一开始直接喂字段效果很差，后来花大量时间和业务专家一起筛字段、统一口径、梳特征，才把关键数据沉淀出来。很多 AI 项目最后拼的根本不是模型，是数据。

知识库也一样。很多企业以为，把文档全丢进知识库就完事了。真做的时候才发现，最难的是这几件脏活：

哪些文档过期了，得清掉？

哪些规则互相冲突，得裁决？

哪些内容要拆成最小知识单元，AI 才检索得准？

哪些问题该按“用户怎么问”来组织，而不是按原始文档目录？

哪些知识有适用边界，超出边界就不能用？

这些才是真正耗时间的地方，这块我在第 6 章会专门用一个案例来讲解。这里先记住：知识库的难点从来不是模型，是知识治理。

在组织上，通常要 IT 或数据团队牵头，但业务部门必须深度参与。IT 管系统、接口、权限、数据底座；业务管字段含义、业务规则、知识内容、使用场景。少了任何一方都做不成。我见过太多“技术团队建了个知识库，业务部门不爱用；数据平台搭起来了，没人知道怎么用”，最后系统又变成摆设。

有几个行业数据，仅 7% 的企业认为自己的数据“完全 AI 就绪”，数据孤岛问题占 56%、缺数据战略占 44%（Cloudera×HBR 2026）。斯坦福一份报告跟踪了 51 个成功项目的报告，得出几乎一样的数字：真正“数据完全就绪”的，只有 6%，注意，这是在已经做成的项目里。Gartner 甚至预测，到 2026 年，组织会放弃 60% 缺乏“AI 就绪数据”的项目。也就是说，数据这一层没补，大半项目会死在这儿。

但斯坦福那份报告还补了一句：数据没就绪，不等于不能上 AI。那 6% 之外的大多数成功项目，恰恰是一边上 AI、一边借它把脏数据洗出来的，大模型本身，就成了清洗数据的工具。这正是下一章要专门讲的“AI 倒逼数字化补课”。所以在第二层，你应该挑个场景，让 AI 和数据治理一起做、互相带，而不是等数据全做完再上 AI。

我最近帮一家区域餐饮连锁搭智能体，第一目标不在这个“智能体”本身多炫，真正想干的，是借这个项目把它的底层数据线上化、把数据治理做起来。等于借 AI 这把火，先把第二层的地基打了。地基打好，后面做经营分析、做业务模式升级才有可能。

第三层：岗位与流程的 AI 化，AI 成为流程里的固定节点

员工会用 AI 了（第一层），企业有了数据知识底座（第二层），下一步是真正进业务，把一个岗位、一条流程，改造成“AI + 人”的新模式。

以前是人找资料、人写总结、人做分析、人审材料、人填表、人跟客户。现在变成：AI 先做第一轮（初筛、生成、总结、分类、提醒、校验、推荐），人负责判断、修正、处理例外、最终拍板。

作为老板一定要记牢一点：AI 不是简单替代人，是成为流程里的一个固定节点。AI 可以干很多活，但最终结果仍然要有人负责。

这一层的表现，是一句我经常听到的话：“我们员工个人效率是提上去了，但公司整体好像没什么变化。”，这恰恰说明你停在第一层和第三层之间：AI 散落在个人手里，没沉淀成企业流程。

我在某大厂做过一个风控监控预警流程的改造，原来这活靠高级分析师的经验：盯监控、发现异常、归因分析、判断风险、出预警报告、流转处置。问题是每个人做法都不一样，经验全在脑子里，新人很难快速上手，这是“经验锁死在个人身上”。

我们第一步是先把流程拆成标准步骤：每一步的输入是什么、输出是什么、判断标准是什么。拆清楚之后才看得出来，哪些环节适合 AI、哪些环节必须人来判断。

这件事让我得出一个判断：很多企业 AI 推不动，表面看是 AI 不够强，实际是原来的流程就不清楚。流程不清，AI 不知道该嵌在哪、按什么标准干。所以这一层有固定的三步走：

流程系统化：先把靠经验、靠口头、靠个人习惯完成的工作，整理成 SOP。

工作流 AI 化：把 SOP 里适合 AI 的节点嵌进去，初审、分类、总结、提取、生成、校验、推荐。

岗位职责重构：重新定义这个岗位需要什么能力、人负责什么、AI 负责什么。

拿合同审核举个例子，方便你区分层次：如果 AI 只是帮法务先审一遍合同、指出风险条款、给修改建议，这就是第三层，改变的是一个岗位、一条流程的工作方式。

第三层是承重墙，我在第 7 章会用一个“流程大幅压缩”的完整案例，把“先 SOP、再 AI、最后才是智能体”这个顺序来讲，这里你先把这一层的位置记住就行。

第四层：组织协同方式的变化，从个人工具到组织能力

第三层解决的是单岗位、单流程。第四层开始进组织层面，解决的是：怎么让 AI 从个人工具，升级成组织能力。

这一层的表现，是“新型 AI 孤岛”：市场部有自己的 AI 工具，销售部有自己的 AI 表格，财务部有自己的分析模板，客服部有自己的知识库。每个部门都在试，但彼此不打通。结果就是各用各的、拼不成整体，这是比“没人用”更隐蔽的坑，因为表面看大家都很积极。

真正的第四层，是让 AI 进入跨部门协作、知识流动、经营分析、管理决策。AI 跨部门汇总信息、自动生成经营分析、辅助管理层发现异常、缩短汇报链路、把原来散在各部门的信息重新组织起来。

还拿合同审核说，如果 AI 只帮法务审合同，那是第三层。但如果它打通了销售、法务、财务、审批系统，让合同风险提前暴露，让销售在提交前就知道问题，让财务提前识别付款条款，让审批路径根据风险自动调整，这就进了第四层，改变的不再是一个岗位，是整个跨部门协同方式。

这一层动的是部门协作、知识流动、信息透明度、管理方式、决策速度。AI 转型真正有价值的地方，不是某个员工变强，是组织整体变强。第六篇我会专门讲组织怎么转，这里先把这一层标在你的地图上。

第五层：业务模式与经营方式的变化，用 AI 重新设计业务

最后一层，最深，也最少有企业够得着：企业能不能从“用 AI 优化原业务”，升级成“用 AI 重新设计业务”。

大多数企业用 AI 都是从提效起步：用 AI 写内容、做客服、辅助销售、分析数据、审合同。这些都重要，但本质上还是在优化原有业务。

再深一层，是 AI 改变企业的产品形态、服务方式、获客路径、销售方式、客服体系、供应链、内容生产、决策模式。

举个跨境电商的例子。原来做内容生产靠人工写脚本、拍摄、剪辑、发布。后来用 AI 写脚本、AI 配音、AI 剪辑、人工审核，把内容生产做成了一条流水线，这很有价值，但还只是“用 AI 重做了一条流水线”，本质还是第三、四层。

更深一步，是商业模式也变了：从按视频数量收费，变成按转化效果收费；从卖人力服务，变成卖 AI 驱动的增长系统；从单次交付，变成持续运营，这才接近第五层。

第五层不是每个企业一上来就能做的，它更多是一个方向。但它极其重要，长期看，真正拉开差距的，不是谁多用了几个 AI 工具，是谁更早用 AI 重新设计了自己的业务。

大多数企业在前两层，别幻想一步到位

一张速查表，你拿它对着自己企业一眼对号入座：

企业 AI 转型白皮书

层级	你大概率在这层，如果...	别误判：过了这层的真标准
一·人	买了工具、做了培训，但只有几个爱好者在用	普通中坚员工，把 AI 用进了每天的工作
二·数据与知识	demo 好看，一进真实业务就发现数据乱、系统不通	关键数据能被调用、知识治理成了 AI 查得到的资产
三·岗位与流程	个人效率提了，公司整体没什么变化	有岗位/流程被改成了“AI + 人”的固定模式
四·组织协同	各部门各用各的 AI，拼不成整体（新型孤岛）	AI 打通了跨部门，进了经营分析和管理决策
五·业务模式	还在用 AI 优化原来的业务	用 AI 重新设计了业务、甚至商业模式

看完这张表，你不用急着给自己贴一个精确标签，更合适的做法是看哪一层的问题最多。

如果第一层、第二层的问题最多，说明你现在最该补的不是智能体，而是人和数据：员工有没有真正用起来，关键数据和知识能不能被 AI 调用。这个阶段，先把培训、知识治理、数据打通这些地基补上。

如果第三层的问题最多，说明你已经过了“会不会用 AI”的阶段，下一步要拆流程。重点不是再买一个工具，而是把某个岗位、某条流程拆成清楚的输入、输出、判断标准，再看 AI 应该嵌在哪个节点。

如果第四层、第五层的问题最多，说明你开始碰到组织级协同和业务重构的问题。这个阶段，别再只看单点提效，要看跨部门数据怎么流动、责任怎么分、经营方式会不会被重新设计。

一句话：第一、二层补地基，第三层拆流程，第四、五层重构协同和业务。

对完这张表，再给你一个我的观察：绝大多数企业，现在还卡在第一层到第三层之间。少数开始摸第四层，第五层，大多还在纸面上。

这不是说你落后，这是常态，问题不在你在哪一层，问题在两件事：一是别误判自己的位置，二是别想跳层。

很多老板会高估自己，做了培训，以为过了第一层（其实员工没真用起来）；买了系统，以为过了第二层（其实数据根本没治理）。误判位置，比落后更危险，因为你会把钱投到错误的层级上，地基没打就盖楼。

斯坦福那份报告里有个数字，正好印证跳层的代价：61% 的成功项目，在成功之前都至少失败过一次。我见过的失败，相当一部分就是跳层跳出来的，底下两层还晃着，就奔着最炫的第四、五层去，撞墙、推倒、再来。那些失败的钱，最后都没写进对外宣传的成功故事里，但它们是真金白银烧掉的。与其跳上去摔一次再退回来补课，不如老老实实从脚下这层往上走。

也别想跳层，我见过太多老板看了个组织级 AI 协同的案例（第四层），或者听了个“AI 重塑商业模式”的演讲（第五层），就想直接上。我一般会先把他们拉回来问几个问题：你员工用起来了吗？你数据打通了吗？你有哪条核心流程是拆清楚、能写成 SOP 的吗？三个问题问下来，大部分老板会发现，自己连第二层都还没站稳。

所以，搞清楚自己的位置之后，你真正要做的就三件事，这也是这本书后面所有章节服务的目标：

我们现在在哪一层？（这一章帮你量）

下一步应该进哪一层？（顺着递进走，别跳）

哪件事最值得先做？（第 15 章会给你一套“给问题标价”的诊断方法）

这一章，你该带走的判断

第一，用这把五层尺子，先量准自己的真实位置，不是你希望的位置。做了培训不等于过了第一层，买了系统不等于过了第二层。诚实地量，比盲目乐观值钱。

第二，五层是递进的，跳层必出问题。底层没夯实就盖上层，楼越高塌得越快。大多数企业该做的，是把脚下这一两层站稳，而不是去够最炫的那一层。

第三，别因为“还在前两层”而焦虑。绝大多数企业都在这儿。AI 转型是马拉松不是百米冲刺，走对顺序、一层一层往下走，比一步到位的幻想稳得多。

一句话

这章收束成一句话：**AI 转型不是看你用了多炫的工具，是看你这家企业，扎扎实实走到了第几层。**

先量准自己在哪儿，再决定往哪走。这比看一百个别人家的案例都有用，因为别人在第四层做成的事，搬到还在第二层的你这儿，大概率水土不服。

第 3 章 AI 转型和数字化转型：不是一回事，但前后相连

一个让很多老板纠结的问题

我跟传统行业的老板聊 AI，聊到一半，常会被同一句话拦住：“曾老师，我们数字化都还没做完呢，能搞 AI 吗？是不是得先把数字化补齐了再说？”

问这话的老板，分两种心态。一种是真焦虑，觉得自己连 ERP、CRM 都没用利索，落后了一大截，AI 这班车是不是赶不上了。另一种是找台阶，潜台词是“数字化都没做完，AI 先缓缓也合理”，给自己一个不动手的理由。

这一章我就回答这个问题。结论先放这儿，省得你绕：AI 转型和数字化转型不是一回事，但前后相连。你不必等数字化全做完再上 AI；但数字化欠下的债，迟早要补，而且 AI 会逼你提前补。

我既不想吓退那个真焦虑的老板，也不想纵容那个找台阶的老板。下面分三层讲清楚。

AI 不等于数字化，但它踩在数字化的肩膀上

先把两个概念分清楚，不然后面全是糊涂账。

数字化转型，干的是“把线下搬到线上、把模拟变成数据”。你原来纸质开的单，变成系统里的电子单；原来靠人记的库存，变成 ERP 里的数字；原来口头交接的流程，变成 OA 里的审批流。它的核心成果是：业务被记录成了数据，流程被搬上了系统。

AI 转型，干的是“让数据和流程自己产生判断、产生动作”。在数字化攒下的数据和流程之上，AI 来做总结、分类、生成、预测、推荐、决策辅助。它的核心成果是：数据和流程开始“动脑子”了。

你品一下这两句话，关系就出来了，AI 是踩在数字化的肩膀上的。数字化负责把东西变成数据、搬上系统；AI 负责让这些数据和系统聪明起来。

那如果肩膀不结实会怎么样？我用回上一章那个画面：数据没打通、流程没线上化的时候，AI 进来就抓瞎。

这不是比喻，是我反复见到的实情。AI 要做事，得有东西可吃：它得能调到数据，得能读到流程，得能进到系统里触发动作。可很多传统企业是什么状态？历史报价散在不同人的 Excel 里，客户信息记在销售的手机通讯录里，业务规则压根没写下来、全在老员工脑子里，几个系

统之间互相不认识。这种地基上，你接个再强的大模型，它也只能干瞪眼，没有数据可调、没有流程可嵌、没有系统可连，AI 就是个空有本事使不上劲的人。

我前面提过的那家制造业集团，调研时财务负责人讲了个场景：老板想看经营看板，收入、成本、利润、现金、应收应付，哪个项目有异常。数据其实都在系统里，但分散在 ERP、财务系统、协同系统、还有一堆 Excel 里，口径还不统一。想做这个 AI 看板，绕不开的第一步，是先把数据口径理清、把系统对接打通，这恰恰是数字化的活。所以我当时没把经营看板列为第一期，就是因为它会先变成一个数据治理项目，重、慢、老板还看得不勤（一个月甚至一个季度才看一次）。

这就是“AI 踩在数字化肩膀上”最直白的意思：**有些 AI 场景，你想做也做不了，因为数字化这一步的债没还。**

数字化没做完，AI 也能先做

如果我只讲“AI 踩在数字化肩膀上”，你很容易得出一个错误结论：那我得先花两年把数字化全做完，再来碰 AI。

这个结论是错的，而且会害了你。

为什么错？两个原因。

第一，数字化“全做完”这件事，根本不存在。数字化没有终点，再大的企业也有没线上化的角落、没打通的系统。你要等它“全做完”，就永远等不到上 AI 的那天。等你慢悠悠补完数字化，AI 红利期早过了，你的对手已经用 AI 跑出去几条街。

第二，AI 和数字化可以一边做一边补，不必排队。你不需要把数字化当成 AI 的前置作业全部交完再交卷。完全可以挑一个数据相对齐、流程相对清的小场景，先把 AI 跑通，拿到结果、建立信心，同时借这个 AI 项目，倒逼着把它涉及的那一小块数据和流程补上。

我前面说的那家米线餐饮连锁，就是这个打法。它的数字化也没全做完，但我没让它先去补两年数字化课。我们直接上一个智能体项目，而这个项目的第一目标，恰恰是借机把底层数据线上化、把数据治理做起来。等于用 AI 这把火，烧着烧着就把数字化的地基也打了。数据治理好了，后面做经营分析、做业务升级，路就通了。

所以正确的姿势不必按“先数字化、再 AI”串行排队，而是“挑个小场景，AI 和数字化一起做、互相带”。AI 给你一个具体的、能见到结果的目标，逼着你把这块的数据和流程顺手补齐，这比你干巴巴地“为了数字化而数字化”，动力大得多，也精准得多。

这个餐饮项目是简述，再给你看一个例子，一家区域特色连锁餐饮，做区域特色餐饮，跨区域经营，有一定门店规模。

这家公司找我，表面上是想做一个“点餐推荐智能体”：顾客点菜时，帮忙推荐高毛利、库存充足、还能搭上“健康餐饮”养生概念的菜品。听起来是个点餐工具，但我一头扎进去就发现，真问题不在点餐。

它原来用的业务系统出了问题，已经处于半瘫痪、数据散乱的状态。数据是有，但乱七八糟、散在各处，对不上、也用不了。这种状态下，别说做智能体，就连最基本的经营数据都理不清。

所以这个项目从第一天起，我们定成了一次借智能体倒逼的数字化重构。点餐推荐只是表层动作，背后真正要补的是菜品、供应链、库存、会员、销售和健康餐饮知识这一整套数据底座。

这就是“AI 倒逼数字化补课”最完整的样子：老板要的是一个看得见的智能体，真正沉淀下来的，是一套能长期用下去的数据底座。智能体是面子，数据治理是里子；面子让老板愿意掏钱、看得到进展，里子才是这家企业 AI 转型真正的地基。

（这家公司的故事，后面讲“单点怎么长成组织能力”那一章，我还会从另一个角度再讲一次。这里你先记住这一层：它本质上，是一次用 AI 倒逼出来的数字化重构。）

AI 是放大镜，会把你的数字化欠债照出来

现在轮到那个想找台阶、想缓一缓的老板了。我前面给了台阶，但这台阶不是让你逃避，是让你别被吓住、赶紧上路。因为 AI 这件事，有个你躲不掉的特点，

AI 是放大镜，你一上 AI，它会立刻把你欠的数字化债照得清清楚楚。

这就是第 1 章那个判断在数字化上的具体表现：AI 一进来，最先照出来的，往往不是机会，而是你原来藏在数据、流程和系统里的欠债。

你数据散在十几个 Excel 里？AI 一调用，立刻发现口径对不上、同一个客户三个名字。

你业务规则全在老员工脑子里，没写下来？AI 一上，发现根本没有“标准”可学，老师傅说东新人说西。

你几个系统互不打通？AI 想跨系统取数，立刻卡死在接口上。

这些问题，是你本来就有的，只是 AI 把它照出来了，只不过平时被人肉填坑掩盖着、没人捅破。AI 一进来，等于做了一次全身体检，把这些藏在角落的欠债，一次性照到台面上。

我知道很多老板第一反应是难受：你看，还没见着 AI 的好处，先把我家底的乱摊子全抖出来了。

但我想让你换个角度看这件事：这件事可以当成一次体检。

你想想，这些数据混乱、流程不清、经验锁死的问题，是 AI 来了才有的吗？不是。它们一直在那儿，一直在悄悄给你放血，只是平时靠老员工的经验、靠人肉协调、靠“大家都懂的潜规则”硬撑着，你看不见、也算不出损失。AI 的价值之一，恰恰是逼着你把这些一直存在、却一直被忽视的问题，提前暴露、显性化。

照出来，你才有机会治。照不出来，它就一直慢性失血，等到哪天那个唯一懂规则的老员工离职了，整块业务直接断档。

而且我还要给你一个比几年前更乐观的消息：AI 不只把债照出来，它现在还能帮你还债。过去

你要把那些散在 Excel、扫描件、乱表格里的数据理干净，得靠人一份份手抠，慢得要命；现在大模型本身就成了清洗数据的工具。斯坦福那份报告里有个数字很能说明问题：在他们研究的成功项目里，88% 都是靠大模型，解锁了原本根本没法用的数据，扫描的图片、乱七八糟的表格、各种格式的历史文档，这些两年前还得堆一屋子人去手工整理的东西，现在 AI 能直接读、直接理。我在第 6 章会用一个制造企业案例，把“AI 怎么帮你把一锅粥一样的脏数据洗成能用的知识”讲清楚。

所以“AI 倒逼数字化”其实有两层意思：它一边逼你正视欠债，一边又递给你一把比过去快得多的铲子去还债。这也是为什么我说，现在恰恰是传统企业补数字化课最好的时机，逼你的人和帮你做事的工具，是同一个。

这件事我在某大厂踩过一个很深的坑，正好印证。当时我推一个大模型探索项目，季度末没拿出阶段性结果，绩效不被认可。复盘下来，根因是：探索型的成果难验证、成本高，但底层的“基础搭建”没人愿意先做。后来我想明白一件事，得先把可见的基础搭好（把策略线上化、监听邮件结构化、建表、建知识库），这些“不性感”的地基活先干了，后面做 AI 的人才不用重复踩坑。这其实就是在补数字化的债。AI 想跑得快，前提是有人先把路修平。

所以我对那个想缓一缓的老板的回答是：你可以不等数字化做完，但你不能用“数字化没做完”当借口一直不动。你越早上一个小 AI 项目，越早把欠债照出来，越早开始还。拖得越久，债越滚越大，等真要用 AI 的时候，地基的窟窿能把你拖垮。

这一章，老板该带走的判断

第一，AI 转型和数字化转型不是一回事，但 AI 踩在数字化肩膀上。数字化把业务变成数据、搬上系统；AI 让这些数据和系统聪明起来。地基不结实，AI 站不住。

第二，不必等数字化全做完再上 AI，那一天永远不会来。正确的打法是挑个小场景，让 AI 和数字化一边做一边互相带：AI 给你一个看得见结果的目标，倒逼你把这块的数据和流程顺手补齐。

第三，AI 是放大镜，会把你的数字化欠债照出来，这件事可以当成一次体检。那些数据乱、流程乱、经验锁死的问题，一直在给你放血，只是平时看不见。AI 逼它们提前暴露，是给了你治的机会。越早照、越早治，越主动。

一句话

这章收束成一句话：**别等数字化做完再搞 AI，也别拿数字化没做完当不搞 AI 的借口。挑个小场景上手，AI 会替你那些藏着的欠债照出来，照出来，你才治得了。**

数字化的债，AI 不会替你免掉，但它会逼你正视。这对一个想把企业往前推的老板来说，未必是坏事。

第 4 章 为什么 AI 转型是一把手工程

一句话先放这儿：你不下场，这事必死

这一章我说话会直一点，因为它太重要了。

我见过的企业 AI 项目，做成的和做砸的，差别往往不在预算、不在团队，也不在选了哪个供应商。差别在一把手，老板到底有没有自己下场。下场的，项目能往前拱；没下场的，项目八成原地打转，最后不了了之。

但我这里说的“下场”，不是开会强调两句、批个预算、转发几篇 AI 文章到管理群。这些只能叫“重视”，不叫“下场”。真正的下场，是你亲自去动那些别人动不了的东西：全面转型、跨部门资源协调、组织激励考核、流程与责任重划、数据权限开放、项目优先级、失败容忍度。

这些事情有一个共同点：都不是中层能单独拍板的事。中层可以做小工具，IT 可以接系统，HR 可以组织培训，业务部门可以提需求，外部供应商可以交方案。但真正到了“要不要跨部门调资源”“要不要改流程”“要不要开放数据”“要不要重设考核”“第一次失败后还要不要继续扛”这些问题上，最后都得回到你这里。

这就是 AI 转型为什么是一把手工程。因为 AI 转型要动的，恰恰是只有一把手才有权力动的那几根柱子。

有些数字也能说明这个问题，麦肯锡 2025 年的研究提到，AI 高绩效企业里，高层主动倡导的可能性是其他企业的 3 倍，近 30% 的企业，是 CEO 直接负责生成式 AI 的治理。还有一个被广泛转引的数字：CEO 持续参与的项目，成功率 68%，一旦失去高管支持，成功率掉到 11%，68% 对 11%，这就是“一把手不下场”的分量。

但数字只是表象。真正的底层原因是：AI 转型会碰到组织里最硬的几块骨头，而这些骨头，只有一把手啃得动。

AI 转型不是工具项目，是组织重塑

很多老板一开始会把 AI 转型想简单：买几个工具，做几场培训，搭几个智能体，让员工效率高一点。这些都对，但它们只是表层动作。真正的 AI 转型，往深了走，一定会变成组织重塑。因为 AI 不只是帮人快一点，它会重新切开企业里的工作方式。

以前一个工作靠人从头到尾做完，现在可能变成 AI 先做一轮，人再判断。以前经验锁在老员工脑子里，现在要写成流程、沉淀成知识库。以前一个部门内部就能完成的事，现在要跨部门调数据、接系统、改权限。以前靠人肉协调的流程，现在要变成系统里的固定节点。到了这一步，你动的就不是工具，而是组织。

岗位会变，原来一个员工的价值，是亲手把事情做完。现在 AI 可以做初稿、初筛、提取、汇总、分析，员工的价值就变成判断、复核、兜底、处理例外。这个岗位到底还需要几个人？需要什么能力？绩效怎么评？这些是组织问题。

流程会变。原来一个审批、一个审核、一个客诉，要在人和人之间流转。AI 介入后，可能会变成系统先判断，AI 先生成建议，人只处理例外。流程节点变了，责任边界也要变。谁有权让 AI 先审？谁负责复核？错了谁兜底？这些是管理问题。

数据要流动，AI 要真正做事，必须读数据、查知识、接系统。可是企业里的数据往往被不同部门握在手里。销售有客户数据，财务有收款数据，客服有投诉数据，供应链有库存数据。AI 想把一件事做完整，就得跨部门拿数据。谁能决定这些数据能不能打通？不是 IT，是你。

利益也会被重新分配，AI 把一个部门的人效提高了，原来的人还留不留？省下的人力怎么安排？哪个部门配合了算功劳？哪个部门被改造了会不会觉得自己被削弱？这些事如果不说清楚，中层一定会犹豫，员工一定会害怕。

所以，AI 转型表面看是工具，往深处看是组织重塑。而组织重塑这件事，从来不是一个项目经理、一名 IT 负责人、一个外部供应商能完成的。它必须由一把手压住方向、利益和节奏。

老板焦虑、中层犹豫、员工迷茫：三股劲不往一处使

为什么 AI 转型非得一把手亲自压？因为组织里不同的人，对这件事的心思根本不一样，甚至是拧着的。

老板是焦虑的，你坐在最高处，看得见外部变化。同行在用，客户在问，政府在推，媒体天天讲 AI。你怕错过，怕被颠覆，所以你最想要的是结果：快点见效，证明这事值。

中层是犹豫的，这是最关键、也最容易被忽视的一环。中层不是不懂 AI 的价值，是他算的账和你不一样。AI 改造流程，往往意味着动他管的那摊子。可能要裁掉他手下的人，他的“地盘”缩水；可能要暴露他这块业务过去的低效，他的“绩效”打折；可能改到一半出了问题，他要背锅。对中层来说，推进 AI 是“风险大、收益不确定、还可能砸自己饭碗”的事。所以他表面配合，实际能拖就拖、能避就避。

员工是迷茫的，最底下的员工，第一反应是怕，怕 AI 把自己替代了。你让他用 AI 提效，他心里嘀咕：我把活儿干得越快、越能被 AI 取代，岂不是越早失业？你让他把自己的经验交给 AI，他会本能地防御：我把经验都交出去了，我还有什么价值？

你看，这三股劲：老板要结果，中层怕担责，员工怕被替代。三股劲不往一处使，项目就散了。

只有一把手能跟中层把话说清楚：我不是要裁你的团队，是要让你这个团队干更值钱的事；配合 AI 改造，算你的功劳，不算你的过。只有一把手能给员工吃定心丸：用 AI 不是为了立刻替代你，而是把你从重复劳动里解放出来，让你去做更需要判断、更有价值的事。

这些话，中层说没用，因为他自己都犹豫；HR 说没分量，因为他不掌握业务生杀；外部顾问

说更没用，因为他不掌握组织利益。只有你这个拍板的人说，才压得住。

最隐蔽的坑：大家都在做小工具，没人敢碰流程

另一个更普遍的坑，它直接解释了为什么很多公司“看起来在做 AI，其实啥也没变”。

现象是这样的：你去看一家“很重视 AI”的公司，会发现到处是 AI 小工具。这个部门用 AI 写周报，那个部门用 AI 做表格，市场部用 AI 出文案。热热闹闹，人人都在用。但你再看它的核心业务流程，纹丝不动。

为什么会这样？因为组织的激励，是奖励“看得见的产出”的。用了多少次 AI、做了几个小工具、培训覆盖了多少人，这些好统计、好汇报、好邀功，做了立刻有人看见。而真正值钱的事，啃下一条核心流程的 AI 改造，又难、又慢、又得罪人、还可能失败。

一个理性的中层，会怎么选，当然选做小工具。投入小、风险低、还能交差。没人傻到去碰那块又硬又烫、搞砸了要背锅的流程改造。

结果就是：整个组织把劲儿都使在最容易、最不痛不痒的地方，真正能改变业务结果的硬骨头，谁都绕着走。一年下来，AI 工具一大堆，业务流程一动没动。老板还纳闷：钱花了，人也用上了，怎么生意没什么变化？

我在某大厂亲历过这个坑的另一面，能帮你看清它的本质。当时我想做一件“难而正确”的事：建风险手机号画像。这事价值大，但研发不愿意做。他们觉得这是脏活累活，不信任准确率，而且它不在研发的 OKR 里。做了没考核、没奖励，凭什么帮你？

这就是激励错位的现场版：一件对公司整体有大价值的事，因为不在某个团队的考核里、短期看不见产出，就推不动。我当时怎么破的？先自己筛了一部分数据、构建策略、跑出可见的应用效果，证明它真有用。然后在季度末和领导沟通，把这件事正式放进 OKR。放进 OKR 之后，资源才到位，事情才真正往前走。

我想让你从这个例子里看到的是：“难而正确”的事，靠基层自觉是推不动的，因为激励不在那儿。必须有人有权力去重设激励，把它放进考核，给它配资源，给做的人撑腰。在一个部门，这个人是部门负责人；在一家公司，这个人只能是公司一把手。

你不亲自去碰那块流程、不亲自去重设激励，你的组织就会本能地躲开所有硬骨头，只在边角料上热闹。

为什么这些事只有一把手能做

为什么必须一把手下场？

AI 转型不是一个工具项目，而是一场组织变革



到这里，我们可以把“一把手工程”讲得更具体一点。AI 转型里，有七件事，几乎天然只能由一把手推动。

第一，全面转型。AI 如果只停在某个部门、某个工具、某个小助手上，它当然可以由部门自己做。但只要你想让 AI 真正进入业务流程、管理方式和经营系统，它就不再是局部项目。它会牵动人、流程、系统、数据、考核、预算。它要求公司从“局部试用 AI”，走向“系统性重构工作方式”。这个判断，只有一把手能定。因为只有你看得见全局，也只有你能决定：这件事不是某个部门的兴趣项目，而是公司级的转型动作。

第二，跨部门资源协调。AI 项目一旦进入真实业务，很少只属于一个部门。做客服 AI，要客服、产品、IT、法务、数据一起配合；做财务 AI，要财务、业务、系统、权限一起配合；做经营分析，要销售、供应链、财务、数据全部打通。问题是，每个部门都有自己的目标、资源和边界。你让一个部门把数据开放给另一个部门，让一个团队为另一个团队的项目投入人力，让 IT 优先排这个需求，凭什么？中层协调不了这个，外部供应商更协调不了。只有一把手能把部门墙打穿，把人、钱、系统、数据重新调到同一张桌子上。

第三，组织激励考核。组织不会自动做正确的事，组织只会做被考核的事。你嘴上说 AI 重要，但考核里没有，奖金里没有，晋升里没有，那它就永远排在后面。中层为什么喜欢做小工具？因为好汇报。为什么不碰流程？因为流程改造难、慢、容易失败，还不一定算自己的功劳。所以，AI 转型要真正推进，必须进入激励体系。做 AI 改造的人，功劳怎么算？配合开放数据的部门，收益怎么算？流程被重构后，人效提升算谁的？试点失败了，是不是一票否决？员工把经验沉淀给 AI，是不是会反过来威胁自己？这些问题不解决，组织就不会真动。而激励考核这件事，只有一把手能改。

第四，流程与责任重划。AI 不是贴在流程外面的装饰，它一旦进入业务，就会改变流程本身。以前是人初审，现在 AI 先初审；以前是员工查资料，现在 AI 先检索；以前是主管逐单判断，

现在 AI 先给建议，人只处理例外。这时候必须回答几个问题：AI 处理到哪一步？人从哪一步接手？什么情况必须转人工？错了谁负责？谁有权推翻 AI 的建议？谁负责持续维护规则和知识？这些是流程和责任问题，如果不重划责任，AI 项目就只能停在“辅助看看”的阶段，不敢真正进入生产。而流程与责任怎么重划，只有一把手能拍板。

第五，数据权限开放。AI 要做事，必须有数据。但企业的数据往往散在各个部门、系统和人手里。财务有财务的数据，销售有客户的数据，客服有投诉的数据，供应链有库存的数据，老板想看的经营分析，往往需要把这些数据串起来。问题是，数据开放从来不只是技术问题，它背后是权限、安全、利益和责任。谁能看哪些数据？哪些数据可以给 AI 用？哪些必须脱敏？哪些只能本地部署？数据开放后出了问题谁负责？部门愿不愿意把自己手里的数据拿出来？这些事，不是一句“IT 打通一下”就能解决的。没有一把手拍板，数据权限很容易卡死在部门边界里。最后 AI 没数据可用，只能变成一个聊天框。

第六，项目优先级。AI 能做的事情太多了。会议纪要能做，知识库能做，客服能做，销售能做，财务能做，经营分析也能做。问题是资源永远有限，不可能什么都同时做。所以老板必须定优先级：先做哪个？先不做哪个？哪个小场景先跑通？哪个大系统暂时不碰？哪个项目哪怕看起来很炫，也先放一放？这件事不能交给部门自己决定，因为每个部门都会觉得自己的需求最重要。也不能完全交给供应商决定，因为供应商天然倾向于做他最擅长、最容易卖、最容易交付的东西，不一定是你公司最该做的东西。项目优先级必须回到经营视角：哪个问题最值钱？哪个最容易落地？哪个能建立组织信心？哪个失败了代价最低？哪个跑通后能复制？这件事，只有一把手能从全局拍板。

第七，失败容忍度。AI 项目一定会有失败。可能模型答错，可能员工不用，可能数据太乱，可能流程接不进去，可能第一个试点没有达到预期。真正拉开差距的，不是会不会失败，而是失败之后组织怎么反应。如果第一次失败，老板就开始追责、换人、停项目，组织马上会学会一件事：AI 项目是高风险事项，最好别碰。从此以后，大家只敢做最安全、最边缘、最容易汇报的小工具，不会再碰核心流程。但如果一把手提前说清楚：试点允许失败，失败要复盘，复盘后继续迭代；只要方向对、问题真、过程透明，就不因为一次失败否定团队。组织才敢去啃真正有价值的硬骨头。失败容忍度，本质上是一种组织安全感。这种安全感，也只有一把手给得出来。

一把手自己得是深度用户：否则你定不了方向

前面讲的，都是“一把手要推动组织”。但还有一件事更底层：一把手自己得是 AI 的深度用户。为什么？因为你要定方向、配资源、压组织，而这些动作都建立在一个前提上：你得对 AI 有判断力。

一个不亲自深度用 AI 的老板，对 AI 的判断大概率会两头错。要么高估，你听了几个炫酷案例，以为 AI 什么都能干，上来就要做一个“AI 大脑”统管全公司。结果发现数据没打通、流程没理顺、责任没人兜底，钱砸进去只做出一个漂亮 Demo。要么低估。你觉得 AI 不就是个聊天机器人、不就是高级搜索，没什么大不了。于是你不敢投、不愿动，眼睁睁错过窗口期。

这两种误判，都来自同一个根源：你没真正用过，所以不知道 AI 的能力边界在哪里。过去的软件，功能是固定的。你学会 Excel，大概就知道 Excel 能干什么、不能干什么。但 AI 不一样，它是泛化的。同一个工具，不同的人用出来差别巨大。有人拿它问问题、写文案，有人拿它写代码、做调研、搭 workflow，最厉害的人拿它重构整条业务流程。

所以，对 AI 的认知，没法靠听报告、听分享、上培训获得。真正的认知只来自一件事：你自己在工作里持续用它，用到知道它能干什么、不能干什么、哪里容易错、怎么给它上下文、怎么拆任务、怎么验证结果。

我给你一个粗糙但有效的自检指标：你敢不敢看自己一个月的 token 账单。token 可以粗略理解成 AI 处理文字的计量单位，烧的 token 越多，说明你让 AI 干的活越多、越重。一个只会聊天式使用 AI 的老板，手动一问一答，问个问题、改段文案、查点资料，一个月烧不了多少 token，因为人手动一问一答，速度有上限。一个真正深度用 AI 的老板，会用 Agent、用 workflow 批量解决问题，让 AI 一次性跑几十上百个任务、连续工作几十分钟。这个时候，token 消耗会比浅层使用高几个数量级。

我不是要你纠结“多少 token 才合格”。随着模型越来越便宜，这个数字早晚会变。但它背后的逻辑不会变：深度用 AI 的人，一定是在用它批量地、连续地、重构式地解决问题，而不是手动一问一答。token 账单不是门槛，是镜子。它照出来的是：你到底是在浅层体验 AI，还是已经真正潜进去了。

你不潜进去，就没有判断力。没有判断力，你定方向、配资源、压组织，全是凭感觉拍脑袋。这才是“一把手自己要用 AI”的真正原因。不是为了显得先进，而是因为你不亲自用，就不知道该把公司带到哪里。

光关心项目还不够，要把 AI 放进组织系统

讲到这里，你可能会问：那一把手到底要做到什么程度，才算真正下场？我给你一个简单分层。

第一级，是被动审批。批了预算，然后全权下放，基本不再过问。这种叫“我支持”，不叫“我下场”。

第二级，是周期监督。每月看一次进展，出了大问题才出面。这比第一级好，但本质还是被动。

第三级，是主动驾驶。每周盯进度，主动扫清障碍，参与关键决策。这个阶段，项目往前走的概率已经大很多。

第四级，是战略整合。把 AI 采纳、流程改造、效率提升、知识沉淀，直接写进公司的目标、OKR、考核和奖金里。让 AI 不再是“某个项目”，而是变成组织运行方式的一部分。

真正能做到全组织级转型的，一定不是停在“主动驾驶”，而是进入“战略整合”。区别在哪里？主动驾驶是：我很关心这个项目，我帮你扫障。战略整合是：我把 AI 变成了衡量组织成功的指标。做不做、做得好不好，直接关系到每个人的考核和利益。

你回头看前面那个风险手机号画像的例子，它为什么从推不动变成推得动？转折点就是把它放进了 OKR。这说明什么？光老板关心没用，得让 AI 进入考核、挂钩利益，组织才会真的动起来。真正的一把手工程，不是你坐在上面批钱、表态、偶尔问问，而是你把 AI 变成全公司要一起背的目标，而且做好了“第一次可能失败、你还得继续扛”的准备。

只有一把手能解的三件事

讲了这么多，我把“一把手工程”最终落到三件具体的事上。这三件事，组织里谁都替代不了你。

第一，定方向。AI 转型要往哪走，先打哪个场景，要什么、不要什么，这是战略选择。中层定不了，他只看得见自己那一摊；IT 定不了，他容易从系统和工具出发；外部顾问能帮你分析，但最终拍板的是你。你要决定的不是“用不用 AI”，而是先做提效，还是先做增长；先做培训，还是先做诊断；先做知识库，还是先做工作流；先做低风险小场景，还是啃核心流程；什么事情坚决不让 AI 直接拍板。方向定错了，后面越努力，偏得越远。

第二，配资源。AI 转型要钱、要人、要时间，还要从现有业务里抽调骨干。这些资源的调动，只有你能拍。很多真正值钱的项目，往往不是某个部门当前最愿意做的事。它可能跨部门，可能短期不显功，可能会暴露旧流程的问题，可能会让某些人不舒服。没有你配资源，硬骨头永远没人啃。

第三，压组织。老板焦虑、中层犹豫、员工迷茫，这三股拧着的劲，只有你能把它们拧到一处。你要跟中层把利益说清，给他重设激励；给员工吃定心丸，打消“被替代”的恐惧；该调整的考核去调整，该重划的责任去重划，该开放的数据去推动开放。这种“压”，不是简单施压，而是带着权力、承诺和兜底的组织推动。

定方向、配资源、压组织，这三件事，外包不出去，下放不下去，全压在你一个人肩上。这就是“AI 转型是一把手工程”最实在的含义：不是让你喊口号、表态度，是让你亲自下场定这三件事。

这一章，老板该带走的判断

第一，AI 转型不是工具项目，而是组织重塑。它会动岗位、流程、数据、权限、考核和责任，所以天然不是某个部门能独立完成的事。

第二，一把手不下场，组织会本能地躲开硬骨头。中层会选择小工具，员工会害怕被替代，部门会守住自己的数据和资源。你不重设激励、不压组织，AI 就只能停在边角料上热闹。

第三，只有一把手能动七个关键杠杆：全面转型、跨部门资源协调、组织激励考核、流程与责任重划、数据权限开放、项目优先级、失败容忍度。这些事，谁都替不了你。

第四，一把手自己必须深度用 AI。不是为了显得先进，而是为了形成判断力。你不亲自用，就

不知道 AI 的能力边界，定方向就一定容易偏。

一句话

AI 转型是一把手工程，是因为它要动全面转型、资源、考核、流程、责任、数据、优先级和失败容忍度；这些关键杠杆，只有一把手能拍板、能统筹、能兜底。

第 5 章 从 AI 培训到员工真正用起来

培训当天热闹，两周后没人用

先看一个很常见的场景。

你花钱请了讲师，给团队做了一场 AI 培训。当天气氛很好，大家听得频频点头，现场演示一个 demo，有人惊呼：“原来 AI 还能这么用”。

培训结束，你觉得这钱花得值，团队总算“懂 AI 了”。然后呢？两周过去，你回头看，几乎没人在真用。还是老样子，该写的报告手写，该做的表格手做，AI 那点新鲜劲过去了，什么都没变。

这一章主要讲：培训这件事本身没错，错的是你以为听完课就等于用起来。这中间隔着一道大多数老板没看见的鸿沟。

我做企业 AI 落地这大半年，培训是接触老板最多的入口。我自己下场带的培训不少，有大厂的、有银行的、有传统行业的。我越来越确信一件事：绝大多数 AI 培训之所以听完就忘，很多时候是培训目标一开始就定错了。目标定成了“让大家知道 AI 很厉害”，而不是让大家手里多一个能天天用的工具。

这两个目标，差着十万八千里。

我用两家做过的对照案例，来跟你讲清楚这道鸿沟到底是什么、该怎么跨过去。一家是某头部电商平台，财务团队；一家是某四大行之一，体制内银行，覆盖多个部门和网点。两家的出发点完全不同，但要跨的那道坎，是同一道。

在讲怎么做之前，得先讲清楚“员工到底卡在哪一层”。因为如果你不知道员工现在在哪，你就不知道培训该把他们送到哪，也就听不懂这一章真正想给你的判断。

原来什么样：大部分员工卡在“只会聊天”那一层

把员工用 AI 的能力，分成十级，从 LV1 到 LV10。

AI 使用水平 10 级

Lv	称号	能力标准
0	旁观者	知道 AI 但从没用过 · 全球 80% 在这
1	尝鲜者	帮我总结、帮我写邮件，你不会追问，不会补充背景，一般只用豆包或 deepseek，把 AI 当高级搜索引擎
2	对话者	会追问 + 补背景，知道提示词，开始意识到怎么问非常重要，工作中两三个场景应用，如写周报
3	驯化师	主动立规矩 + 调 Prompt，开始给上下文、拆分任务，AI 产生结果从随机变成大部分可用 超 70%
4	越境者	用 AI 做专业外的事，做过海报、做数据分析，为 GPT、Claude 付费，发现能力边界大幅扩张 超 90%
5	织网者	AI 嵌入 workflow、有固定用法、开始学会搭建 workflow 和智能体、开始设计自己和 AI 之间的协作方式
6	召唤师	从 ChatBot 跨到 Agent，开始尝试 ClaudeCode、skills，用 AI 做出自己的产品 超 97%，体验差最震撼
7	铸造师	设计 Agent，自己造 Skill、AI 反馈循环、复利曲线启动
8	造物主	开始研究 AI coding、CC、Codex，搭建自己工作台 超 99%
9	觉醒者	AI 成为思维方式的一部分，遇到问题第一反应怎么和 AI 一起做好 0.01%
10	一人军团	产出能力一个人对标几十人团队 · 系统能力 >> 个人能力 顶级 0.001%

来源 · 数字生命卡兹克

LV1-3：只会聊天。打开一个对话框，像问搜索引擎一样问 AI 一句，得到一段话，复制出来用。问什么、怎么问、得到的东西准不准，全凭感觉。这是绝大多数人的水平。

LV5-7：会把 AI 嵌进自己的活里。知道哪个环节该用 AI、怎么把一件复杂的事拆给 AI 一步步做、怎么搭一个能反复用的小 workflow、怎么判断 AI 给的东西能不能信。AI 成了他做事的一部分，不是偶尔玩一下的玩具。

LV8-10：能造工具给别人用。能把一类反复发生的工作，做成一个别人也能用的 Agent 或软件。这是少数人。

这把尺子上，你公司里绝大多数员工，都卡在 LV1-3。不管他嘴上说自己用过 ChatGPT，用过豆包，真上手看，基本都是只会聊天那一层。

这里我想插一个判断，它能帮你看清培训到底该图什么：AI 真正擅长的，是把下限抬上来，重点不在顶高上限。一个本来只能做到 50 分的普通员工，AI 能稳定地帮他做到 60、70 分；但从 85 分到 95 分那一截，靠的还是人自己的判断力和认知深度，AI 帮不上太多。这对你判断培训的价值太关键了：培训的目标，是把那批卡在 LV1-3、产出参差不齐的普通员工，整体抬到 LV5-7 的及格线以上，也就是把下限拉齐。而一个组织的产出，往往不是被那几个最强的人决定的，是被这条下限决定的。你把最普通那批人的下限抬上来，整个组织的效率水位就涨了一截，这件事的价值，恰恰被很多只盯着打造几个 AI 高手的老板低估了。

这不是你公司的问题，这是普遍现象。麦肯锡 2025 年的调研里有个数字很能说明问题：高管严重低估了 AI 在一线的渗透，高管以为员工的 AI 使用率只有 4%，但实际员工的使用率约 13%。注意，这个数字不是说员工用得更多，恰恰相反，13% 这个绝对值已经很低了，而且这里说的用，大部分就是 LV1-3 那种打开对话框聊两句。

而你请讲师做的那种“听课式培训”，绝大多数最多把员工从 LV1 推到 LV3。讲师讲了概念，演示了 demo，学员“看到了”“知道了”，但没“做到”。看到和知道，留不住。过两周，LV3 又掉回 LV1。

我把员工学 AI 这件事拆成五层，你一看就明白培训卡在哪：

看到：看了 demo，知道 AI 很强。

知道：了解基础概念、工具地图、能力边界。

做到：自己上手，跑通一个基本流程。

探索：自己能找新场景，同事之间互相帮。

判断：知道 AI 什么时候该用、什么时候不该用、用到什么程度可信。

普通讲师能做到的，就是第 1、2 层，顶多带一点第 3 层。而真正能让员工用起来的，是第 4、5 层，让他自己会找场景、会判断。这两层，听课听不出来，得靠现场带练、课后陪跑、持续答疑一点点磨出来。

我希望你先停在这里，看清一件事：听完课没人用，很多员工没用起来，原因不在懒，而在于他们没有真正带过那道从“知道”到“做到”的坎。他回到工位，面对自己手头那堆具体的活，根本不知道哪一件该用 AI、该怎么用，培训里讲的都是通用案例，跟他的活对不上。于是他放弃了，回到老路。这道坎，才是 AI 培训真正的命门。下面这两家公司，出发点天差地别，但都得过这道坎。

两个不同的起点：一个怕被取代，一个纯提效

我先讲清楚某头部电商平台和某四大行之一为什么找我做培训。因为老板做培训的动机不同，培训该怎么设计就完全不同，这是你作为老板第一个要想清楚的事。

某头部电商平台：财务团队不进化，就会被业务部门接管

某头部电商平台找我，是给规模不小的财务团队做培训。表面上的动机和所有人一样：提效，让财务人员学会用 AI 做事。但跟 CFO 深聊之后，我看到底下还有两层动机，这两层才是真正驱动他们花钱的东西。

第一层，是大厂的社会责任感。这种体量的平台公司，管理层是真心希望员工别在 AI 浪潮里掉队。让员工学到真东西，有能力适应新时代，不至于哪天被淘汰了还不知道怎么回事。这是一种“我得对我的人负责”的心态。

第二层，是财务部自己的危机感，而这一层，才是最锋利的那把刀。财务的工作有个特点：高度标准化。报销、对账、核算、出报表，流程清清楚楚，规则明明白白。标准化意味着什么？意味着这些活恰恰是 AI 最容易切进去的。

那么问题来了：如果财务部自己不先把这套 AI 化，会发生什么？业务部门会先学会。业务部门一旦掌握了用 AI 处理这些标准化财务流程的能力，财务部存在的必要性就要被打个问号。

这个恐惧，比提效这个动机强十倍。它核心不是我想干得更快，核心是我不进化就会被取代。所以这家平台的财务团队学 AI，是一种主动进化，趁着还有时间，用 AI 把自己武装起来，既

提效，又把财务专业知识和 AI 结合，变成业务部门替代不了的角色。

我跟 CFO 沟通时，把这层话挑明了说。这也是很多老板请外部讲师的一个隐性目的：有些话老板自己不方便对员工说，得借讲师的嘴说出来。比如公司给了你这么多学习机会和工具，你还不珍惜，那就算公司不动你，你自己在这个岗位上也待不下去了。这话从老板嘴里说出来像威胁，从外部讲师嘴里讲出来，是行业趋势，是善意提醒。

某四大行之一：体制内，就是单纯想提效

某四大行之一完全是另一个起点。

覆盖多个部门和网点的培训。它是体制内，没有某头部电商平台那种不进化就被取代的市场化危机感，岗位相对稳定，不存在业务部门把某个部门干掉的故事。

所以这家银行的动机很纯粹：就是提效。让员工学会用 AI，把手头的活干得更快、更省力。没有那么多弯弯绕，就是想让这套先进工具在体制内的日常工作里真正跑起来。

这一单还有个特点：它走的是聊天式 AI 到执行式 AI 的升级。很多体制内单位接触 AI，停在问一句答一句的聊天层面。我这次带的核心，是把他们往让 AI 真正帮你执行任务那一层推，以 Claude Code 这类能真正动手做事的工具为核心，而不只是停在对话框里聊天。

具体的培训对象，主要是客户经理和一些基础岗位的工作人员。他们的核心诉求很实在：一是基于银行已有的各类文件，再二次加工、写出新的文件；二是每个月那些固定要做的事项，能不能自动化、或者至少快速化。这两个诉求，在我看来正好对应 Claude Code 的两个能力：前者靠搭一个知识库解决（让 AI 基于银行自己的规范和模板来写），后者靠做成一个 skill 解决（把每月重复的固定动作，沉淀成一个能反复调用的能力）。

某四大行之一这一单我连续做了几场，覆盖一批员工。时间更长的那场，原本对方只想要一天，但我跟他沟通课程大纲时，他觉得内容太满、压缩讲不透。我干脆打电话给他，把我打算怎么讲从头到尾过了一遍。他听完说，这套东西一天压不下来，需要补一段实操。后来加出来的那段时间加在哪？加在把工作任务拆开、评估哪些该用 AI 哪些不该用这件事上，他认同这一步，比多教几个工具更重要。

我讲这个细节，是想让你看到一个反差：很多 AI 培训是甲方压时长、压价格，讲师配合着把课往薄里讲；而真正有料的培训，常常是反过来的，甲方一开始想要的时长，根本装不下真正该讲的东西。区别就在于，你请的人是来讲一场热闹的，还是来让你的人真能上手做事的。

为什么要把这两家放一起讲

我把这两类企业放一起，是想让你看清一件事：同样是“做 AI 培训”四个字，背后的老板动机可能完全不同。一类是市场化大厂，带着社会责任感和“被取代”的危机感，员工学习有内在驱动力；一类是体制内单位，纯粹为了提效，需要外部把“用起来”的氛围带进去。

动机不同，培训的设计、话术、节奏都得不一样。给某头部电商平台，我会重点打“财务不进化就被业务取代”这根弦，因为这根弦一拨，学习的主动性自己就上来了；给某四大行之一，我会重点降低使用门槛、把工具实操做扎实，因为这里缺的不是危机感，是“原来这东西我也能用起来”的体感。

作为老板，你做培训之前，先问自己一句：我到底想要什么？是想让团队认知升级、跟上时代？还是想让某个具体部门马上提效、产出能用的工具？这两个目标对应完全不同的培训方案。想不清楚就上，大概率两头不靠，认知没拉齐，工具也没落地。

怎么做的：把人从“看到”推到“做到”

讲清楚了起点，再讲动作。培训要让人真用起来，核心就一件事：把员工从 LV1-3 推到 LV5-7。从“只会聊天”推到“会把 AI 嵌进自己的活里”。这件事，光靠站在台上讲是做不到的，我的做法是这几步。

第一步：先摸清楚学员现在在哪一层

我做任何一场培训，第一个动作不是备课，是搞清楚这帮人到底什么水平。这里有个默认假设我必须告诉你：默认大部分人都很差。不管公司给我的简介里写得多漂亮，真上手一看，大多还停留在比较初级的水平。判断方法也简单：

公司内部如果已经在用一些 AI 工具（比如内部搭的对话平台），说明用过，但多半不深，LV2-3。

如果完全没有任何 AI 工具，基本是零基础，LV1。

还有一种最麻烦的情况：老板自己很懂 AI（LV7-8），但员工 LV1-2，中间差着巨大的鸿沟。这种时候老板的期望和员工的现实严重错位，培训方案最难设计。

关于这个错位，我还要多说一层，因为它牵出一个更要命的问题。我见过不少老板，自己不怎么用 AI，觉得“我派几个骨干去上个课、或者请人来给员工培训一下”就行了。但行业里有个反复出现的现象：那些教企业做 AI、做智能体的课，老板派 CIO、CTO 或者几个中层去听，听完回到公司，几乎都推不动，因为 AI 转型一动就牵扯流程、分工、考核，这些事中层定不了。最后，往往是老板自己不得不再学一遍、亲自下场，事情才动得起来。这一点我在讲“一把手工程”那章已经展开过，这里你只要记住一句：培训能拉齐员工的认知，但拉不齐老板自己的认知。老板自己那一课，谁也替不了。

而且，参差不齐是常态。同一个班里，有人完全不会，有人已经能搭简单工作流。课程设计必须同时照顾最差的（不能让他跟不上）和最好的（不能让他觉得无聊）。我的办法是实操环节用分层任务，基础任务保证人人能完成，进阶任务让快的人有得挑战。

第二步：不讲通用案例，拆他们自己的场景

这是听课式培训和我做的培训，最大的区别。

普通讲师讲的是通用案例：你看，AI 能写文案、能做表格、能总结会议纪要。学员当场觉得

“哇好厉害”，回到工位一看，自己手头那堆活跟讲的案例对不上，不知道从哪下手，就放弃了。

我的做法是反过来：直接拆学员自己的工作场景。在那家电商平台，我拆的是财务的活，报销审核里哪一步能让 AI 做初筛、对账时怎么用 AI 比对、月度报表的文字说明怎么让 AI 起草。在那家银行，拆的是他们各部门的日常，材料怎么写、流程怎么走、哪些重复劳动能交给 AI。拆场景这一步，是把“AI 很强”翻译成“AI 能帮我干我这件具体的活”。只有翻译到这一层，员工才会真的动手。因为他看到的不再是别人的案例，是自己明天上班就要面对的那张表、那份报告。

第三步：培养火种，别指望一次性教会所有人

一个团队规模不小，你不可能靠一两场培训把每个人都拉到 LV5-7，这不现实。

正确的做法是培养火种用户。在培训里找出那些学得最快、最有热情的人，重点把他们往上推。这些人推到 LV5-7 之后，会变成团队内部的“种子”，他们在自己工位上用出了效果，身边同事看见了，自然会去问、去学。这种同事之间的自发传导，比任何讲师都管用。

为什么？因为同事用的是同一套系统、干的是同一类活。火种用户做出来的东西，旁边的人一看就懂、就能抄、就能改成自己的。这是讲师给不了的，讲师讲完就走了，火种用户天天坐在他旁边。

这套“火种带动”的逻辑，不是我一个人的土办法，它背后是一条很硬的人性规律。有个做企业 AI 转型的同行讲过一句话，我很认同：这世上 99% 的人，是因为看见才相信，不是因为相信才看见。你在台上把 AI 讲得天花乱坠，大部分员工是将信将疑的；但当他亲眼看见隔壁工位那个和自己干一样活的同事，因为用了 AI 早下班两小时、或者被老板当众表扬，他立刻就信了，而且是那种会马上动手去学的信。火种用户真正的价值，就是把“AI 有用”这件事，从“讲师嘴里的道理”，变成“身边同事眼前的事实”。

斯坦福那份报告里有个案例，把这套打法落得更具体：一家半导体公司推 AI，发现工程、IT 这些技术部门接受得快，但法务、HR 这些部门一直滞后。他们的解法很具体：少开大会、少靠通知硬压，在每个部门内部找一个 AI champion（也就是火种），靠这些火种在部门内部做点对点的带动。靠人带人，比靠制度压人有效得多，因为人信的，是身边那个活生生的同事，不只是一纸文件。所以你培养火种，别只培养一两个总部明星，要尽量让每个部门、每个工种里都有一颗种子。

第四步：课后陪跑，把“知道”变成“做到”

这是最关键、也是最被低估的一步。

培训当天再热闹，只要课后没有跟进，两周必然打回原形。所以我做的培训，课只是入口。课后会有一段陪跑期，答疑、解决具体问题、帮他们把第一个真正能用的工作流跑通。

这一步的意义，是把“知道”变成“做到”。一个员工只有亲手跑通了一个属于他自己的工作流、用它真省下了时间，这个习惯才算在他身上扎下根。在这之前，他学的所有东西都是悬空的，随时会掉。

陪跑也是过滤器。哪些人真的用起来了、哪些人学完就丢了，陪跑期一清二楚。真用起来的，就是你接下来要重点扶持的火种；学完就丢的，你也别浪费资源硬推，组织手段（考核、激励）比培训管用。

现场到底卡在哪、又是怎么跑通的。

把人从“看到”推到“做到”，听起来是一句话，真到现场，全是琐碎又具体的坎。我带这两场培训，学员卡住的地方高度一致。我挑几个讲，因为这些坎，恰恰是“听课式培训”永远碰不到、也解决不了的。

最开始，卡在工具本身。装环境反复报错、API 怎么配、key 填在哪里，这些对搞技术的人是常识，对财务、对客户经理是天书。这一关过不去，后面什么都谈不上。所以我的实操环节，第一件事就是手把手帮每个人把环境装好、API 配通，把这道物理门槛先清掉。

往后，卡在不会拆活。给 AI 一个笼统的大任务（“帮我把这个月的对账做了”），AI 当然做不好，学员就得出结论“这件事不行”。其实问题不在 AI，在他没把活拆开。我在现场反复带的就是这个：把一件复杂的事，拆成 AI 能一步步接住的小任务。

再往后，卡在“说不清自己要什么”。一个学员把需求描述得很简短、模糊，AI 理解出了偏差，给的结果自然不对，他就说“AI 给的东西不理想”。但根子是他没把话说清楚。这一点我得专门点出来，因为它正是这本书反复在讲的那件事，用不好 AI，很多时候问题往往在于你没能力把自己要什么说清楚。这个能力，恰恰是培训最该练、也最难练的。

Claude Code 这类工具做事时，会自己尝试很多种方案。一个本该一分钟解决的问题，它可能吭哧吭哧试上十分钟还没出结果。学员盯着屏幕，既不知道是卡死了还是在正常工作，也不知道自己该干嘛，其实正确的动作，是按下 **Ctrl+C** 把它停下来，让它回过头看看到底卡在哪、再重新来过。但没人教过他这个，他就只剩干着急：怎么这么久还不出来？这种细节，是真正下场用过、带过人才知道的，讲 PPT 的讲师讲不出来。

最容易让人放弃的，是 AI 犯错那一刻。AI 偶尔会改错东西，比如把一个文件名改错了，或者犯个别的小错。学员一看 AI 出错了，又不知道怎么撤回、怎么恢复，当场就给 AI 判了死刑：AI 会犯错，我这活不能用它。

这是整场培训我最想掰过来的一个认知。AI 会犯错，很多时候是因为你没把话说清楚；而人和 AI 的协作，本来就是一个磨合、逐步建立信任的过程。正确的姿势是先给它一些简单的、甚至专门用来测试的小任务，确认它在这类活上靠谱了，再一点点加复杂度。不要一上来就把最复杂、最重要的活丢给它，然后它一出错就全盘否定。就像带一个新人，你不会第一天就让他独立负责最关键的项目。这个磨合的耐心，比任何工具技巧都重要。少数脾气急的学员，AI 一次没满足就开始上火，说这东西一点用都没有，从头到尾不愿意多问几句、逐步逼近他要的结果；而那些最后真正用起来的人，恰恰是愿意跟 AI 来回磨、把它当个会成长的助手来调教的人。

至于“我年纪大了、学不会”，这种担心，在这类资历较深的员工偏多的地方，会更明显一些。但真教起来，这些工具的上手门槛比他们想象的低得多。而且这场培训是自上而下、组织推动的，大家有共识、有统一节奏，抵触情绪反而不大。这也印证了后面会讲的那条：让员工用起

来，组织的推动力，往往比工具本身更关键。

讲到这里，可以看一个具体画面。现场总有那么几个学得快的人，基础案例一跑通，立刻就不满足于练习题了，他们直接把自己手头的工作流程拿出来，当场就要让 AI 帮他跑。跑的过程中冒出各种实际问题，我一个一个现场给他解决、给他讲清楚背后是怎么回事。这种现场，普通 AI 讲师是接不住的。因为我自己在大厂做过很多年企业级的工程落地、亲手交付过给企业用的产品，所以学员一说他的业务、他卡在哪，我们之间是有共同语言的，我能真正听懂他那行活是怎么回事，而不只是教他几句通用的提示词技巧。这一点，是我和大多数只会讲工具的讲师最不一样的地方，也是这些人最后能真正从“只会聊天”跨到“能用 AI 做事”的关键。

拿到什么结果

先给你一组我在多个培训项目里反复验证过的数字，这是这套打法的稳定产出：

一个月后，不少的日常工作能用 AI 提效。经过培训加陪跑，大部分被推到 LV5-7 的员工，能把自己手头不少的日常活接上 AI。

两到三个月后，效率提升明显提升。重点在于整体工作节奏起来了，不只是一件事变快，重复的活交给 AI，人腾出手干更值钱的事。

具体到这两家：某头部电商平台这边，培训现场的一批人，基本都跑通了，装好 Claude Code、配好各种 API，然后当场动手把自己的工作拆解成一批任务。财务的活前面说过，高度标准化，这一拆就看得很清楚：大部分人手里很多活，都能拆成清晰的 SOP 交给 AI 协助。拆得清，提效就是自然结果。而现场学得快、用得活的那批人，就是后面财务部继续往 AI 化走的火种。

某四大行之一这边，落地的场景就是前面说的那两类：一类是基于银行文件快速二次写文档，客户经理还把跟客户沟通的内容快速做成总结、做项目推进记录；另一类，是每个月固定时间要做的那些事项，从手动变成自动或半自动。对体制内来说，这种“日常重复劳动被接管”的体感，就是最实在的提效。

但我想让你看的，不只是这两个数字。真正的价值在数字背后，有两层。

第一层，员工的能力真的迁移了。很多员工听完 AI 培训之后，没办法把课上学到的东西关联到自己的实际岗位。也就是说，学是学了，跟自己的活接不上，很快就放下了。而我这套“拆场景 + 陪跑”的打法，核心解决的就是这个“接不上”的问题，让学的每一样东西，都长在他自己的活上。

第二层，这类头部平台的打法，撬动的是一个圈层。这家平台的 CFO 在财务这个圈子里认识非常多同行的 CFO。一场培训做好了，做出效果、攒下能传播的案例，这个口碑会在 CFO 的社群里自己跑起来。老板和老板之间是有社群、有联系的，一家做了见效，其他家会跟进。这是 B 端培训一个很实在的逻辑：做一家标杆客户，等于敲开了它背后整个圈层的门。

老板该学到什么

如果你是企业老板，从这个对照案例里，我希望你带走三条判断，它们能用到你公司任何一次 AI 培训上。

第一，培训不是终点，是入口。

这是最重要的一条，你花钱做的那场培训，本身解决的只是“认知”和“基础使用”，把员工的水平从 LV1 拉到 LV3 的及格线上。它解决不了三件事：场景诊断、流程改造、系统落地。这三件才是真正让 AI 在你公司产生价值的部分，而它们都在培训之外。

所以，别把“做完培训”当成“完成了 AI 转型”。这是我见过老板最常犯的认知错误之一。培训只是把人领进门，真正的活在门里头。把培训当终点，你会发现钱花了、人散了、什么都没留下。

这里有个反差，你必须心里有数：让一个员工用 AI 提效 50%，不难；但让整个组织提效，极难。为什么？因为个人提效，只要这个人自己愿意学、愿意用就行；而组织提效，要动流程、动分工、动考核、动上下游的协同。举个最常见的例子：一个客服用 AI 把工单处理快了三倍，听着很好，可如果他的考核还是“每天处理多少单”，那他快了之后只是更早闲下来，组织整体的产出一点没多，甚至老板还会开始盘算“是不是可以裁人了”。你看，个人那一端的效率是实打实涨了，但它没有自动变成组织的效率，中间卡着考核、流程这些更硬的东西。所以你千万别把“培训完员工都会用了”当成“组织转型成功了”，前者是个人的事，后者是组织的事，这本书后面好几章（尤其讲组织那一篇），都在帮你把这中间的距离补上。

第二，先想清楚你要的是“认知”还是“提效”，别两个都想要。

做培训之前先问自己：我是想让团队认知升级、跟上时代（像某四大行之一那样的提效铺底、像大厂的责任感），还是想让某个部门马上产出能用的工具（像某头部电商平台财务团队那样的硬提效）？

这两个目标对应完全不同的方案：要认知，半天到一天的课程冲击一下就够；要提效，必须是带实操、带场景拆解、带课后陪跑的深度培训，而且得明确知道这件事不便宜。想不清楚目标就上，两头不靠。你以为你买的是提效，讲师给你的是认知冲击，散场时大家很嗨，但手里什么都没多。

第三，培训能解决认知，解决不了主动性和效率。我跟所有要做培训的老板都会画一张图，这里也给你：

认知转变，培训能解决。靠认知冲击、案例、demo。

能力提升，培训能解决一部分。靠实操和方法论，剩下靠员工自学。

主动性，培训解决不了。靠组织机制：考核、激励、淘汰。

效率提升，培训不能直接解决。它是能力提升加上主动性之后，自然出来的结果。

理解上面，你对培训的期望就摆正了：培训负责把“不会”变成“会”，负责把“不知道”变成“知道”；但“会了愿不愿意天天用”，那是组织的事，不能只靠讲师解决。一个员工学会了却不用，

问题往往不在培训，在你公司有没有一套让他不得不用、用了有好处的机制。这是后面讲“怎么让员工持续用起来”那一章会专门展开的。

一句话

这章收束成一句话：**培训能解决“认知”，解决不了“落地”**。听完课只是站到了起跑线，真正的活，拆场景、改流程、搭系统、养习惯，全在课后。把员工从 LV1-3 的“只会聊天”，推到 LV5-7 的“会把 AI 嵌进活里”，这才是培训该有的目标。而要送他们过这道坎，光讲课远远不够，得拆他们自己的场景、培养火种、课后陪跑。这些动作，决定了你那笔培训费，到底是买了一场热闹，还是买了一支真能用 AI 做事的队伍。

第 6 章 知识库为什么难点不是模型

一个被所有人想简单了的项目

先把一个很多老板不爱听的判断放在前面：搭知识库这件事，最难的部分，跟模型一点关系都没有。

我们做过一个项目，是给一家制造企业做 ISO 审核智能体，这家公司是行业里做精密制造的大厂，我们服务的是它一家区域子公司。用这个项目来讲清楚一件事，它戳破了企业里关于知识库最大的一个误解。

这个误解是：很多老板，包括很多技术团队，一想到做知识库，脑子里的画面是，把公司的文档一股脑传上去，接个大模型，AI 就能回答问题了。听起来简单，买个工具、传一批文件、配一下，搞定。

然后真做了，结果是 AI 答得驴唇不对马嘴。问 A 答 B，要么答非所问，要么一本正经地胡说，要么干脆说“我没找到相关内容”，可那内容明明在文档里。老板一脸困惑：文件我都传进去了，怎么还是不行？是不是模型不行？要不换个更贵的模型？

换模型没用。因为问题压根不在模型。

做知识库，模型只占很小一块，真正的难、真正的脏活累活、真正决定成败的，是把你那一堆乱七八糟的历史数据，洗成 AI 能用的知识。这件事行话叫数据清洗，说人话就是：你那些文档，现在根本不是给 AI 准备的，你得先把它们改造成 AI 读得懂的样子。

这一步做好了，后面接模型、做问答，反而是相对轻松的部分。这一步做不好，后面再贵的模型都白搭。在讲怎么做之前，得先了解“原来什么样”，也就是这家公司的数据，到底烂成什么样子。你不亲眼看一眼那堆烂数据，你就理解不了为什么“传文档”这条路必死。

原来什么样：数周填一张表，数据烂成一锅粥

这家公司它是做精密制造的，要给下游的大客户供货。不同客户进来前都要做一次 ISO 审核，简单说，就是客户拿一张很长的审核表过来，上面几百个问题，挨个问你：你的质量管理体系是怎么做的？这道工序的标准是什么？出了不良品怎么处理？相关记录在哪？

这个部门要做的，是把这张审核表填完，每一个问题都需要生产线上下游十几个部门，从自己公司海量的质量管理文档、流程记录、历史档案里，找到对应的依据，填进去。

这件事原来要花两三周时间，因为上下游十几个部门，每一个部门的员工还不一定熟悉，还需要去资料堆里翻，还有可能填错，还需要来回修改，一份份翻、一点点找、一条条对。慢，而

且累，而且换个新人根本干不了。

这听起来，是不是一个适合上 AI 的场景？知识密集、重复发生、有大量文档可查，简直是为知识库量身定做的。在他们给我知识库之前，我也是这么判断的。当他们把数据样例给我时，我看到的是另一幅画面。

坦白说，是一锅粥。**脏、乱、差，而且高度非结构化。**什么叫非结构化？我给你具体看几样它资料里存在的东西：

扫描的图片。很多关键记录、签字盖章的表单，是纸质文件扫描成的图片。图片里是字，但对 AI 来说，那就是一张图，不是文字，它读不出里面写了什么。

五花八门的 PDF。一堆 PDF，有的是文字版，有的是扫描版（本质又是图片），格式各不相同。

复杂的表格。大量信息塞在 Excel 或 Word 的表格里：合并单元格、跨页、嵌套；表格里还内嵌着图片、写满了公式；有的一个文档因为做过表格合并，塞了好几张表挤在一起。人看着都费劲，AI 更搞不清哪个值对应哪个项。

流程图和思维导图。质量流程、工艺流程画成了流程图，有的知识点整理成了思维导图，框、箭头、分支、节点。这些图里藏着关键的逻辑，但对 AI 来说，又是一张看不懂的图。

光是格式上的乱，就已经够呛了。但更麻烦的还在后头，**内容本身也是乱的：**

有的数据已经过时了。流程改过、标准换过，但旧文件还躺在资料堆里，新旧混在一起，你不深入了解，根本分不清哪份是现行有效的。

压缩包套压缩包，有的还带密码。资料层层打包，有些压缩包加了密，我们还得专门找他们要解压密码，才能打开看里面是什么。

逻辑互相冲突。不同文件对同一件事的说法对不上，到底以哪个为准，得人去判断、去跟他们确认。

满是专有名词。一部分是公司内部的叫法，一部分是这个行业的术语，而且很多术语，在别的行业里是完全不同的意思。你不懂这一行，就会理解错、切错、归错类。

你把这些东西原封不动丢给大模型，会发生什么？它要么读不到（图片、扫描件里的字它抓不出来），要么读错（表格的对应关系、相互冲突的说法它理不清），要么读了个寂寞（流程图、专有名词它看不懂）。然后你问它问题，它当然答得驴唇不对马嘴，不是它笨，是你喂给它的根本就是它消化不了的东西。

我希望你在这里停一下，看清一件事：这家公司不是没有知识，它的知识多得很，但那些知识是“写给人看的”，不是“给 AI 准备的”。一个老员工凭经验能在这堆乱资料里翻出答案，因为他脑子里有一张地图，知道什么东西大概在哪、那张模糊的扫描件上写的是是什么。但 AI 没有这张地图，也没有这双能“看懂”图片和流程图的眼睛。

这就是企业知识库最普遍、最致命的真相：你公司里所有的文档，几乎都是为人组织的，不是为机器组织的。而做知识库，本质上就是把“给人看的”重新整理成“给 AI 用的”。这件事，跟模型聪不聪明没有半点关系，它是个又脏又累的整理活。

我们怎么做的：大部分功夫花在洗数据上

很多人做知识库项目，一上来就问：用什么向量数据库？选哪个嵌入模型？RAG 怎么搭？prompt 怎么写？这又是把顺序搞反了。我们做这个 ISO 审核项目，精力分配是这样的：绝大部分功夫，我估计占了大头，花在了数据清洗上。真正搭检索、写 prompt、做自动回填的部分，反而是相对省力的后半段。你看看下面步骤你就明白这部分功夫花在哪了。

第一步：把“看不懂”变成“读得出”

第一关，是把那堆 AI 读不了的东西，变成 AI 能读的文字。

扫描的图片、扫描版 PDF，要先把里面的文字识别出来，变成真正的文本。

复杂的表格，要拆开、理顺，把“这一行这一列对应的是什么值”重新组织成 AI 能对上号的结构。

流程图里的逻辑，什么情况走哪条分支、出了问题流转到哪，要把它从图里“翻译”成文字描述。这一步很笨，很多地方机器自动识别的结果不准，还得人去核、去补、去改。但它是地基。这一步不做，后面全是空中楼阁。AI 连你的文档“读”都读不全，谈何回答问题。

第二步：把“一锅粥”拆成“一条条能查的知识”

文字提出来了，还远远不够。一大段一大段连在一起的文字，AI 检索的时候照样抓不准。第二步，是把这些文字按合理的颗粒度切开，整理成一条一条边界清晰的知识。一个质量标准是一条，一个流程节点是一条，一个不良品处理规则是一条。每一条都要清清楚楚、能被单独检索到、能对应回它原本出自哪份文件。这一步考验的根本不是技术，是对这家公司业务的理解，你得知道审核表会怎么问、一个问题需要哪几条知识拼起来才能答全，才知道该按什么逻辑去切、去归类。切错了，AI 检索的时候就拼不出完整答案。

第三步：这时候，才轮到模型出场

把脏数据洗成了一条条干净、结构化、边界清晰的知识之后，这时候，RAG、prompt、模型，才开始登场。而且你会发现，到了这一步，事情突然变得不难了。

搭检索（RAG）：知识已经切干净、标清楚了，让 AI 根据审核表的问题，去这些干净知识里检索匹配的内容，这一步是相对标准化的工作。

写 prompt：告诉 AI 该怎么理解审核问题、检索到内容后怎么组织成审核表要的答案格式，有了干净知识打底，prompt 也好写。

自动回填：让 AI 把检索、组织好的答案，自动填回审核表对应的位置，临门一脚，水到渠成。

我们还用了多个 Agent 分工协作来干这件事，有的负责理解审核问题、有的负责检索知识、有

的负责组织答案回填。但这些“高科技”的部分，全都建立在前面那些脏活的基础上。地基打好了，盖楼是快的；地基没打好，楼盖得越高塌得越快。

这就是我想让你记住的那个比例：一个知识库项目，大部分功夫在洗数据，剩下那部分才是模型和技术。而老板的注意力，常常 100% 放在模型和技术上，对那些脏活视而不见，这正是知识库项目最常翻车的地方。

拿到什么结果

这套智能体做出来之后，效果很明显：原来要花两三周才能填完的 ISO 审核表，智能体能自动检索、自动回填，把人从那堆翻找的苦活里解放出来，AI 能填写 80%+，准确率也能达到 95%，最终缩短到两三天能完成。

人力上的变化更直接，原来这张表要生产线上下多个部门、每个部门派一个人各自填自己那块；现在不用这么兴师动众了，质量部门让 AI 先把大部分的内容自动填好，剩下的再交给上下游相关部门快速确认一下就行。从十几个部门各填一摊、耗时数周，变成一个部门主导、较短时间搞定。

但结果背后更重要的是它的逻辑，因为它对你判断自己公司的知识库项目，价值远大于一个数字。

真正被解决的，不是“填表慢”这个表面问题，而是“知识散在一堆烂数据里、只有老员工捞得出来”这个底层问题。数据洗干净、结构化之后，这家公司多年积累的质量管理知识，第一次变成了一个机器能用、能查、能复用的资产。老员工脑子里那张“东西在哪”的地图，被显性化、固化下来了。这才是这个项目最深的价值，它顺带解决了填表慢，但它本质上是把一堆死的、散的历史数据，盘活成了活的知识资产。

在这儿我想多说一层，因为它能彻底改变你对“洗数据”这件苦差事的看法。很多老板把数据清洗看成纯粹的成本，又费钱又费时，还不出彩，恨不得让供应商赶紧跳过。但换个角度想：你洗出来的这套干净知识库，恰恰是别人抄不走的东西。模型是公开的，谁都能买；但你这家公司几十年攒下的质量管理经验，被一条条洗干净、结构化好，是你独有的。斯坦福那份报告专门统计过：在他们研究的成功项目里，75% 都把“专有数据”列为 AI 战略的关键，47% 直接把自己积累的数据称为“护城河”。报告里有句话我很认同，开源模型迟早会把性能差距追平，到那时候，真正拉开差距的，不再是你用哪个模型，而是你喂给它什么数据。所以洗数据这件脏活，它是在帮你筑一道随时间越来越深的护城河。你越早开始洗、洗得越干净，这道护城河就越宽。

对于上面项目，光是清洗这家公司的知识库，我们就花了一段不短的时间。要知道，这家公司的质量管理文档，前前后后堆了很多。而且这一段时间，绝不只是埋头处理文件，为了把前面说的那些过时数据、互相冲突的说法、一堆行业专有名词搞清楚，我们还得一边深入研究他们这个行业，一边反复跟他们公司懂行的人沟通确认。

洗数据从来不是一个纯技术、纯体力的活，它是个需要懂业务、需要不断跟人对齐的判断活。

哪份文件是现行有效的、这个术语在他们这行到底指什么、两个冲突的说法以哪个为准，这些 AI 自己判断不了，只能靠人去把它一条条理清楚。这也正是为什么我们反复强调：知识库的成败，根本不在模型那一端，而在有没有人愿意、也有能力把这堆脏数据啃下来。

老板该学到什么

如果你是企业老板，从这个案例里，我希望你带走三条判断，它们能帮你看穿任何一个知识库项目靠不靠谱。

第一，知识库的难点不是模型，是数据。

你的直觉会告诉你：AI 答不好，是不是模型不够强？换个更厉害的、更贵的模型试试？九成情况下，换模型没用。因为问题出在你喂给它的数据，那些数据又脏又乱又看不懂，再强的模型也消化不了一锅粥。

这件事不是我一个人的判断。全球做 AI 落地的机构得出的结论高度一致：真正卡住企业 AI 的，绝大多数不是技术。有个广为引用的拆解叫 10-20-70，一个 AI 项目里，算法模型只占 10% 的功夫，技术和数据搭建占 20%，而剩下 70% 全是人和流程、数据治理这些“脏活”。知识库项目尤其如此。你以为你在买一个聪明的模型，实际上你真正要付出的，是把自己家那堆历史数据洗干净的代价。

斯坦福那份报告里有两个数字，一个：在他们研究的成功项目里，数据“完全就绪”的只有 6%，剩下 94%，数据都有这样那样的毛病，照样做成了，靠的就是把洗数据当成项目的一部分认真干，而不是指望模型替你扛。另一个是：42% 的项目里，用哪个模型完全可以互换，换个模型，结果差不多，因为决定成败的根本不是模型。持久的优势在编排层，而不在基础模型。翻译成人话就是，你把数据洗干净、把知识组织好、把流程接顺，这些模型之外的功夫，才是真正值钱、别人学不走的地方。所以下次再有人跟你说“换个更强的模型就好了”，你心里要清楚：九成是数据的问题，不是模型的问题。

第二，“上传文档”不等于“建知识库”。

把文件传进系统，这只是第一个动作里最不值钱的那一下。文档不等于知识，上传不等于治理。你那些 Word、PDF、扫描件、表格、流程图，是给人看的原始材料，不是 AI 能直接用的知识。中间隔着一整套又脏又累的活：识别、提取、拆分、结构化、标边界、对应来源。

所以，当有供应商跟你说“你把文档给我们传上来，我们接个大模型就能用了”，你要立刻警惕。这句话基本等于告诉你：他没打算干那些脏活，或者他根本不知道这些脏活的存在。这种项目，做出来大概率就是答得驴唇不对马嘴的样子。

第三，知识库本质是个管理问题，不是技术问题。

把数据洗干净、把知识按业务逻辑切好、还要定清楚以后谁来持续维护它（资料会变、流程会改，知识库不养就会过期），这些全是管理动作，不是写代码能解决的。它要求一个既懂你这行业务、又懂 AI 边界的人，坐下来一条一条地把你的知识抠出来、理顺、固化。

所以做知识库，别只盯着技术团队问“模型选好了没”“RAG 搭好了没”。你更应该问的是：谁来负

责把我们这堆脏数据洗干净？谁懂业务、能判断知识该怎么切？上线之后谁来养它？这几个问题答不上来，技术再好，知识库也是个会越用越歪、最后没人信的摆设。

一句话

这章收束成一句话：**做知识库，难的不是模型，是把你那堆脏数据洗成 AI 能用的知识。**模型只占很小一块，数据治理才是大头。老板的注意力如果全放在“选哪个模型”上，这个知识库项目大概率要翻车。真正决定成败的，是有没有人愿意、也有能力，把你那堆给人看的、又脏又乱的历史文档，一条一条洗成机器读得懂、查得到、用得上的干净知识。

第 7 章 从岗位经验到 AI workflows

一个被重复电话逼出来的项目

这是我在某互联网大厂做支付风控时，一个客诉处理项目的结果。一件原来要消耗几小时的事，最后压到分钟级。

我想用这个项目，跟你讲清楚一件事，AI 真正落地的入口，往往要先把一个岗位上反复发生的重复劳动，变成一条能跑的工作流，而不是一上来就做能自主决策的多智能体大系统。

这件事听起来不性感，比起“我们搭了一套自主决策的 AI 大脑”，说“我们把客服的重复问题自动化了”显得太小、太朴素。但我做了这么多企业 AI 项目，越来越确信一件事：能落地、能算账、能让老板真看到结果的，几乎都是从这种“不性感”的小切口开始的。那些一上来就喊着要做大系统、要全自动的，大半死在半路上，这一点，后面讲“坑”那一章我会用一个失败案例给你看。

在讲怎么做之前，得先讲清楚原来是什么样。因为如果你感受不到原来有多痛，你就感受不到这种效率变化的分量，也就抓不住这一章真正想给你的判断。

原来什么样：一个问题要穿过三道人墙

某互联网大厂支付业务的体量，意味着每天有海量用户在用。用户用支付时遇到问题，交易突然被拦了、账户异常、提现失败，第一反应，就是找客服。

我给你还原一个最普通的场景。一个用户，大额转账被风控拦了，急得不行，打电话给客服：我这笔正经交易，为什么不让我付？客服一看，系统显示“风险拦截”，但具体为什么拦、该不该放，客服自己也懵，因为支付风控这块的问题有个特点：相当一部分，客服根本判断不了。

一笔交易为什么被拦，背后可能是几十条风控策略里的某一条命中了，牵扯到用户画像、设备指纹、行为序列、历史关联，这些东西客服看不懂，也无权查。客服能看到的，只有一个冷冰冰的结果：这笔，被拦了。至于为什么、该不该解、怎么跟急眼的用户解释，他答不上来。

客服解决不了，怎么办？走流程，找风险分析师。这个流程是这样的：客服在 ERP 系统里提交一张工单 → 工单流转风控团队 → 风险分析师接单 → 登录后台、拉数据、查命中的是哪条策略 → 判断这笔到底是误拦还是真风险 → 把结论写回工单 → 客服收到结论 → 客服再回头答复用户。

一个问题，要穿过“用户 → 客服 → ERP 工单 → 分析师 → 回流 → 客服 → 用户”这么长一条链。用户在那头干等着，中间隔着三道人墙。一来一回，几个小时就过去了，运气不好，跨了个班次，用户第二天才等到一个答复。

慢，只是你最先看到的那层问题。真正要命的，是底下两点。

第一，分析师被重复问题淹没。

我把一段时间的客诉工单拉出来分析，发现一个问题：分析师每天处理的问题里，有一大半是重复的。同一类拦截、同一种异常，换一个用户，又来一遍；换一天，还是这些。来来回回，就那么几类问题在循环。

风险分析师是风控团队里最贵、最稀缺的人。他们本该去干什么？去研究新冒出来的欺诈手法、去和黑产正面对抗、去设计能拦住下一波攻击的策略，这些才是只有他们能干、且真正创造价值的事。结果呢？大量工时，被消耗在回答“这笔为什么被拦”这种已经回答过几百遍的问题上。

这是一种极其昂贵的浪费。你花高薪请来的专家，有一半工时，在做一件实习生水平、但实习生又干不了的事。它干不了，是因为它需要专业判断；可它又简单重复，根本不配占用一个专家的时间。这个又简单又只有专家能做的尴尬地带，就是 AI 最该切进去的地方。

第二，经验全锁在分析师的脑子里。

怎么判断一笔交易是误拦？老分析师扫一眼就知道。但你让他说清楚“你到底是怎么知道的”，他说不清楚，这是干了几十年练出来的手感，是一种“只可意会”的判断力。

这就带来一个组织上的脆弱点：整个团队的判断力，押在那几个老分析师身上。新人接手，要么判断不准，要么不敢拍板，只能再回头去问老人。老人一忙、一请假、一离职，这块能力就出现一个空洞，而且这个空洞，平时被几个老人硬扛着，你根本看不见，等到他们真走了，你才发现整块业务突然没人能做主了。

慢、重复、经验锁死，这三个问题都不是“AI 不够聪明”造成的。它们是业务流程太长、加上经验没有被沉淀下来造成的。这一点，你在这本书里会反复撞见。企业 AI 做不起来，绝大多数时候卡的不是模型，是这些“AI 之外”的东西。斯坦福那份跟踪 51 个成功项目的报告，把这件事说得很直白：他们回头看，77% 最难啃的挑战都在技术之外，流程、数据、组织。技术，反而是最容易的那部分。

我们怎么做的：第一步根本没碰 AI



很多人一上来就问：用什么模型？搭几个智能体？接不接知识库？这是把顺序彻底搞反了。我们在做这个项目时，第一步，是把分析师判断一个客诉的过程，拆成一条人能看懂的工作流。

把“手感”翻译成“能写下来的逻辑”

我做的事很笨：找团队里最资深的分析师，搬个凳子坐他旁边，一笔一笔地过他处理过的真实客诉。他每判断一笔，我就在旁边追问，

“你第一眼看的是哪个字段？”“看到这个值，你为什么立刻觉得是误拦？”“什么情况下你会判定它是真风险、维持拦截？”“什么情况下你拿不准、必须自己深挖、不敢直接给结论？”

一开始他答得很笼统：凭经验啊，看多了就知道了。但你别接受“凭经验”这三个字，那是一堵墙，墙后面才是真东西。追问得足够细、足够多，那层“凭经验”就开始裂开，露出底下真正的判断结构：原来他第一眼看的是某几个关键字段的组合，原来某个特征一出现他就基本能定性，原来某类情况他从来不自己拍板、一定上报。

把这些答案一条条攒起来、归类，一个原来“只可意会”的判断过程，慢慢变成了一套“可以写下来”的规则：

遇到 A 类拦截、且满足 X 条件 → 结论是误拦，可放行，这样答复用户；

遇到 B 类、命中 Y 特征 → 是真风险，维持拦截，这样答复用户；

遇到 C 类、或关键信息不足 → 智能体不下结论，直接转人工。

这一步，把老师傅脑子里的手感，翻译成能写下来的判断逻辑，才是整个项目最难、也最值钱的部分。

AI 不会凭空懂你的业务，它不知道你的业务，不知道你公司那些没写在任何文档里、只活在老

员工脑子里的判断标准。你得先有人把这些隐性经验挖出来、显性化，AI 才有东西可学。这件事做完，后面才轮到 AI 出场。

跳过这一步直接上 AI 的企业，我见过太多。结果都一样：模型很聪明，答得很流利，但答得不对，因为它压根不知道你这行“对”的标准是什么。它不知道你要的是什么，这就是为什么很多公司买了最贵的模型，出来的东西还是不能用：你没喂给它判断的标准，它只能给你一个看起来很专业、实际靠不住的答案。

有了逻辑，才搭智能体

把判断逻辑和相关知识沉淀成一个知识库之后，我们才开始搭那个客服能直接用的问答智能体。

这里的活，大致是三块。

一是知识库：把那套判断规则、常见问题的标准答复、各类拦截的解释口径，整理成 AI 能检索、能对上号的结构化知识，这块的脏活累活，和我在第 6 章讲的那个制造企业洗数据是一个道理，不展开，你记住它不轻松就行。

二是检索和应答：让智能体接到客服的问题后，能准确地从知识库中找到对应的判断逻辑，再组织成一段客服能直接拿去答复用户的话。

三是转人工的口子：在智能体判断“这个我拿不准”的时候，能干净利落地把问题交还给人，而不是硬答。

逻辑跑起来之后是这样的：客服遇到自己搞不定的支付风控问题，不用再提工单、不用再排队等分析师，直接问这个智能体。智能体基于沉淀下来的判断逻辑和知识库，当场给客服一个能用来答复用户的结论。

原来那条穿过三道人墙、动辄几小时的长链，被压成了一步：客服直接问、当场得到答案。

上线不是终点：得有人“养”它

还有一步，很多企业会忽略，但它决定了这套东西能不能长期活下去：上线之后，得有人持续地养它。

智能体刚上线时，不可能一步到位。总有它答错的、答得不够好的、或者拿不准却硬答了的情况。关键在于，有没有一套机制，把这些错误收集起来，哪些问题它答砸了、为什么砸、是知识库缺了内容还是判断逻辑没覆盖到，然后一轮一轮地补进去、调出来。

风控这行还有个特殊性：黑产的手法每周都在变，旧的判断逻辑会过期，新的拦截类型会冒出来。所以这套智能体不是“做完就不管了”，它得跟着业务一起持续更新。没有这个养的机制，再好的智能体，几个月后也会慢慢变得不准、然后没人信、最后没人用。这一点，斯坦福那份报告也印证了：他们研究的成功项目，100% 都用的是迭代的方式，边跑边改，而不是憋一个完美版本一次性上线。那些把 AI 项目当成一锤子买卖、上线就撒手的，基本都没成。

我们没追求 100%

这里有个关键设计，你作为老板一定要记在心里，因为它会决定你公司几乎所有 AI 项目的生死：我们没有追求让 AI 解决 100% 的问题。

智能体的目标：解决大部分高频重复的问题。少部分真正复杂、信息不足、需要深度研判的，智能体不硬答，直接转人工，转回给风险分析师。

为什么不追求 100%？

因为最后那部分，恰恰是最容易判错、风险最高的部分。让 AI 硬扛这部分，一旦判错，把一笔真风险当成误拦放行了，损失是真金白银，而且可能是大钱。把这部分留给人，你既守住了风险底线，又把分析师从大量重复劳动里彻底解放出来，让他们专心去干真正需要人脑的复杂问题。

这个“大部分自动处理、少部分交还给人”的设计，在斯坦福那份报告，专门研究了不同程度的人工参与跟产出的关系，得出一个很有说服力的结论：那种“AI 自主处理 80% 以上、人只复核例外”的模式（他们叫“升级模式”），带来的中位效率提升是 71%；而那种“每一步都要人点头同意”的模式（他们叫“审批模式”），只有 30%。在合适的场景里，给 AI 多一点自主权、人只守住例外和高风险那部分，效果反而显著更好。我那个 80/20，踩中的正是这个最优区间。

但反过来也要清醒：报告同样强调，不是监督越少越好。像医疗、金融这种错一次就出大事的场景，人工审核是法规和风险逼出来的硬要求，该让人逐个把关就得逐个把关。关键不是“要不要人”，是“在你这个具体场景，人该守在哪个位置”。我们做的是风控，容错空间相对大、错了能补救，所以敢让 AI 承担大部分高频任务；你的场景如果错一次代价极高，人工兜底比例就必须更高。

我见过太多企业 AI 项目栽在“贪 100%”上。老板觉得既然上了 AI，就该全自动、全替代，留着人显得不够先进。结果为了那最后一小部分复杂问题反复折腾，系统越做越复杂，越复杂越不稳，最后整个项目拖垮。在企业里，一个能稳定解决大部分高频问题、并且老老实实把复杂问题交还给人的 AI，远比一个号称什么都能干、但你不敢完全信它的 AI 有价值。

拿到什么结果

这个项目上线后：

处理时长：数小时 → 分钟级处理。一个客诉问题从提出到拿到结论，原来要走完整条人墙链路、动辄一整套流程，现在客服当场就能解决大部分。

覆盖率：xx%（绝大部分问题）。客服遇到的支付风控类问题，有一大半不用再惊动分析师，智能体直接搞定。

准确率：xx%（达到较高可用水平）。在它给出结论的那些问题上，判断准确率达到较高水平；拿不准的，它不逞强，转人工。

AI 带来的不只是这三个数字。真正的价值，在数字背后。

分析师被解放了。那些被重复问题占掉的工时，还给了分析师。他们重新有时间去做真正该做的事，研究新欺诈手法、对抗黑产、设计新策略。你那笔最贵的人力投入，终于花在了刀刃上。

把分析师从一件配不上他能力的重复劳动里捞出来，放到更值钱的事上去。斯坦福报告里有一个我印象极深的案例，和我这个项目几乎是一回事：一家公司的安全运营团队，六个人，每月处理一千五百条告警，大部分是误报，全靠人工一条条筛。上了 AI 之后，告警处理能力涨到每月四万条，需要的人力从六个降到一点五个，但没有一个人被裁，空出来的四点五个人，转去做了威胁狩猎、安全架构这些更高阶的事。报告里那位高管的一句话，我觉得每个老板都该记住：AI 不是替代你已经有的人，是替代你本来要招的人。你现有的人，现在能干两个、三个、四个人的活。

这话用在我们这个项目上一样成立：智能体接走了重复劳动，不是让分析师失业，是让现有的几个分析师，顶得上原来好几倍的人，你省下的，是未来本来要扩招的成本，不是现在这帮人的饭碗。这个区别，你想清楚了，推 AI 时来自团队的阻力会小一大半；想不清楚，你一开口“上 AI”，底下人就跟你离心离德。

经验被沉淀下来了。原来锁在几个老人脑子里的判断力，现在变成了一套写下来的逻辑、一个跑着的智能体。老人离职，能力不再随之消失；新人上手，有一个随时可问的“老师傅”。这套智能体，本质上是把团队最值钱的经验，变成了公司能留住、能复用的资产。

这才是这个项目最深的价值：它不只是省了时间，它把一个组织押在个人身上的脆弱能力，变成了留在系统里的稳定能力。

为什么这个项目能做成：四个特征，帮你判断自己的场景

讲到这里，你大概会问一个最实在的问题：这个项目能做成，是不是因为它特殊？我公司的活，适不适合也这么干？

这是个好问题，而且有答案，主要在以下四条。你也可以拿这四条，去量你自己手上的活，能对上几条，就大致知道它适不适合做成 AI 工作流：

第一，高频、重复。这件事得反复发生、量足够大，投入做一套自动化才划得来。我们客诉，一天海量、且大半重复，完美符合。如果一件事一个月才发生两三次，那不值得为它搭智能体。

第二，有清晰的成功标准。AI 得能判断自己干得对不对。误拦还是真风险，是个有明确答案的判断题，对就是对，错就是错。如果一件事连人都说不清怎样算干对了，AI 更没法干。

第三，错误可以补救，不是一锤子致命。智能体偶尔判错，后面能被人发现、能纠正、能转人工兜底，而不是错一次就酿成无法挽回的大祸。我那个设计里，拿不准就转人工，就是给出错了留了退路。反过来，如果一个场景错一次就是人命、就是巨额损失，那就要慎之又慎，大幅提高人工把关的比例。

第四，数据能跨系统拿到。AI 得能够到它判断所需要的数据。我那个项目，智能体要查策略、

查画像、查历史，这些数据是能调到的。如果你要做的事，关键数据压根没线上化、散在纸上或某人脑子里，那对不起，先回去补第 2、3 章讲的数据课。

这四条，就是一张筛子。你公司里那些反复发生、有明确对错、错了能补救、数据拿得到的活，就是你最该优先下手的 AI 场景；反过来，那些偶发、说不清对错、错一次就完蛋、数据还散着的活，你先别碰。老板不需要懂技术，但你得会用这把筛子，筛出哪些活值得让 AI 去啃。

老板该学到什么

如果你是企业老板，从这个案例里希望你带走四条判断，它们能用到你公司任何一个 AI 项目上。

第一，顺序是先流程，再 AI，最后才是智能体。

不要一上来就问“我要不要做个智能体”。先问：这个岗位上，反复发生、占人最多的事是什么？这件事的判断逻辑，能不能拆清楚、写下来？流程没拆清楚之前，任何 AI 上去都是乱的，它不知道该嵌在哪、按什么标准干。流程不清，AI 项目几乎一定失败。这是我见过最普遍、最致命的坑，后面讲“坑”那一章还会专门说。

第二，最值钱、也最难的活，是把隐性经验显性化。

你公司里一定有这样的人：活干得最好，但说不清自己怎么干的。这些“说不清的手感”，就是你最大的隐性资产，也是 AI 能不能落地的关键。把它们挖出来、写下来，AI 才有的学。这件事没法外包给纯技术的人，他不懂你的业务；也没法只靠业务的人，他不知道怎么变成 AI 能用的东西。它需要一个既懂业务、又懂 AI 边界的人，坐下来一笔一笔把经验抠出来。

第三，别贪 100%，守住 80/20。

让 AI 处理高频重复问题，把高风险、高价值的复杂问题留给人。这不是退而求其次，这是企业 AI 最稳的落地姿势，这往往是效率最高的姿势。追求全自动、全替代的项目，死亡率最高。

第四，用那四个特征，挑你该先做的场景。高频重复、有清晰对错、错了能补救、数据拿得到，能对上的，优先做；对不上的，先别碰。这是你不用懂技术、也能做的判断。

一句话

这章收束成一句话：**AI 落地不是从买一个多聪明的 AI 开始，是从把一个岗位的重复经验，拆成一条 AI 能跑的工作流开始的。**

而拆经验这件事，慢、笨、磨人，却是整件事里最值钱的部分。它做不好，后面再强的模型也白搭；它做好了，流程大幅压缩，只是顺带的结果。

第 8 章 AI 内容工厂如何改变生产方式

一条短视频的成本，从 100 美金降到 15 美金

同样是一条用来带货的短视频，放在跨境电商这门生意里，这两个数字之间差的不是十几倍成本，是两种完全不同的生意做法。100 美金一条的时候，你在做内容；15 美金一条、一天能发一两百条的时候，你在做的是内容工厂。

这是我们团队给一个跨境电商团队做的项目，他们在 TikTok 上带货卖鞋，常见的铺量打法：靠海量短视频去砸自然流量，谁的内容铺得多、跑得勤，谁就有更大概率被算法选中、跑出爆款。问题是，内容这东西原来是靠人一条一条做的，而靠人做，这门生意从一开始就算不过来账。

我想用这个项目讲清楚一件很多老板对 AI 内容有误解的事：AI 在内容上真正值钱的地方，不是帮你写一条更好的文案、剪一条更炫的视频，是把整条内容生产流程重构成一个能自动跑的工厂。价值在流程，不在单点。

这个判断，和上一章是一脉相承的。上一章讲的是把一个岗位的重复经验拆成一条工作流，这一章是把整条内容生产线变成一条工作流，只不过场景从风控客诉换成了带货短视频，规模从一个岗位放大到了一条产线。但底层是同一件事：AI 落地的入口，是流程的重构，不是某个单点的替换。

在讲怎么做之前，先得让你看清楚，靠人做内容，这账到底烂在哪。

原来什么样：一门从第一条视频就开始亏的生意

跨境电商带货，尤其是在 TikTok 这种平台上靠短视频走量，生意的底层逻辑是这样的：你不知道哪条视频会爆。没有人知道。你能做的，是大量地发，用数量去博概率，发得越多、测得越勤，跑出爆款、被算法推上量的机会就越大。这是个用产量换概率的游戏。

这个逻辑本身没错。错的是，当你的内容靠人一条一条做的时候，用产量换概率这条路，从一开始就走不通。

我们来算一笔账。

第一，单条成本压不下来。一条带货短视频，靠人做要做哪些事？要去看现在什么内容在火、什么钩子能抓住人，要写脚本、要找素材或者拍摄、要剪辑、要配上字幕和音乐、要适配平台的调性，最后才能发出去。这一整套下来，一条的成本在 100+ 美金。

这个数字单看还好，吓人的是它乘上“铺量”两个字。

铺量打法的本质是，你必须接受绝大多数视频是不会爆的，可能一百条里跑出来三五条，剩下九十几条就是炮灰，是你为了博那几条爆款必须付出的测试成本。如果一条炮灰的成本是 80 美金，你一天想发一百条，光内容生产这一项，一天就要烧掉八千美金，而这里面绝大部分钱，是花在那些注定没有水花的视频上的。这不是贵，这是这门生意在靠人做内容的前提下，根本算不过来账。

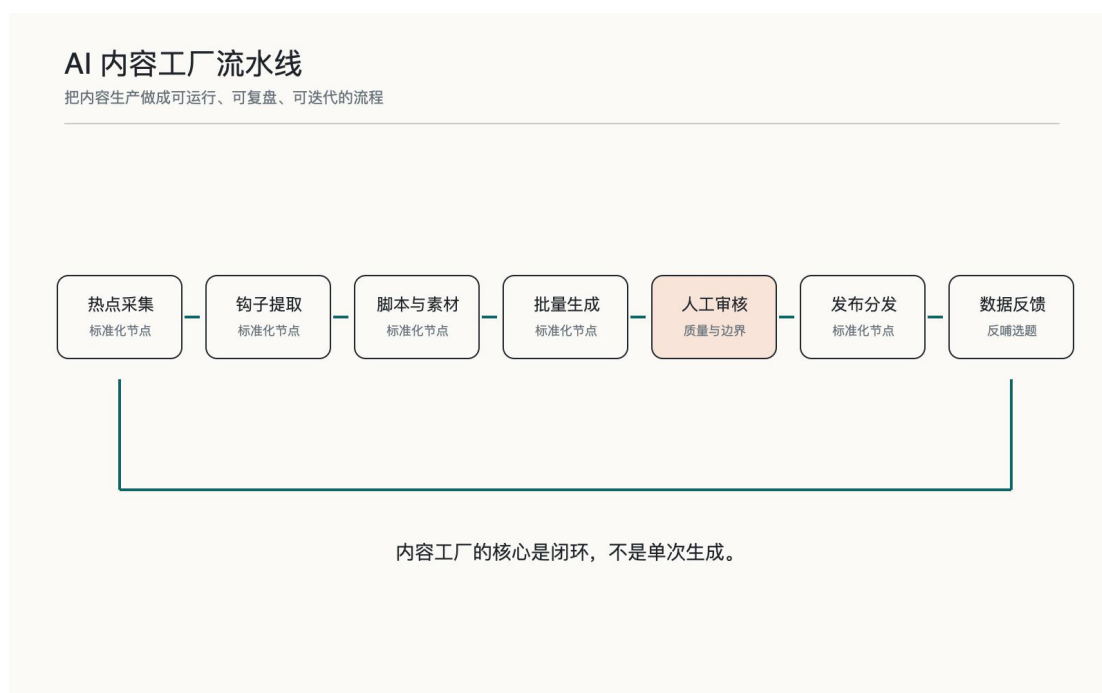
第二，产量天花板被人手死死焊住。就算你不在乎成本，愿意烧钱，人做内容还有一道过不去的坎：慢，一个熟练的内容人员，一天能做多少条像样的带货视频？几条、十几条到头了。你想一天发一百条、两百条，就得堆人。堆到几十个人的内容团队，管理成本、协同成本、人员流动，全冒出来了。而且人多了，内容质量和风格还会散，这个人理解的钩子是这样，那个人理解的是那样，产出参差不齐。

铺量这门生意要的是“量大、稳定、快”，而人手恰恰是“量小、不稳、慢”。你想要的和人能给的，方向是反的。

第三，跟不上热点的速度。TikTok 上的内容流行是以小时、以天为单位变的。今天火的钩子、火的玩法、火的 BGM，可能后天就过气了。靠人去追，发现一个热点、研究怎么用、做出内容、发出去，这个链条走完，热点窗口往往已经关了。

但到这里，成本高、产量低、追不上热点，这三个问题，没有一个是靠让 AI 写一条更好的文案能解决的。你哪怕请来全世界最会写带货脚本的人，他一天还是只能写那么几条，成本还是降不下来，热点还是追不上。单点再强，救不了一条算不过来账的产线。这门生意的病根，不在每一条内容做得好不好，在整条生产线的成本结构和产能结构。要治，就得动整条线，不是动某一个环节。

我们怎么做的：把内容生产，拆成一条能自动跑的流水线



很多人一听 AI 做内容，第一反应是：那不就是用 AI 生成视频、生成文案吗？这又是把事情想小了。我们做的不是用 AI 生成内容这一个动作，是把原来靠人串起来的整条内容生产链，从找热点到发出去，拆成一条能自动运转的流水线。AI 生成，只是这条流水线上的其中一道工序。这条流水线，我们拆成了四段。

第一段：热点采集，让机器去盯着平台的风向

人追热点追不上，那就让机器去盯。这条产线的最前端，是一套自动采集热点的机制：持续地从平台上抓正在跑量、正在被推的内容，看什么样的钩子、什么样的玩法、什么样的呈现方式现在有流量。这件事机器做有两个人比不了的好处：一是不知疲倦，可以 7×24 小时盯着；二是快，热点一冒头就能被捕捉到，而不是等人某天刷到了才发现。这一段解决的，是前面说的第三个病，**追不上热点**。把靠人偶然刷到变成机器持续监测，热点的捕捉从被动变成了主动。

第二段：钩子提取，把“为什么火”从热点里抠出来

抓到了正在火的内容，不等于你就会做内容了。关键是：它为什么火？一条带货视频能跑起来，核心往往就在前几秒那个钩子，是一句话、一个画面、一个反差、一个痛点，让用户划到这儿停下来、看下去。铺量打法里，你要复制的不是某一条具体视频，是它背后那个有效的钩子结构。所以流水线的第二段，是从采集到的热点内容里，把钩子提取出来：这条内容是用什么方式抓住人的？这个钩子能不能套用到我们自己的鞋上？把热点从“一条没法直接抄的视频”，变成“一个可以复用的钩子模板”。这一段是整条线里最像在沉淀经验的一环，它做的事，本质上和上一章里把老分析师的判断逻辑抠出来是一类动作：把一个原来藏在爆款里、说不清道不明的“为什么火”，变成可以被机器批量复用的结构。

第三段：AI 生成变体视频，一个钩子，批量裂变成上百条

有了钩子模板，有了我们自己的商品，这一段才轮到 AI 生成正式上场。做的事是生成变体。基于一个验证过的钩子结构，套上不同的商品、不同的画面组合、不同的呈现方式，批量生成大量“同一个钩子、不同表现”的视频。一个有效的钩子，可以裂变出几十上百条变体，撒到平台上去测，看哪个变体、哪个商品组合能跑出来。这一段直接解掉了前面两个病，成本和产量。单条的生产成本被压下来，产量则从“一个人一天几条”变成“机器一天上百条”。铺量打法最需要的量大、快、稳，到这一段才真正拿到手。

第四段：自动发布与数据反馈，让产线自己越跑越准

最后一段，是把生成好的内容自动发布出去，形成持续的内容供给。但这一段真正的价值，不在“发”这个动作，在“发完之后”。

内容发出去，平台会给你反馈：哪条跑量了、哪条没水花、哪个钩子的变体表现好、哪个商品组合点击高。这些数据回流到产线最前端，就构成了一个闭环：机器盯热点 → 抠钩子 → 生成变体 → 自动发布 → 数据回流 → 指导下一轮盯什么、抠什么、生成什么。

这个闭环是整条产线最关键、也最容易被忽略的一环。它意味着这条线会根据数据自己调整、越跑越准，而不是按固定模板机械地产出。上一轮哪个钩子跑出了爆款，下一轮就往那个方向多生成；哪个商品组合没人看，下一轮就少做。铺量打法的本质是用产量博概率，而这个数据闭环，做的事是让你每博一轮，命中的概率都比上一轮高一点，产量在涨，精度也在涨。

四段连起来，你会发现关键不在任何单独一段，而在于它们被串成了一条不用人一步步推、能自己转、还能自己越跑越准的线。这才是工厂和工具的根本区别：工具是你用它做一条内容，做完它就停在那儿；工厂是它自己源源不断地产出、还会根据反馈自我校准。一个是死的，一个是活的。

真正的坑，藏在“每个 SKU 都得单独调”里。

讲到这儿，听起来好像挺顺。但我们也踩了很多坑，这条产线不是一次就调通的，是调了很多个版本才做出来的。而最大的坑：每一个 SKU，都得单独做精细化调整。

什么意思？你卖的不是一双鞋，是一个一个不同的款。运动鞋和小皮鞋，卖点不一样、适合的钩子不一样、画面该怎么呈现不一样、用户被什么打动不一样。你不可能用一套参数、一个模板，套到所有鞋上还都能跑得好。哪个钩子配哪个款、画面怎么生成、节奏怎么处理，都得针对这个具体的 SKU 去调、去试、去校准。

这件事没有捷径，它意味着这条产线每上一个新品类、新款式，都得有人沉下去重新调一轮，不是搭一次就能永久通用的。这是 AI 内容工厂里最被低估、也最容易让人栽跟头的部分，很多人以为做个内容工厂就是搭个能批量生成的系统，真做起来才发现，大量的功夫花在针对每个具体商品反复调试这种又细又笨的活上。

具体要调什么？就拿鞋子来说，这个款防不防水？鞋面是真皮还是合成革、皮质会不会脱胶？这些卖点维度，直接决定了钩子怎么写、画面该突出哪儿、节奏怎么配。换个品类，差异更大：真要换成首饰，该调的维度就完全是另一套了，材质是银还是合金、会不会含镍过敏、是日常款还是送礼款、突出的是设计感还是性价比，跟鞋子那套几乎没有一个能复用。所以这条产线，我们光把鞋子这一个品类调顺，前前后后就改了一二十版。每上一个新品类、甚至每上一批新款，这个又细又笨的调试过程，基本都得从头再来一遍。

把这件事单独拎出来讲，是因为它对应着上一章那个同样的判断：AI 项目里最值钱、最难的活，往往是那些又细又笨、没法跳过的“脏活”。上一章是把老师傅的手感一笔笔抠出来，这一章是把每个 SKU 一款款调出来。形态不同，本质一样，都是 AI 项目里那个性感不起来、却决定成败的环节。

拿到什么结果

这条产线跑起来之后：

单条成本：大幅下降。直接降到原来的十分之一左右。原来一条炮灰视频的成本能让你心疼，现在一条的成本低到你可以放心大胆地用产量去博概率。

日产量：稳定批量生产。从一个人一天几条，到一条产线一天上百条。铺量打法最核心的量，彻底解决了。

自然流量：持续增长。持续的内容供给，换来了持续的、不靠付费投放的自然流量。

原来：单条成本 80 美金，一个人一天做几条。你想铺量，要么烧不起钱，要么堆不起人，铺

量这条路在经济上根本不成立。

现在：单条成本 15 美金，一条产线一天上百条。铺量需要的量大、低成本、快速度，同时拿到了。那个原来算不过来的账，现在算得过来了。

这就是“重构流程”和“优化单点”的根本区别：优化单点，是让你那几条人工视频做得更好，改变的是质量，改变不了生意能不能成立；重构流程，是把整条产线的成本结构和产能结构换掉，改变的是这门生意本身能不能做。

AI 在这里带来的不是内容更好了，是一种原来做不了的生意，现在能做了。

AI 创造价值最高的那一批，往往把原本做不到的事变得可能，而不只是把成本降下来，有的公司用 AI 做成了过去因为太贵、太慢而根本没人敢碰的业务。那些把 AI 用在营销内容上的公司，真正的收获普遍来自内容的量和精准带来的增长，不只是省几个文案的工资，有家企业的内容平台，AI 生成的活动点击率直接翻了一倍，上线周期从七周压到六小时。注意这里的逻辑：省下来的那点制作成本是小头，产能和精度被解放之后带来的增长才是大头。跨境电商这个项目也一样，单条从 80 美金降到 15 美金当然爽，但真正改写这门生意的，是铺得起量之后跑出来的那些爆款和源源不断的自然流量。所以你看 AI 内容的价值，别只盯着省了多少制作费，要盯着它能不能让你做成原来根本做不成的增长。

老板该学到什么

如果你的生意里有内容这一环，不管是带货、是营销、是品牌、还是获客，从这个案例里，希望你带走三条判断。

第一，别把 AI 内容理解成生成一条文案/视频，要理解成重构一条产线。很多老板对 AI 内容的期待，停在帮我写篇推文、剪个视频。这个用法不能说错，但它只动了一个单点，改变不了你内容生产的成本结构和产能结构。真正能改变生意的，是问自己：我的整条内容生产流程，从找选题到发出去，能不能变成一条机器能自动跑的线？单点生成省的是一个人的几分钟，流程重构省的是整条产线的成本量级，这两者差着一个数量级。

第二，先看你的生意算不算得过来账，再决定 AI 该动哪。这个案例的要害在于：AI 把一门算不过来账的生意，变得算得过来了。你判断该不该上 AI、该上在哪，核心问题应该是：我这门生意的账，卡在哪个环节算不过来。卡在成本结构上，就动成本结构；卡在产能上，就动产能。AI 要动的，永远是那个让你账算不过来的环节，不一定是那个最容易上 AI 的环节。（这条判断，后面讲哪些场景值得做那一章，我会展开成一个完整的方法，给每个问题标个价，看 AI 该砸在哪。）

第三，真正的功夫，都在针对你自己业务的精细化调整上。这条产线最重的活，不是搭一个能批量生成的系统，那部分有现成的能力可用。最重的活，是针对每一个 SKU 单独调。任何号称通用、一键、搭起来就能跑的 AI 内容方案，你都要打个问号：它有没有针对你的具体业务、具体商品、具体场景去精细化调过？没调过的，大概率跑不出结果。AI 内容工厂不是买来即用的标准件，是要沉下去、针对你的业务一点点调出来的，这部分调试的功夫，才是它能

不能为你赚到钱的真正分水岭。

一句话

这章收束成一句话：**AI 内容真正值钱的地方，不是帮你做出一条更好的内容，是把整条内容生产线变成一座能自动跑的工厂，它改变的不是内容的质量，是这门生意能不能成立。**

单条 15 美金、日发上百条，这些数字本身不重要。重要的是它们合在一起，把一门原来算不过来账、铺不起量的生意，变成了一门账算得过来、量铺得起来的生意。这才是流程重构和单点优化的根本差别。

第 9 章 从单点工具到组织级协同

一个点餐推荐，背后是整个集团的数字化重构

先说一件听上去很小的事：让服务员在点餐时，主动给客人推荐一道菜。

就这么一件小事。但在一家区域特色餐饮连锁企业里，我们为了让这个推荐真正跑起来、跑得稳，最后动的是整个集团的数据和系统。

这是我们团队正在做的一个项目，客户是一家区域特色连锁餐饮企业，主打健康餐饮概念。它有多品牌、多门店、跨区域运营。我们的角色是智能体开发和技术总协调，和另外几方一起在做。

我把这个案例放在第二篇的最后一章，是因为它要回答的，是这一整篇案例里层级最高的一个问题：AI 在一家企业里，是怎么从一个单点工具，一步步长成一种组织运行方式的？

前面几章讲的都是单点：一个客诉智能体、一条内容产线、一个岗位的工作流。它们都很有价值，但它们都是点。这一章要讲的，是怎么从点连成线、从线长成面，从一个服务员手里的点餐推荐，到一个店长的经营管理，再到整个集团总部的品牌管控。这条从“个人提效”到“组织能力”的路，是这本书的主线，也是它的终点。

在讲这套分层体系怎么搭之前，得先讲清楚：这家餐饮企业原来到底卡在哪。因为这个项目最关键的判断，恰恰藏在它真正要解决的问题是什么里，而这个问题，和大多数人第一眼以为的，不一样。

原来什么样：三个看起来不相干的痛，根子是同一个

这家餐饮企业找过来的时候，摆在桌上的痛点有三个。这三个乍一看互不相干，但你顺着往下挖，会发现它们的根子是同一个。

第一个痛：服务员推不动该推的菜。

餐饮的利润，很大程度上藏在推荐里，同样一桌客人，服务员引导他们点什么、不点什么，直接影响这一桌的毛利。这家餐饮企业的管理层很清楚这一点，他们希望服务员在点餐时能主动去推那些毛利高、同时库存又充足的菜。

但这件事在现实里几乎做不稳，因为：服务员流动大、整体素质参差不齐。餐饮业的用工现实就是如此，人来人往，今天培训好的人，过两个月可能就走了。而什么时候该推哪道菜这件事，本质上依赖经验和判断，要同时考虑这道菜赚不赚钱、现在库存够不够、客人是什么口味、当下是什么时令。这套判断，靠人去记、去权衡，老员工或许做得到，但你没办法让一支

高流动、参差不齐的队伍稳定地做到。

第二个痛：养生知识太多，人脑根本记不住。

这家餐饮企业主打健康餐饮这个概念，这意味着它的菜不只是好吃，还带着养生属性，不同人群适合什么搭配、不同产品有什么功效、哪些人群需要避开，这是它区别于普通火锅店的卖点。

但这个卖点有个要命的落地问题：这套健康餐饮的知识库太大了，大到没有任何一个服务员能记得住。你不可能要求一个流动性极高的服务员队伍，人人都成为半个养生专家。结果就是，这个本该是核心竞争力的卖点，在一线根本传递不到客人那里，知识躺在那儿，但用不出来。

第三个痛：老系统散乱，还塌了。

更现实的是，这家餐饮企业原来用的那套业务系统，服务商爆出资金危机，服务稳定性和后续支持都成了问题。系统散乱、各自为政。这已经不是升级一下能解决的事了，整个系统必须重构。

现在，把这三个痛放在一起看：推荐推不动、知识记不住、系统要重建。

大多数人看到这三个痛，第一反应会是：那就做一个聪明点的点餐推荐工具嘛，能推荐高毛利的菜、能讲养生知识、界面做得好用一点。把它当成一次点餐系统的升级。

但我们对这个餐饮项目的本质判断是：它要的不是一个更聪明的点餐工具，它要的是一次完整的餐饮企业数字化重构。

为什么这么说？你回头看那三个痛：服务员推不动，是因为推荐该依据什么这套逻辑没有沉淀在系统里；知识记不住，是因为健康餐饮的知识没有被治理成机器能调用的形态；系统散架，是因为各个数据原本就没打通。这三个痛，没有一个是“点餐界面不好用”造成的，全都是“底层数据和经验没有被组织起来”造成的。它们指向的是同一个根：这家企业的数据、知识、经验，是散的。

所以我们做这个项目的办法，从一开始就反着来，不去先设计一个点餐 App，而是先问一个智能体要真正跑起来，需要哪些数据、哪些标准、哪些流程，再用这个去倒推整个企业该怎么把它的数据和系统重新组织起来。

这是这一章第一个、也是最重要的判断：当你看到几个看起来不相干的业务痛点时，别急着一个一个打补丁。先往下挖，看它们是不是同一个根。如果是，那你要解决的就不是几个工具问题，而是一次底层的重构。

我们怎么做的：分层智能体体系，一层一层往组织里长

想清楚了这是一次数字化重构，不是工具升级，接下来的问题就变成了：这么大的重构，怎么落地？

答案不是一口气把整个集团的系统全推倒重来，那是贪大求全，死亡率最高的做法。我们的做法是分层：把这次重构拆成分层智能体体系，从最贴近一线、最容易出成果的那一层做起，一

层跑通了，再往上长一层。

这套分层体系，恰好对应了 AI 在一家企业里成长的三个阶段：面向顾客、面向门店、面向总部。从一个点，到一家店，到一个集团。

第一阶段：菜单营销智能体：面向顾客，先在一线把单点跑通

第一层，叫菜单营销智能体，面向的是点餐这个最一线的场景。它要解的，就是前面那两个最直接的痛，推荐推不动、知识记不住。

它的核心，是一套多因子的推荐逻辑。服务员该给这桌客人推哪道菜，不再靠他自己的经验去权衡，而是由智能体根据多个因素综合算出来。这套推荐逻辑会把几类因素放在一起综合判断：客人的需求和口味、菜品毛利、库存情况、时令特点、历史表现。

你注意这个排序逻辑，它把客人真正适合什么放在前面，再兼顾企业经营目标。推荐这件事一旦本末倒置，先想着赚钱、不管客人想不想要，推荐就变成硬推，客人反感，长期反而伤生意。把需求放在前面，赚钱放在后面，这套逻辑本身就是一个经营判断的显性化。

至于那套人脑记不住的健康餐饮知识，在这一层被治理成了智能体可以随时调用的知识库，服务员不需要再去背，需要讲养生功效的时候，智能体直接给到。那个原来用不出来的核心卖点，现在能稳定地传递到每一桌客人面前了。

但第一阶段真正的分量，不在推荐这个动作本身。在于：为了让这个推荐算得准，我们必须把这家企业多类原本散开的核心数据，第一次围绕“菜单”这个枢纽打通。

包括菜品、原料供应链、健康餐饮知识库、销售、会员、门店等。前面那套推荐权重要算得出来，就必须同时知道：这道菜是什么（菜品）、成本和库存够不够（供应链）、有什么养生功效（知识库）、卖得好不好（销售）、这个客人是谁（会员）、这家店什么情况（门店）。这套推荐逻辑，逼着这些核心数据域必须以一套统一的标准接到一起。

这就是第一阶段最深的价值：它表面是个点餐推荐，骨子里是借点餐推荐这个一线最容易出成果的场景，把整个企业的数据底座给搭起来了。这个底座一旦搭好，后面的第二阶段、第三阶段才有地基可站。

这正是我反复强调的那条落地原则：从最贴近一线、最容易出成果、最能算清账的那个单点切进去，但切进去的同时，要为后面的“面”埋好数据的种子。单点跑通是为了拿到第一个成果、建立信心，而打通多类核心数据域，是为了让这个单点不只是一个孤立的工具，而是组织能力的第一块地基。

第二阶段：店面经营智能体：面向店长，从单点长到门店

第一阶段把一线的点餐跑通、把数据底座搭好之后，第二层往上长一级，到门店，店面经营智能体，面向的是店长。

店长是干什么的？一家店的人、货、钱，全压在他身上：排班用人、采购进货、出品管理、财务对账……一摊子事。而这一摊子事里，有相当一部分是数量和价值的管理，这个月该备多

少货、各项成本控制在多少、对账对得平不平。这部分工作，高度依赖数据、依赖标准、依赖反复核算，恰恰是 AI 最擅长、最该接手的。

所以第二阶段做的事，是把店长身上那部分“数量和价值管理”的活，剥离出来交给智能体，覆盖人力、采购、出品、财务、对账等一批经营场景。智能体把这些算清楚、管起来之后，店长就被解放出来，可以专注于智能体干不了、也最需要人的那部分，行为管理：怎么带团队、怎么抓服务、怎么处理现场、怎么管人。

这一层和第七章那个客诉智能体的逻辑，是完全相通的：让 AI 干那些高频、重复、靠数据和规则就能算清的活，把人解放出来，去干那些真正需要人的判断、人的温度的活。只不过这一次，被解放的不是风险分析师，是店长；场景不是一个岗位，是一整家门店的经营。

这里其实藏着一个关于“AI 怎么改变组织”的更深的道理。你有没有想过，一家企业里那么多层级、那么多管理动作，是怎么长出来的？很大一部分，是为了“管住”那些靠数据和规则就能算清、但人工算起来又慢又容易出错的事，备货多少、成本控在哪、账对不对。AI 真正在做的，是把这一层为了管住可量化的事而存在的活接走，让组织回归到它真正的核：判断、创造、和人打交道。店长就是个缩影，智能体接走了他那部分数量和价值的管理，剩下的带团队、抓服务、处理现场这些必须靠人的判断和温度的事，才是他这个岗位真正的核。AI 不是把人换掉，是把组织里那些本不该靠人脑硬算的部分剥走，让人回到只有人能干的那部分上去。一家企业 AI 转型转到深处，组织往往会变得更瘦、也更实，瘦掉的是那些可被算法替代的中间管理动作，留下的是判断、创造和人际这些真正的核。

到这一层，AI 已经从一个一线工具长成了一家门店的经营助手。它管的不再是一桌客人点什么菜，而是一家店怎么运转。从点，到面了。

第三阶段：品牌管理智能体：面向总部，从门店长到集团

后续更高一层，品牌管理智能体，面向的是集团总部。这一层规划在后续阶段落地。

当多家门店的数据，都通过第一阶段的底座和第二阶段的经营管理沉淀上来之后，总部就第一次有可能站在集团的高度，去做原来根本做不了的事：

跨门店诊断评优，把多家门店放在一起比，谁经营得好、谁出了问题、好在哪、差在哪，一目了然。

集团目标的层层管控，总部定的目标，能拆解、下发、追踪到每一家店，而不是发个文件就石沉大海。

统一供应商、统一研发、统一套餐，在集团层面统一管控供应链、菜品研发和套餐设计，把连锁的规模优势真正发挥出来。

走到这一层，AI 管的已经不止是一桌客人、一家店，而是整个集团怎么运行。它成了总部管理这多家门店的方式本身。

从点（一桌客人的推荐），到面（一家店的经营），到体（整个集团的管控），这就是 AI 在一家企业里，从单点工具长成组织运行方式的完整路径。

一个老板必须看懂的设计：这套分层体系为什么是同一棵树

要把这件事点出来，因为它是整个项目最值钱的设计，也是最容易被忽略的地方。这第一阶段、第二阶段、第三阶段，不是三个独立的项目，是同一棵树的根、干、冠。

它们靠什么连成一棵树？靠第一阶段那个以菜单为枢纽打通的多类核心数据域。第一阶段搭好的那套统一数据底座，是第二阶段店长经营管理的数据来源，也是第三阶段总部跨店诊断的数据来源。没有第一阶段这个底座，第二阶段就是无米之炊，第三阶段更是空中楼阁。

这就是为什么一开始说它要的是一次数字化重构，不是一个点餐工具。如果只把第一阶段当成一个孤立的点餐推荐去做，做完了就完了，那它永远是个单点工具，长不出第二阶段和第三阶段。正因为从第一天起就把它当成“整棵树的根”来设计，让一个点餐推荐去承担“打通多类核心数据域”的任务，这个单点才具备了往上生长成组织能力的可能。

这件事还涉及一个协同问题，这个餐饮项目不是我们一家在做，是多方协同：我们（智能体开发和技术总协调）、原来的业务软件方、飞书管理平台方、财务系统方，加上甲方自己。

在这些协作方里，我们承担的是最核心的两块：一是开发智能体本身，二是帮他们把底层的数据地基搭起来，梳理知识库、定义系统需要哪些表、哪些字段、每个字段该怎么用。这个定义字段的活看着不起眼，其实是整个协同的枢纽：我把需要哪些字段、怎么用讲清楚，甲方的其他部门才知道该往里填什么数据，业务软件方和财务系统方才知道该按什么口径来对接。我们最后把智能体和知识库都部署在云平台上，对外提供一个 API，其他各方按这个 API 把数据传进来，我们运算完，再把结果输出回去。

这么多方要协同，靠的也正是那套统一的数据标准，所有系统都得按这个标准对接，数据才能在这棵树里流得通。数据标准不统一，各方就是一个个孤岛，这棵树就长不起来。

（这个项目目前还在进行中，近期才启动，第一阶段在落地、第二阶段在规划、第三阶段还在更远的路线图上。我讲它，不是因为它已经跑出了多漂亮的成果，那些还没到兑现的时候。我讲它，是因为它的设计思路最完整地展示了单点怎么长成组织能力这条路。成果要等时间，但思路现在就值得你看懂。）

拿到什么结果

因为项目还在进行中，没办法晒提升了百分之多少这种还没兑现的数字。希望你看的，是这个项目已经带来的、结构性的变化。

第一，那个原来传递不到客人面前的核心卖点，落地了。健康餐饮的知识，从躺在那儿没人记得住，变成了智能体随时能调用、能稳定讲给每一桌客人。一个高流动、参差不齐的服务员队伍，现在能稳定地输出原本只有专家才讲得清的养生推荐。**这家企业最大的差异化卖点，第一次真正能在线兑现了。**

第二，推荐这件事，从“靠人的手感”变成了“靠系统的逻辑”。该推什么菜，不再取决于今天上

班的是老员工还是新人，而是由那套多因子权重稳定地算出来。企业最怕的“能力押在个别人身上”，在点餐这个环节被解掉了。

第三，也是最关键的，这家企业第一次有了一套打通的数据底座。菜品、供应链、知识库、销售、会员、门店等原本各自为政的数据域，围绕菜单被统一组织了起来。这件事的价值，远远超出点餐本身：它是这家企业从“一堆散系统”走向“一个数字化组织”的第一块地基。有了它，门店级的经营管理（第二阶段）、集团级的品牌管控（第三阶段）才谈得上。

这三个结果，没有一个是点餐推荐准确率提升了多少这种单点指标。它们全都是结构性的、组织级的变化，能力沉淀进了系统、数据打通成了底座、卖点真正落了地。这正是组织级协同和单点工具的根本区别：单点工具给你一个更好的指标，组织级协同给你一个不一样的组织。

老板该学到什么

这是第二篇案例的最后一章，希望你带走的三条判断，分量也最重。

第一，几个不相干的痛点，往下挖往往是同一个根，别一个个打补丁，要看是不是该重构。

这家餐饮企业那三个痛，推荐推不动、知识记不住、系统散架，表面上是三件事，根子上是同一件：数据和经验是散的。如果你只盯着每个痛点单独打补丁，做一个推荐工具、配一个知识手册、修一修旧系统，你会得到三个互不相干的小工具，问题一个没真解决。真正该问的是：这些痛点背后，是不是同一个底层问题？如果是，那要做的就不是几个工具，是一次重构。这个判断，是区分小修小补和真转型的分水岭。

第二，从单点切入，但要让单点为“面”埋种子，切的时候就想好它怎么长。

我们没有一上来就推倒重来，而是从最一线、最容易出成果的点餐切进去，这是稳。但我们没有把点餐做成一个孤立的工具，而是让它去承担打通多类核心数据域的任务，这是远。好的 AI 落地，既要近（从能出成果的单点切入，快速建立信心），又要远（让这个单点为后面的组织能力埋好地基）。只近不远，你会得到一堆永远长不大的孤立工具；只远不近，你会陷进贪大求全、迟迟出不了成果的泥潭。又近又远，单点才能长成组织。

第三，AI 转型的终点，不是人人会用工具，是 AI 成为组织运行的方式。

前面几章的案例，客诉智能体、内容工厂，都很有价值，但它们改变的是“一个岗位、一条产线”的效率，是个人和单点层面的提效。这家餐饮企业这个案例不一样，它展示的是 AI 怎么从一桌客人的推荐，长到一家店的经营，再长到一个集团的管控。走到第三阶段那一层，AI 已经不是员工手里的一个工具，它就是这家企业管理多家门店的方式本身。

这就是前面讲过的那个区别：单点提效解决的是“一个人干得更快”，组织转型解决的是“这家公司换了一种运行方式”。前者靠工具就能启动，后者必须把数据、流程和经验一层层沉淀进系统里。

我特别想强调这个区别，因为它直接关系到你这家企业的护城河。有个判断我越做越认同：工具是最容易被复制的，组织能力是最难被复制的。你买了个智能体，你的对手明天也能买一个

一模一样的；但你把多年的经营经验、把多个核心数据域、把店长的管理逻辑，一层层沉淀进一套打通的系统里，这套东西对手抄不走，他就算买了同样的模型、同样的工具，也长不出你这套组织能力。还有句话说得更直接：AI 项目可以复制，AI 能力很难复制。一个项目、一个工具，是可交付、可复制的；但一家企业把 AI 真正长进自己数据、流程、组织里的那种能力，是独一份的。所以这个餐饮项目最值钱的，从来不是那个点餐推荐有多准，是它正在帮这家企业，长出一套对手抄不走的组织能力，而这，恰恰是你掏钱做 AI 转型，最该买到的东西。

一句话

这一章，也是第二篇案例，收束成一句话：**AI 转型的终点不是人人都会用工具，是 AI 从一个单点工具，长成你这家企业运行的方式本身，从一桌客人的推荐，到一家店的经营，到一个集团的管控，一层一层长上去。**

这家餐饮企业的分层智能体体系之所以能连成一棵树，是因为从第一天起，那个最一线的点餐推荐，就被设计成了“整棵树的根”。这就是单点工具和组织能力的分水岭：单点是用完即止的工具，组织能力是能往上生长的根。

第 10 章 企业 AI 落地最常见的 20 个坑

先看一个数字

一份 2025 年企业生成式 AI 研究提出，约 95% 的试点没有产生可衡量的损益回报。这个数字说的是研究样本里的损益结果，不该被读成所有 AI 项目都没有价值。它提醒你：钱花了、demo 跑通了，生意上依然可能看不到回报。

我服务过几十家企业，自己在某大厂做了几年风控 AI，这个数字我一点都不意外。绝大多数 AI 项目都死在一堆“AI 之外”的地方。这一章我们把 AI 落地拆分成二十个坑，按最容易踩的顺序排好。

有研究报告发现，77% 最难啃的挑战都在技术之外，变革管理、数据质量、流程重做。所以你接下来看到的这二十个坑，绝大多数都是人、流程、组织的坑。这其实是个好消息：技术你可以买、可以外包，但这些坑，只能靠你这个老板自己看清、自己躲。

我建议你这么读：每看一个坑，停一秒，问自己一句，我公司是不是也这样？如果你看到一半就开始冒汗，说明这章对你有用。

二十个坑分四组：

认知类：老板自己最容易犯的，放最前面（坑 1-5）

执行类：项目启动和落地阶段栽跟头的（坑 6-10）

运营类：东西做出来了、却没人用的（坑 11-15）

合作类：被供应商、被合作方坑的（坑 16-20）



一、认知类的坑：老板脑子里就错了

认知类的坑最隐蔽，因为它发生在项目还没开始的时候，发生在老板自己脑子里。前面五个坑，你一个都躲不过，因为它们都是想当然。

坑 1：把买工具当成了转型

最普遍的一个坑：老板觉得，我给全公司开通了某个大模型的会员，或者采购了一套 AI 办公系统，这事就算转型了。

工具是入口，不是转型。我见过太多企业，账号买了一大批，三个月后一看后台，活跃用户不到一成。工具静静躺在那儿，没人用，也没改变任何一条业务流程。

为什么？因为工具解决的是“能不能用”，转型解决的是“业务怎么变”，这是两件事。你买一台机床回来，不等于你会造零件；你得有图纸、有工艺、有工人、有质检，机床才产生价值，AI 工具就是那台机床。

有个数据可以印证：到 2026 年，Gartner 预测组织会放弃六成缺少“AI 就绪数据”的项目。工具买回来，发现喂不进去东西，最后只能放弃。买工具这个动作，是整个转型里最便宜、最不重要的一步，但很多老板把全部认知都停在这一步。

坑 2：以为培训完就结束了

第二个坑紧跟着第一个：花钱做了一次 AI 培训，员工听得热火朝天，老板觉得认知拉齐了，这事算成了。

培训当天确实热闹。但你过两周回去看，该用的还是不用，该怎么做事还怎么做事。培训能解决知道，解决不了做到，更解决不了持续做。

我把员工用 AI 的能力拆成五层：看到、知道、做到、探索、判断。培训最多把人推到第三层，做到，跑通一个基本流程。但真正决定落地效果的是第四、第五层，那靠的是组织机制（考核、激励、有人带），不是靠两天课。

大量员工听完 AI 培训之后，没办法把课上学到的东西关联到自己的实际岗位。课听了，回到工位发现用不上，于是就放下了。把培训当终点的老板，钱花在了认知上，却没花在落地上。培训是必要的，但它只是第一步里的一小步。

坑 3：被 AI 博主洗脑，预期被拉到天上

现在打开任何一个平台，都是一个人干一个团队，AI 让我月入百万，三天搭一个数字员工。老板刷多了，预期就被拉到天上去了。

我接触过一个做女装选品的客户，聊下来发现，她对 AI 的预期完全是被短视频喂出来的，她要的是一个全自动、像人一样思考的选品大脑。我们给了 demo，质量其实不差，但她回了一

句：不够 AI native，就没下文了。她心里那个标准，是博主造出来的幻象，现实里没有任何一家公司能交付。

博主卖的是流量，你买的是落地，这两件事的标准差着十万八千里。博主的视频里，AI 永远第一次就成功、永远不出错、永远不用清洗数据。企业项目里，大部分时间在干脏活累活，AI 经常答错，数据乱得一塌糊涂。被博主洗脑的老板，会用一个不存在的标准去要求供应商，最后要么觉得所有人都是骗子，要么砸钱去追一个追不到的东西。

坑 4：贪大求全，一上来就要做大系统

这个坑专坑有钱、有魄力的老板。他一旦决定做 AI，就想做个大的，一个平台，覆盖全公司，所有部门都能用，一步到位。

但客观的说，一上来就做大而全平台的项目，死亡率最高。大系统意味着长周期、高投入、多部门协同，任何一个环节卡住，整个项目就拖死了。而且大系统做出来，往往沦为摆设，因为它什么都想覆盖，结果什么都不好用。

我自己在某大厂做风控时，哪怕是大型、有数据有团队，一个策略分析的子项目都投入了很长时间，背后还牵涉大量调用、策略和智能体协同。一个中小企业，凭什么觉得能一口气吃下一个覆盖全公司的大平台？

正确的姿势是反过来的：找一个小而痛的场景，先跑通，再复制。智能体不是第一步，是第三步。这条后面讲经验那一章我会反复讲。

坑 5：以为 AI 能 100% 替代人

老板觉得既然上了 AI，就该全自动、全替代，留着人显得不够先进。

我在某大厂做 AI 风控时，我的组长选了全自动替代那条路线，让 AI 完全接管、无人值守。团队投入了非常大精力，最后失败了。我当时提的是另一条路：AI 干大部分高频重复的，少部分复杂、高风险的，交还给人。这条路后来跑通了。

为什么 100% 替代几乎必死？四个原因：

第一，大模型的底层架构决定了它的输出永远不可能 100% 准确；

第二，业务对手（黑产、竞争对手）每周变手法，没有固定答案可学；

第三，边界场景太多，自动化容错的难度是指数级上升的；

第四，没有人工兜底，错误会持续放大却没人发现。

记住这个识别信号：当有人跟你说用 AI 干掉某某岗位，完全自动化，无人值守，警铃就该响了。一个能稳定处理大部分高频任务、并且老老实实把复杂问题交还给人的 AI，远比一个号称什么都能干、但你不敢完全信它的 AI 有价值。

二、执行类的坑：项目启动和落地阶段栽的跟头

认知对了，到了真正动手的阶段，还有一堆坑等着。

坑 6：把 demo 当成了上线

这是企业 AI 项目里，我见过最贵、最普遍、最让老板措手不及的一个坑。

供应商来给你演示，画面漂亮，AI 对答如流，你问什么它答什么，老板当场拍板：就是它了，这就是 demo。然后呢？然后你以为再花一两个月、把这个 demo 接进公司系统，就能上线用了。

这中间隔着一条巨大的鸿沟，大到能吞掉整个项目。

一个全球数据给你看这条鸿沟有多宽：Gartner 预测，约 30% 的生成式 AI 项目会在 POC（概念验证，也就是 demo 阶段）之后被放弃。S&P Global 的调查更直接，企业在 POC 之后放弃项目的比例，从一年前的 17% 涨到了 42%。也就是说，差不多一半的项目，demo 跑得好好的，一进真实环境就死了。

为什么 demo 到上线这么难？拆开看这几条：

第一，demo 用的是干净数据，上线用的是脏数据。演示的时候，供应商喂给 AI 的是精挑细选、格式工整的样本。你公司真实的数据呢？散在十几个系统里，格式不统一，有错漏，有过期，写得乱七八糟。同一个 AI，喂干净数据答得漂亮，喂你的真实数据就开始胡说八道。

第二，demo 不碰权限，上线全是权限。演示环境里没有权限这回事，AI 想看什么数据都能看。真实环境里，谁能查什么、哪些数据不能出本地、哪个部门的东西另一个部门看不到，这一套权限体系，能把一个简单的 demo 拖成一个复杂的工程。

第三，demo 不接流程，上线必须嵌进流程。demo 是个孤岛，你单独问它、它单独答。上线意味着它要嵌进你真实的工作流，上游谁给它喂数据，下游谁接它的结果，出了错往哪报，这些 demo 里全没有。

第四，demo 不考核准确率，上线天天考核准确率。演示的时候你看的是它能做，上线之后你看的是它做对了多少。一个在 demo 里答对八成的 AI，听起来不错，但如果这八成里有一成是关键业务、错了要赔钱的，那它在生产环境里就是个定时炸弹。

我们做风控项目时，深知这条鸿沟。所以做策略部署系统，第一期根本没碰花哨的 AI，先做的是把整条流程线上化，监听邮件、结构化存储、状态机管理、自动留痕。这些不性感的基础建设做完，AI 才有地方嵌。后来才上 AI，识别出一部分邮件可以自动部署（而且还要人工 double check）。真落地是从脏活累活开始的，demo 是从漂亮画面开始的，方向正好相反。

给你一个老板能直接用的判断：看到一个让你眼前一亮的 demo，先别拍板，先问一句：从这个 demo 到真正能在我公司用，还差哪些事、要多久、谁来做。如果对方答不上来，或者轻描淡写说：很快就能上，这个坑你大概率要踩。

坑 7：业务方不参与，技术单方面硬干

第二个执行坑：项目交给技术团队（或供应商），业务部门当甩手掌柜，开个启动会就再也不露面了。

结果就是技术按自己的理解做，做出来业务一看：这不是我要的。返工，再返工。

AI 不会凭空懂你的业务，一笔交易该不该拦、一个报销合不合规、一个客户问题该怎么答，这些判断标准全在业务老员工的脑子里，没写在任何文档上。技术团队再聪明，也猜不出你这行的门道。业务方不参与，等于让一个不懂行的人替你做决定。

我接触过一个字节的 HR 系统项目，里面招聘这个模块做得动，因为业务方目标清晰；但休假模块就是做不动，因为业务方很佛系、说不清自己到底要什么。业务方没目标、不参与，AI 项目就是空中楼阁。

坑 8：流程都没理清，就急着上 AI

这个坑和上一个是连体的。很多企业的业务流程，本身就是乱的，没人说得清完整流程，全靠几个老员工的默契在跑。这种状态下硬上 AI，AI 根本不知道自己该嵌在哪、按什么标准干。

我有个判断，后面会反复讲：流程不清，AI 项目几乎一定失败。这是我见过最普遍、最致命的坑之一。

这里要借一个很形象的说法，AI 是放大镜，它不创造复杂，只照亮你原有的混乱。你公司流程里那些“凑合能跑”的乱账，平时藏着没事，AI 一进来，全给你照出来了。你以为买的是 AI 的智慧，结果被迫做了一次组织体检。

这其实是好事，但很多老板没这个心理准备，以为上 AI 是来提效的，结果发现第一步是来“做体检、补欠账”的，心态就崩了。先把流程理清，AI 才有地方落脚。流程理清这件事，本身就是项目最难、最值钱的部分。

坑 9：过早追求全自动，把人踢出去

坑 5 讲的是认知层面“以为能 100% 替代”，这个坑讲的是执行层面“过早全自动”，就算你知道不能完全替代，落地时也容易贪心，一上来就把 AI 的决策权开到最大，把人从流程里踢出去。

这里有个要命的数学：错误是会复合的。假设一个流程有五个环节，每个环节 AI 都有 90% 的准确率，听起来很高了吧？五个环节连起来，0.9 的五次方，整体准确率只剩 59%。一半多一点。你敢把这样一条全自动流水线放进生产环境吗？

Klarna 这家公司很值得参考，他们 2024 年上了 AI 客服，30 天处理 230 万次对话，相当于 700 个客服，自动化率 67%，处理时间从 11 分钟降到 2 分钟，数据漂亮得不得了。到 2025 年，Klarna 又开始重新招募人工客服，并明确提出让客户始终可以选择与人工沟通。公开表述里，AI 提供速度，人提供同理心和复杂场景处理。

正确的姿势是先降级、再放权。先让 AI 当副驾驶（做初稿、初筛、辅助分析），人当主驾驶（判断、拍板）。等跑稳了、有数据了，再逐步扩大 AI 的决策权。

坑 10：让大模型去干它本不该干的活

最后一个执行坑，偏一点技术，但老板必须懂这个判断：不是所有活都该交给 AI 大模型，有些活必须交给确定性的工具。

企业业务里，有些环节需要 AI 的泛化能力，理解、生成、推理、总结，这些 AI 擅长。但有些环节必须确定性，费用计算、权限审批、数据入库、业务规则判断，这些一步都不能错。

让大模型去算钱，是灾难。大模型的本质是生成，它的思维链不稳定，同一个问题问十次，可能给你十个略有差异的答案。这对聊天没关系，对企业生产是致命的。你能容忍你的 ERP 系统今天用算法 A、明天用算法 B 吗？不能。

我在某大厂踩过一个坑：一条策略上线前，金额单位搞错了，请求过来的是“元”，查询接口返回的是“分”，差了一百倍，直接导致误拦截。这种确定性的活，错一次就是事故。真正会用 AI 的人，会把模型、规则、数据库、人工审核组合起来用，每个环节用对工具，而不是一股脑全塞给大模型。

三、运营类的坑：东西做出来了，却没人用

很多老板以为，AI 项目最难的是“做出来”。错了。做出来和用起来是两件事，后者往往更难。这一组讲的就是上线之后才暴露的坑。

坑 11：工具做出来了，没人用

这是运营阶段最让老板心痛的坑：花了大几十万，系统做出来了，验收也通过了，结果三个月后一看后台，没几个人在用。钱打了水漂，还落下一句：AI 不行。

一个客观事实，工具没人用，基本都是卡在和工具本身无关的原因上。

第一，入口不顺。员工要用这个 AI，得先打开一个新网址、再登录一个新账号、再切换到另一个界面。每多一步，就劝退一批人。员工本来手头活就多，你让他为了用 AI 多点五次鼠标，他宁可不用。真正能活下来的工具，都是嵌在他每天已经在用的地方，比如直接在企业微信里、在他天天打开的系统里。

第二，场景不真。这个 AI 解决的不是员工真正头疼的问题。是老板拍脑袋觉得：这个该上 AI，但一线员工根本不在这上面花时间。解决了一个假问题，自然没人用。

第三，知识质量差，答得不靠谱。这个最致命，员工第一次问，AI 答得驴唇不对马嘴；第二次问，还是不靠谱。一次糟糕的体验，就足以摧毁信任，后面你再怎么优化，他也不会回来了。我见过太多知识库项目栽在这，文档一股脑塞进去，没清洗、没治理，AI 答非所问，员工问一次就再也不问了。

第四，习惯没变。人是有路径依赖的，原来怎么做事，他还想怎么干。没有外力推一把，他不会主动改变工作习惯，哪怕新方式确实更好。

第五，没人推、没人养。上线就等于结束了，没有人负责推广，没有人收集反馈，没有人持续优化，没有内部案例，没有答疑入口。一个没人管的系统，自然慢慢死掉。

我把第三和第五条单独强调一下，因为它们最容易被老板忽略。

关于第三条，我自己做风控 Co-Pilot 时，把防幻觉当成硬约束，知识库没覆盖的问题，AI 不许瞎猜，直接转人工。这是设计，不是 AI 无能。宁可它老实说：我不知道、转人工，也不能让它编一个看起来对、其实错的答案。因为编错一次，员工就再也不信它了。

关于第五条，我在某大厂还踩过一个相关的坑：有一次我推大模型探索，做了几个月没有阶段性结果，季度末不被认可。根因就是探索型的东西难验证、又没人持续盯着喂养。我后来的做法是，前期先做能看得见的基础搭建（把流程线上化、建表、建知识库），让这些实实在在跑起来、有人用、能被看见。

给老板的判断：别问“这个 AI 做出来了吗”，要问“这个 AI 上个月有多少人在用、用了多少次、解决多少业务问题”。第一个问题是建设型成果，第二个才是经营型成果，你要的是后者。

坑 12：没有评测集，凭感觉上线

这个坑很技术，但后果老板承担。

什么叫评测集？就是一套标准答案题库，你拿一批业务问题、配上正确答案，定期让 AI 去答，看它答对多少。没有评测集，你就完全不知道这个 AI 到底准不准、这次优化是变好了还是变差了。

没有评测集 = 凭感觉上线。今天觉得它答得挺好，就上了；明天它答错了，你也不知道是一直就错、还是刚刚改坏的。整个项目处在“盲人摸象”的状态。我自己做项目，会明确准确率指标（比如风控拦截类必须达到很高准确率），没有量化标准就不让上线。

坑 13：没有 badcase 机制，错了也不知道、不改

评测集是上线前的体检，badcase 机制是上线后的复诊。

Badcase 就是“答错的案例”，一个负责任的 AI 系统，必须有一套机制：把每天答错的、答得不好的案例收集起来，分析为什么错，然后改进。没有 badcase 机制的项目，等于上线就不管了，错误一直在发生，但没人收集、没人分析、没人修。

没有 badcase 库的项目，就是凭感觉在跑。系统会随着使用不断暴露新问题，你得有个口子去接这些问题、消化它们。没有这个口子，系统只会越用越不被信任。

坑 14：上线没人运营，项目自然死

这条是坑 11 的延伸，但值得单独列，因为太多老板以为“上线 = 完成”。

AI 系统是活的，需要人养。知识会过期，要更新；新问题会冒出来，要补充；员工会有新需求，要响应。这些事必须有一个明确的人负责，这个人就是系统的养护工。

我反复跟客户强调一个问题：这个系统上线后，谁来养它？如果答案是想没想过，大家一起看着，那就等于没人养，这个项目活不过半年。一个有人持续养的系统，会越用越好；一个没人养的系统，会越用越废。

坑 15：把生产环境当成了玩具环境

最后一个运营坑，是一个认知错位。

同一个 AI 工具，普通人用是更好用的聊天机器人，企业员工用是重构业务流程的生产力。差别在哪？差在有没有生产环境，真实的业务流、真实的数据、真实的系统、真实的上下游、真实的交付压力、真实的反馈闭环。

没有生产环境，AI 永远停在玩具层。很多企业上了 AI，但只是让员工拿它当个高级搜索框，问一问、聊一聊，没有嵌进任何真实的交付链条。这种用法，投入再多也产生不了生意上的价值，因为它根本没进生产。

这条解释了为什么那么多企业嘴上说在用 AI，但落不了地，他们把 AI 当玩具用，自然得到玩具的结果。

四、合作类的坑：被供应商、被合作方坑

前面三组，坑的主体是你自己。这一组，坑的主体是别人，你找的供应商、你的合作方。这一组的钱，亏得最冤。

坑 16：找错合作方，能力对不上需求

第一个合作坑：找了一个能力和你需求对不上的合作方。

最常见的是两种错配。

一种是，你的项目核心需要某个稀缺能力，但合作方根本不具备。我们评估项目有个硬指标：先算这个项目需要什么核心能力，如果招这种人市场年薪要 60 万以上，而项目总预算不到 100 万，这项目基本必死。我接触过一个医药 AI 智能体平台的项目，聊下来就是这个问题，它需要的人才，远超项目能养得起的水平，我主动放弃了。

另一种错配是，合作方擅长的和你要的不是一回事。比如你要的是落地陪跑，找了个只会讲课的讲师型团队；或者你要的是一个能持续迭代的系统，找了个只会做一锤子买卖 demo 的外包。找合作方之前，先想清楚你这个项目最核心的那块活是什么，再看对方擅不擅长这块。别被对方的热情和漂亮 PPT 带跑。

坑 17：外包不可控，出了事自己背锅

第二个合作坑，我用自己第一个成交项目的血泪教训来讲。

那是一个 TK 鞋子带货的项目，当时的链路是这样的：合伙人在前面销售，过度承诺了客户；

开发外包给了一个不太靠谱的朋友；交付过程我自己都没法体验、没法验收；最后出了问题，中间层背锅，信任全面破裂。

这是一条必死的合作链路：销售过度承诺 → 开发外包失控 → 自己背锅。

外包本身可以做，失控才是坑。如果你要外包，至少守住三条：第一，交付的命脉不能外包给不可控的人；第二，外包前自己得先跑通技术验证（确认这事真能做），并且定好明确的验收标准（每个里程碑交付什么）；第三，固定的同步节奏（比如每周对齐一次），哪怕进度是零，也得让你知道。

我还总结过一条：信息透明度比交付速度更影响满意度。交付期间最怕的不是慢，是客户被蒙在鼓里，催一下、安抚一下、问开发、开发说在做、两天没动静、再催，这个死循环比延期本身更伤信任。

坑 18：没有验收标准，扯皮没完

第三个合作坑：合作开始时没把“做到什么程度算合格”定清楚，结果交付时双方各说各话，扯皮没完。

我见过一个反面案例（行业里很常见的模式）：客户问“这个能做吗”，对方答“能做的”，然后直接报价、签合同。等到交付，发现边界完全没确认过，客户以为能做 A，对方以为只做 B，谁也说服不了谁，项目卡死。

验收标准必须在动手之前白纸黑字写清楚。我现在接任何项目，都按一个节奏走：先接住需求（这个方向能做）→ 加边界（但能不能稳定跑通，取决于你们的数据、业务逻辑是否成型，得先拿样本验证）→ 讲路径（先拿 5-10 个样本跑通验证 → 再设计工作流 → 确认后再批量交付）。中间那个“加边界”的动作，是最关键、也最容易被跳过的。

对老板来说，判断很简单：签合同前，问一句“交付时，我们拿什么来判断这个项目算不算做成了？”如果答不具体的、可检验的标准，这个合作就埋了扯皮的雷。

坑 19：被供应商的话术忽悠，踩了承诺的坑

第四个合作坑：供应商为了签单，说了一堆它兑现不了的承诺，你信了。

先放一个反常识判断：一个敢主动跟你说“什么做不到”的供应商，比一个满口“都能做”的供应商更可信。我自己跟客户沟通，被问到不确定的事，会直接说：这个我需要研究一下再给结论，而不是拍脑袋说：应该可以。被问到能力上限，我会说：人能做的 AI 能做其中很大一部分，人做不到的 AI 也做不到。这种坦诚，专业的客户反而更买账。

反过来，有几种话术你要警惕（下一章还会专门讲）：承诺“全自动、无人值守”的、避而不谈数据准备的、不肯定验收标准的、报价含糊其辞的，这些都是危险信号。

坑 20：把核心能力完全外包出去

最后一个坑，也是最容易被忽略的：在一个快速变化的领域，把你的核心能力完全交给第三

方，本身就是风险。

如果 AI 要成为你公司的核心竞争力，比如你是靠它做差异化、靠它建壁垒，那这块能力你不能完全握在别人手里。今天供应商配合，明天它涨价、它跑路、它被收购，你的核心能力就悬空了。

这不等于什么都要自己做。判断的逻辑是：核心的、要自主可控的、涉密的，倾向自建；非核心的、标准化的，可以采购或外包验证。你得分清楚，哪些是你的命根子、不能假手于人，哪些是可以买、可以外包的标准件。把命根子外包出去，等于把脖子伸给别人。

这二十个坑，有一个共同点

回头看这二十个坑，你会发现一件事：没有一个坑，是 AI 不够聪明造成的。

认知类的坑，是老板对 AI 的理解错了；

执行类的坑，是流程、数据、人没准备好；

运营类的坑，是没人推、没人养、没机制；

合作类的坑，是找错人、没边界、没控制力。

模型很强，强到已经超过了大多数企业能消化的水平。真正卡住企业的，从来不是模型，是这些 AI 之外的承接能力，数据接不接得住、流程理不理得清、组织推不推得动、责任分不分得明。

这一点不是我一个人的判断。RAND 公司 2024 年访谈了几十位一线数据科学家，总结出 AI 项目失败的五大根因，排在第一的是业务领导层对 AI 的误解，五条里几乎没有一条是纯技术问题。BCG 提出一个 10-20-70 原则：算法只占 10%，技术和数据占 20%，人和流程占 70%。全球顶级机构的结论，和我在一线趟出来的判断，是同一个。

所以下一章，我把这二十个坑翻过来，坑告诉你什么不能做，经验告诉你应该怎么做。二十个坑的正面镜像，就是二十条能让项目活下来的经验。

一句话

这章收束成一句话：企业 AI 项目不是死在模型不够强，是死在你没把它接住，数据、流程、组织、责任，这四样接不住，再强的 AI 进来都白搭。

第 11 章 项目里总结出的 20 条经验

把坑翻过来看

上一章讲了二十个坑，每个坑你看完可能都有点发怵。这一章我把它们翻过来，坑是“这么干”，经验是“该这么干”。二十条经验，对应一个 AI 项目的三个阶段：

启动阶段：怎么开好头，别一开始就跑偏（经验 1-7）

执行阶段：怎么稳着走，别中途翻车（经验 8-14）

运营阶段：怎么让项目活下来，别死在上线后（经验 15-20）

这二十条，每一条后面都压着我自己做过的项目、或者服务过的客户。

一、启动阶段：怎么开好头

绝大多数 AI 项目的成败，在启动阶段就定了七八分。开头错了，后面再努力都是在错误的方向上加速。

经验 1：先对齐问题，再讨论方案

这是启动阶段第一条，也是最容易被跳过的一条。

大多数项目一上来就讨论怎么做，用什么模型、搭几个智能体、接不接知识库。但真正的风险不是技术，是问题没被对齐。同一个需求，不同的人理解完全不同：业务按场景分，技术按架构分，老板按组织分。所有人都在热火朝天地讨论答案，却没有一个人先把问题对齐。

我的经验是：启动前花时间把问题定义清楚、把边界对齐，比后面返工省十倍的成本。这个动作看起来慢，实际上是最快的。

经验 2：先诊断，再动手

启动 AI 项目，正确的顺序是先做一次诊断，系统地看一遍你的业务，找出哪里最该上 AI，而不是凭老板的直觉拍一个场景就上。

为什么？因为你企业最该做 AI 的地方，往往不是你最先想到的那个。老板凭感觉想到的，常常是最显眼的，但不一定是最值钱、最容易落地的。

我们服务客户时，诊断的标准顺序是：先看组织准备度（数据有没有、权限清不清、流程标不标准、责任明不明），再看业务场景选型，最后才是模型和工具选型。这个顺序倒过来就必

死，很多企业一上来就纠结用哪个模型、买哪个工具，方向从根上就错了。

经验 3：先给问题标价，别从“AI 能干什么”出发

诊断的时候，有个视角的转换至关重要。

错误的视角是从 AI 能干什么出发，这是工具思维，你会找出一堆鸡毛蒜皮的小场景。正确的视角是从哪个问题最值钱出发，具体动作是：给每个问题标个价。

会议纪要自动化，一年值多少钱？销售平均成交从 200 单提到 300 单，值多少钱？发货计划优化，能省多少现金流？标完价，你立刻就清楚了，别把一个能改变业务结果的 AI，浪费在一个一年只省几千块的小问题上。

标价这个动作的真正价值在于：它会让你看见那些最值钱、但一直没被看见的问题。那些每天耗掉你最贵的人、最多的时间、最依赖老员工经验的环节，往往才是 AI 该进去的地方，但它们因为“一直就这样”，反而被忽略了。

经验 4：先小场景，别贪大

这条是坑 4（贪大求全）的正面镜像。

值得做的第一个 AI 项目，一定是小而痛的。小，是指范围窄、周期短、能快速跑出结果；痛，是指它解决的是员工真正头疼、真正高频的问题。

最好的启动场景，我给你几个现成的：周报自动生成、会议纪要、客户反馈整理、跨系统状态同步、内部汇报材料。这些活有个共同点，单看不复杂，但长期占用大家的注意力，做掉之后，人的注意力立刻被释放出来。它们风险低、见效快、客户感知强，是零失败启动包。

我们跟客户聊，从不是一上来谈“AI 重塑你的公司”。聊的是：先把每周这些重复的破事清掉，再谈下一步。反而是这种朴实的开场，更容易签单、更容易出成果。智能体不是第一步，是跑通几个小场景之后的第三步。

经验 5：业务方必须深度参与

这条是坑 7 的正面镜像。

AI 落地最关键、也最难的活，是把业务老员工脑子里的隐性经验挖出来、变成 AI 能用的东西。这件事，纯技术的人干不了（他不懂业务），光靠业务的人也干不了（他不知道怎么变成 AI 能用的东西）。必须是业务方深度参与，有人坐下来一笔一笔把经验抠出来。

但这里有个难点：业务方有经验，但他拆不出 SOP，而且他不懂 AI 的边界。你直接问他：你的判断流程是什么，他答不上来，这是干了多年练出来的手感，只可意会。

我有个破局的办法，叫选择题驱动法。不要给业务方出问答题（你能写下你的流程吗），要给他做选择题。具体做法是：先用 AI 梳理出一版稻草人框架（一个粗糙的初版流程），拿给业务方看，让他判断，哪些对、哪些错、哪些缺。他改一版，你再梳理一版新的选择题，来回迭代几轮。

我们在某大厂搭风控的诈骗标签体系，就是这么搭起来的。每个标签都明确定义：什么属于、什么不属于、正例是什么、反例是什么、边界情况怎么算。业务方填空填不出来，但选择题他一看就知道对错。这是把隐性经验显性化最有效的办法。

经验 6：先摸清预算和决策链，别空转

这条偏商务，但直接决定项目能不能真正启动，老板自己也要懂。

刚开始创业时，我吃过几次亏，就是聊得热火朝天，但一直没问两件事：这个项目有多少预算、最后谁拍板。结果方案做了一大堆，最后发现预算差了一个数量级，或者跟我聊的人根本做不了主。

我有个亲身教训：某地一家银行支行的负责人找我做 AI 培训，我按市场化客户的逻辑按普通市场价报了一个偏低的价格。聊了几分钟，对方透露内部预算其实高不少。报价已经发出去了，相当于一开口就把报价空间压没了。体制内客户的预算是流程预留的额度，你不问，他不会主动说；你先报了，他就按你报的批。

这条对老板的启发是：你自己内部立项时，先把预算和拍板人定清楚，再去找供应商。否则你找来的方案，要么贵得离谱、要么便宜得不靠谱，因为对方根本不知道你能花多少、谁说了算。

经验 7：承认边界，是专业不是无能

启动阶段最后一条，关于心态。无论是你内部团队，还是外部供应商，一个敢说这个 AI 做不到的人，比一个满口都能做的人靠谱得多。

客户问我：AI 能不能解决这个，我的标准回答是：人能做的，AI 能做其中很大一部分；人做不到的，AI 也做不到。客户问我某个具体能力，我会先划清楚范围：在这个范围内，AI 能做到这个程度；超出这个范围，AI 做不到。制造业的老板对“AI 能解决一切”这种话早就免疫了，你越承认边界，他越觉得你专业、能合作。

反过来，那种什么都拍胸脯的，要么是不懂、要么是想骗单，迟早翻车。启动一个 AI 项目，要找愿意跟你划边界的人，不是会跟你画大饼的人。

二、执行阶段：怎么稳着走

启动对了，到了执行阶段，核心就一个字：稳。AI 项目最容易在执行阶段贪快、贪全、贪自动，结果翻车。这一组讲怎么稳着走。

经验 8：先 Copilot，后全自动

这条是坑 9 的正面镜像，也是我最想让老板记住的执行原则之一。

别一上来追全自动，先让 AI 当副驾驶。Copilot 模式是：AI 做初稿、做初筛、做辅助分析，人

负责判断和拍板。等这个模式跑稳了、积累了数据、看清了 AI 在哪些地方靠谱在哪些地方不靠谱，再逐步把决策权交给它。

在某大厂做风控时，我的组长选了全自动替代的路线，团队投入了一段时间，失败了。我提的是大部分提效加少部分人工兜底的路线，跑通了。同一个业务、同一拨人，路线不同，结果天壤之别。

先降级不是不创新、不是保守，是给 AI 找对它该待的位置，让整个项目可控。全自动看着先进，但企业场景里风险太大；先 Copilot，你既拿到了明显的效率提升，又守住了风险底线。

有研究分析大量落地项目发现，AI 自主处理八成、人只盯住例外的模式，中位提效 71%，反而比每一步都要人审批的模式（30%）高得多。先 Copilot 不是为了求稳而牺牲效率，给 AI 一个合适的副驾驶位，效率和安全是能同时拿到的。

经验 9：先把流程理清，再上智能体

这条是坑 8 的正面镜像。

顺序永远是：先流程，再 AI，最后才是智能体。流程不清的时候，任何 AI 上去都是乱的，它不知道该嵌在哪、按什么标准干。

我们之前做策略部署系统，第一期叫“策略线上化”，根本没碰 AI。做的事很笨：监听邮件、把策略信息结构化存储、用状态机管理流转、自动路由和留痕。原来流程乱到什么程度？策略常常拖很久没人部署、历史策略散在个人邮箱里没法回溯、相似标题导致回错邮件。

把这套流程整顺了，AI 才有地方嵌。后来上了 AI，能识别出一部分邮件需求可以自动部署，能做需求和系统配置的一致性校验。流程线上化这一步不性感，但它是后面所有 AI 应用的地基。跳过它直接上智能体，就是在沙子上盖楼。

经验 10：让 AI 造工具，而不是让 AI 直接干活

大模型的本质是生成，它的思维链不稳定，同一个问题问十次，可能给你十个略有差异的答案。这对企业生产是致命的。你不可能让一个每次结果都飘的东西，去支撑你天天要跑的业务。

那怎么办？正确的工作流是：AI → 软件 → 结果，而不是直接：AI → 结果。让 AI 这个艺术家去造一台稳定的机床（一个软件、一个工作流、一个 Agent），然后用这台机床去批量产出确定的结果。

举个例子。可以让 AI 写一个自动分析录音、给改进建议的软件，然后团队每天用这个软件；不要让 AI 每天直接去听销售录音、临场给建议，那样每天结果都会飘。前者结果不可控，后者稳定、可复用。

这条和我后面会讲的 token 消耗那个判断是闭环的，真正深度用 AI 的人，正是在用 AI 造工具、批量解决问题，而不是手动一问一答。一问一答是玩具，造工具才是生产力。

经验 11：泛化的活给 AI，确定的活给规则

这条是坑 10 的正面镜像。

企业业务里，要分清两种活：需要泛化能力的（理解、生成、推理、总结），交给 AI；需要确定性的（算钱、审批、入库、规则判断），交给规则、数据库、人工审核。混着用必出问题，让大模型算钱是灾难，让规则引擎理解需求是死板。

我们做 B 端方案，会强制加一节泛化 vs 确定性分工表，把每个环节列出来，标清楚这个环节用 AI 还是用规则，理由是什么。没有这一节的方案，很难交付。

这不只是技术讲究，对老板也是个判断工具：当有人给你一个方案，里面所有环节都让 AI 来扛，连算钱、审批这种一步都不能错的活也交给 AI，你就该警惕了。真正会用 AI 的人，是把 AI、规则、数据库、人工像搭积木一样组合起来，每块用在它该用的地方。

经验 12：用阶段性成果换基础建设的时间

这条讲一个执行阶段的现实难题：AI 落地需要海量底层建设（数据清洗、标签体系、字段治理、SOP、知识库、工程链路），但老板只看 demo、看效果、看页面。底层做得再好，不被识别，项目就难推。

我们在大厂时的解法是三步：

第一，启动时就给老板做预期管理，明确说：前期在打基础、会踩坑，看不到漂亮结果是正常的；

第二，用 AI 快速搭一个能看的阶段性页面，用来汇报，证明项目没在空转；

第三，用这个看得见的小成果去换看不见的基础建设的时间。

这背后是一个对老板很重要的概念，建设型成果和经营型成果的错位。项目方做的是建设型成果（搭了什么能力），老板看的是经营型成果（带来了什么经营结果）。一期搭的能力，和老板心里的终局预期之间，必然有错位。如果一开始不对齐：什么时点交什么成果，老板会用终局的标准去看一期的结果，然后否定整个项目。

所以这条经验对双方都成立：做项目的人要主动管理预期、定义阶段性目标；当老板的，要理解一期看到的是地基不是大楼，别用看大楼的眼光去验收地基。

经验 13：把成果做得“看得见”，别让基础建设憋死

接着上一条，我们单独把成果可见性拎出来讲，因为在这上面摔过一跤。

在某大厂有一次推大模型探索，做了一段时间，没有阶段性结果，季度末不被认可。根因是：探索型的成果难验证、成本又高，憋在那里没有任何看得见的东西冒出来。

吸取这个教训，我后来做项目的原则是：前期一定先做能看得见、能被用上的基础搭建，把流程线上化、把表建起来、把知识库搭起来，让这些实实在在跑、有人用、能被领导看见。看得见的基础建设，既是项目的地基，也是项目的护身符。完全埋头做看不见的探索，再

有价值也容易在中途被砍掉。

经验 14：处置的严重程度，决定准确率的要求

这条偏专业，但它背后的判断逻辑，任何行业的 AI 项目都用得上。

我们做风控时有个原则：一个动作越重，对 AI 准确率的要求就越高。拦截一笔交易（最重）要求很高的准确率，加一道验证（次重）可以低一点，给用户一个提示（最轻）容错最高。因为动作越重，错了的代价越大。

把这条翻译成给老板的通用判断：AI 在你业务里干的活，后果越严重，你越要把准确率的标尺定高、越要留人工兜底。给客户推荐个商品、整理个会议纪要，错了无所谓，可以放手让 AI 干；但只要涉及钱、涉及合规、涉及对客户的承诺，准确率标尺必须拉高，必须有人复核。不要用同一个准确率标准对待所有场景，轻活放手，重活把关。

三、运营阶段：怎么让项目活下来

很多项目死在上线那一刻，以为上线就是终点，结果上线才是真正的开始。这一组讲，怎么让一个 AI 项目在你公司里长期活下去。

经验 15：知识库要持续运营，不是一次性工程

这条是坑 14（没人运营）的正面镜像，尤其针对知识库类项目。

知识库不是搭完就完事的工程，是一个需要持续养的活物。知识会过期，要更新；新问题会冒出来，要补充；员工会有新需求，要响应。这些事必须有一个明确的人负责。

而且知识库最大的难点根本不在技术，在治理。直接把公司文档一股脑丢进去，做不出能用的知识库，因为企业文档是按部门、按模块组织的，员工是按问题来提问的，两者的逻辑天然不匹配。文档必须经过清洗、拆分、重组、标边界四步加工，才能变成 AI 能用的知识。

有个反直觉的做法：做知识库，别从我有多少文档出发，要从员工最常被问到的 100 个问题出发，反向去构造知识。先收集的高频问题，再去配对应的知识，这样做出来的知识库才答得到点上。

经验 16：准确性要靠评测集，不靠感觉

这条是坑 12 的正面镜像。

准确性是设计出来的，不是模型自带的。你想知道 AI 准不准、这次优化是变好还是变差，唯一靠谱的办法是有一套评测集，一批问题配上标准答案，定期让 AI 去答，量化它的准确率。

我做项目会明确准确率指标，没有量化标准就不让上线。有了评测集，你的每一次优化都有据可依：改完一跑，分数涨了就是真的好了，掉了就是改坏了。没有评测集，你就是在黑暗里瞎走，今天觉得好就上、明天出问题也不知道为什么。

经验 17：要有 badcase 机制，让系统越用越好

这条是坑 13 的正面镜像。

评测集管上线前，badcase 机制管上线后。一个负责任的 AI 系统，必须有一套机制：把每天答错的、答不好的案例收集起来，分析原因，持续改进。

我们做风控 Co-Pilot 时，整个闭环设计里专门有 badcase 治理这一环，感知、分析、推荐、验证、上线，每个环节都可视化、可人工介入，错的案例能被接住、被消化。有 badcase 机制的系统，会随着使用越用越准；没有的，会随着新问题不断暴露越用越不被信任。

给老板的判断：问供应商一句“系统答错了，你们怎么收集、怎么改？”答不上来具体机制的，这个项目上线后就是个野生状态，自生自灭。

经验 18：防幻觉，宁可转人工也不许瞎编

这条是运营阶段的硬约束，也是保住用户信任的命门。

知识库没覆盖的问题，AI 不许瞎猜，直接转人工。我把这条当成做 Co-Pilot 的硬规则，禁止模型自行推断、补全、猜测。

为什么这么狠？因为一次糟糕的体验，就足以摧毁信任。员工问一个问题，AI 编了个看起来对、其实错的答案，员工被坑一次，就再也不信它了，后面你再怎么优化他也不回来。与其让它编一个错的，不如让它老实说：这个我不知道，帮你转人工。

需要注意点，AI 说：我不知道、转人工，不是失败，是设计。一个知道自己边界、该转人工就转人工的 AI，比一个什么都敢答、但你不敢全信的 AI，可靠得多。这和上一章讲的 80/20 是一个道理，守住该守的边界，才是企业级 AI 的成熟姿态。

经验 19：上线后要有人“养”，明确到具体的人

这条是坑 11、坑 14 的正面镜像，也是运营阶段最该被钉死的一条。

每个 AI 系统，上线时必须明确一个问题：谁来养它？不是大家一起看着，是具体到一个人，他负责更新知识、收集反馈、持续优化、答疑推广。

我跟客户反复确认这个问题。如果答案是“没想过”“到时候再说”，我会直接说：那这个项目活不过半年。一个有人持续养的系统，会越用越好；一个没人养的系统，会越用越废，最后静静死掉。养系统的人，是 AI 项目里最被低估、但最不可或缺的角色。

顺带说一个让系统活下来的运营打法：火种用户（先培养一批爱用、会用的种子员工）、内部案例（把用得好的案例传播出去）、激励挂钩（让用 AI 和考核、评优有关系）、答疑入口（员工有问题随时能找到人）。这一套机制，比工具本身更决定一个 AI 系统的生死。

经验 20：员工要的不是会用工具，是会和 AI 协作

最后一条，落到人身上。

AI 时代，企业里真正有价值的人，是会提问、会判断、会拆任务、会和 AI 协作的人，不是单纯会用某个 AI 工具的人。工具会一直变，但这套协作能力是通用的、持久的。

这里有个识别真本事的硬指标，我很推崇：看一个人一个月烧多少 token。在理想内部有个观察，token 烧得最多的人，往往是那些脑子快、嘴跟不上脑子、过去吃亏在表达上、但跟机器对话零摩擦的人，不一定是传统意义上的明星员工。他们几乎都在重构自己所在领域的业务，不是在玩 AI。token 消耗量，是判断一个人有没有真正进入 AI 状态的客观信号，比简历和自我描述都准。

但有一条边界要提醒老板：别让员工把判断也外包给 AI。AI 可以帮人做初稿、做分析、做整理，但最后的判断、拍板，得是人来。如果员工连判断都丢给 AI，长此以往，人的核心能力会萎缩，组织反而变脆弱。培训的目的，是让人更会用 AI 放大自己的判断，不是让人退化成一个只会复制 AI 答案的传声筒。（这条边界，以及 AI 时代到底要什么样的人，第 24 章我会专门展开。）

二十条经验，其实是一条主线

回头看这二十条，启动、执行、运营三个阶段，看似散，其实串着一条主线：AI 项目的成败，从来不取决于模型多强，取决于你有没有把它“接住”。

启动阶段，接住的是问题，先对齐问题、先诊断、先标价、先选小场景；

执行阶段，接住的是节奏，先 Copilot、先流程、让 AI 造工具、分清泛化和确定；

运营阶段，接住的是生命力，持续运营、评测、改 badcase、有人养。

BCG 说算法只占 10%、人和流程占 70%；斯坦福跟踪 51 个成功项目之后说，技术是能跑的，难的是技术之外的一切，77% 最难啃的挑战都在变革管理、数据质量、流程重做上。我们在一线趟了这么多项目，得到的就是同一个结论：AI 这部分是现成的、是最不缺的；缺的是你把它接进业务的那套能力。

一句话

这章收束成一句话：**做成一个 AI 项目，靠的不是把 AI 选对，是把它接住，开头接住问题，中间接住节奏，最后接住生命力。这三段功夫做到了，模型用哪个都不太重要。**

第 12 章 老板的 AI 项目风险识别清单

这一章，是给你的一张工具卡

前两章我讲了二十个坑、二十条经验。但我知道，老板读这种东西，最怕的是：道理我都懂，可真到事上还是不知道怎么判断。

所以这一章不讲道理，只给你工具。三张清单，你可以直接撕下来，下次开会带着用：

1. 听供应商讲方案时，哪几句话一出口，你就该亮红灯
2. 决定让 AI 干什么之前，哪三类问题，绝对不能让 AI 拍板
3. 项目立项之前，哪几个问题，必须问到具体答案

你不用懂技术。这三张清单，靠的全是常识判断。但就是这些常识，能帮你拦下大部分会打水漂的项目。

老板的 AI 项目风险识别清单

三张工具卡：听方案、定边界、立项前，一眼看出项目靠不靠谱

1 供应商方案红灯

听方案时，亮三盏就停

- ① 全自动 / 无人值守 / 完全替代
- ② 避而不谈数据准备
- ③ 不肯谈验收标准
- ④ 报价含糊、不敢拆开
- ⑤ 只演示，不跑真实数据

边界不清，就别急着签。

2 AI 不能拍板的事

定场景时，命中就留人工兜底

- ① 需要业务判断或风险承诺
- ② 来源不清、责任不清
- ③ 知识体系本身还没理顺

要拍板、要溯源、没理顺，只能辅助。

3 立项前必问五题

拍板投钱前，必须问到具体答案

- ① 结果能不能被检验？
- ② 数据和流程准备好了吗？
- ③ 上线后谁来养它？
- ④ 出错了谁兜底？
- ⑤ 老板自己下不下场？

问不出答案，先别立项。

不懂模型也能判断：盯住边界、验收、责任和运营。
供应商敢不敢划边界；AI 该不该拍板；项目有没有人养；出错有没有人兜。

第一张清单：供应商方案里的危险信号

供应商来给你讲方案，画面漂亮、术语一堆，你听得云里雾里。别慌，你不用听懂技术，你只要竖起耳朵，听他有没有说出下面这几句话。说出来一句，亮一盏红灯；亮三盏，这个供应商基本可以放一边了。

危险信号 1：全自动、无人值守、完全替代

这是最大的一盏红灯。

一个供应商如果跟你拍胸脯说：上了我们这套，这个岗位就能彻底不要人了，全自动、无人值守，你要立刻警惕。在上一篇我们讲过，我的组长选了全自动替代的路线，团队投入了一段时间，最后失败了。全自动替代之所以几乎必死，是因为大模型的输出永远不可能 100% 准确，业务对手每周变手法，边界场景太多，没有人工兜底错误就会一直放大。

懂行的供应商，会主动跟你讲：哪部分让 AI 干、哪部分必须留人。张口就是全自动的，要么不懂、要么想签单，迟早翻车。

这一点也有研究佐证：分析大量落地项目发现，效果最好的模式，是 AI 自主处理八成、人只盯住例外那两成（这种模式的中位提效是 71%），而不是追求百分百无人值守。一个连这个基本分寸都不讲、张口就是“全自动替代”的供应商，你基本可以默认他没真正落过地。

危险信号 2：避而不谈数据准备

你问：这个项目要我们这边准备什么数据，他轻描淡写说：不用，你们现有的就够了，我们直接接进去就行。

这盏红灯，很多老板听不出来，但它要命。AI 项目大部分功夫，都在数据和业务的脏活累活上，清洗数据、治理知识、翻译流程。一个绝口不提数据准备、把这部分轻轻带过的供应商，要么没做过真项目，要么故意瞒着你后面要补的巨大工作量。

有个数据可以印证：仅 6% 的企业，数据是完全 AI 就绪的。也就是说，大部分的企业，数据都得先收拾一遍才能用。谁告诉你：不用准备、直接用，谁就是在画饼。

危险信号 3：不肯谈验收标准

你问：做成什么样算合格、怎么验收，他绕开，或者给你一些没法检验的话，比如：会很智能，效果很好，提升明显。

做成什么程度算合格，必须在动手之前白纸黑字写清楚。一个不肯定验收标准、只肯谈愿景的供应商，是在给后面的扯皮埋雷。等交付时，你以为能做到 A，他说只承诺了 B，谁也说服不了谁，项目卡死。

懂行的供应商，会主动跟你列清楚：上线后准确率达到多少、覆盖多少场景、响应多快、答错了怎么处理。能给你具体、可检验数字的，才是真做过的。

危险信号 4：报价含糊、不敢拆开

你问：这个多少钱、钱花在哪儿，他报一个笼统的总价，不肯拆开，或者一会儿一个价。

报价含糊，背后往往是项目边界含糊。边界没定清楚，价格自然算不准。真正想清楚怎么干的供应商，能把钱拆给你看，哪部分是诊断、哪部分是开发、哪部分是数据治理、哪部分是后续

运营维护。拆不开的报价，意味着他自己也没想清楚要干什么。

这里多说一句反过来的判断：一个敢主动跟你说“什么我做不到”“哪部分不能保证”的供应商，比满口都能做的可信得多。我自己跟客户沟通，被问到不确定的事会直接说：这个我得研究一下再给结论，被问能力上限会说：人能做的 AI 能做其中很大一部分，人做不到的 AI 也做不到。主动划边界的，是把你当长期合作伙伴；什么都拍胸脯的，是把你当一锤子买卖。

危险信号 5：只演示、不让你碰真实数据

供应商的 demo 行云流水，但你提出：能不能拿我们自己的真实数据跑一遍看看，他开始找各种理由推脱。

这盏红灯，直指上一篇讲的最贵的坑，把 demo 当上线。demo 用的是精挑细选的干净数据，你的真实数据是散在十几个系统里、格式不统一、有错有漏的脏数据。同一个 AI，喂干净数据答得漂亮，喂你的脏数据可能就开始胡说。

一个对自己东西有信心的供应商，会愿意（在脱敏和安全的前提下）用你的小样本真实数据跑一段给你看。一味只让你看预设好的 demo、不敢碰你真实数据的，你要默认：从这个 demo 到真能用，中间还隔着一条能吞掉整个项目的鸿沟。

第一张清单·一句话记法：全自动、不谈数据、不谈验收、报价含糊、不敢用你真数据，五盏红灯，亮三盏就停。

第二张清单：三类不该让 AI 拍板的问题

第二张清单，解决一个老板最容易犯的错，把不该交给 AI 的事，交给了 AI 拍板。

不是所有问题都适合让 AI 直接回答、直接做决定。有三类问题，AI 只能辅助，绝对不能让它拍板。上线前，你把要让 AI 干的事，拿这三条过一遍，命中任何一条，就必须给 AI 配上人工兜底。

第一类：需要业务判断、或者要做风险承诺的

凡是涉及判断和承诺的，AI 不能拍板。

比如：这笔交易该不该放行、这个客户能不能给授信、这个方案合不合规、这单赔不赔。这些问题的背后，是真金白银的风险，是要有人负责、有人兜底的承诺。AI 可以帮你把材料理清楚、把选项列出来、把初步判断给出来，但最后那一下拍板，必须是人。

我做风控时守的就是这条：高频、低风险的判断，放手让 AI 干；但涉及拦截一笔交易、涉及资金安全的，准确率标尺必须拉到很高，而且关键的、拿不准的，一律转人工。动作越重、后果越严重，越不能让 AI 独自拍板。

第二类：来源不清、责任不清的问题

凡是答案从哪来说不清、答错了谁负责定不下来的，不该设成让 AI 直接回答的场景。

举个例子：员工问 AI 一个公司政策，AI 答了，但这个答案是从哪份文件来的、是不是最新版、答错了导致损失算谁的，如果这些都说不清，那这个问题就不该让 AI 直接答。因为一旦它给了一个看起来权威、其实过期或错误的答案，员工照着做出了事，这个责任的窟窿没人补。

正确的做法是：要么让 AI 答的时候必须带上来源出处（让人能核对），要么这类问题干脆不让 AI 直接答、转给有权威、能负责的人。来源不清、责任不清，就是不能让 AI 单独上的场景。

第三类：知识体系本身还没理顺的

凡是公司内部对这件事本身就没有统一说法、没有标准答案的，别指望 AI 能答好。

这个最容易被忽略。很多老板觉得：我们文档一大堆，搭个 AI 问答肯定行。但如果你这堆文档本身就互相矛盾、版本混乱、根本没人理顺过，AI 不会帮你理顺，它只会把这堆混乱原样、甚至放大地吐给员工。

AI 是放大镜，不是整理术。它照亮你原有的混乱，但不替你消除混乱。知识体系没理顺就上 AI 问答，结果一定是答得乱七八糟，员工问一次就不再用了，一次糟糕的体验，就足以摧毁信任。

这三类问题，我用一句话给你串起来：需要人来负责的、需要人来追溯的、需要人先理顺的，都别让 AI 替你拍板。AI 干它擅长的，理解、整理、生成、初判；把判断、追责、定标准这些事，留给人。

第二张清单·一句话记法：要担责的、要溯源的、还没理顺的，这三类，AI 只能打下手，不能当家。

第三张清单：项目立项前必问的几个问题

第三张清单，是立项前的最后一道关。这几个问题，你在拍板投钱之前，必须问到具体的答案。问不出具体答案的，说明这个项目还没准备好，先别上。

必问 1：这个场景的结果，能不能被检验？

这是最重要的一问。

一个 AI 项目值不值得做，第一看它的产出能不能被检验，能不能判断它做得对不对、好不好。能检验，你才有评测集、才有改进方向、才能持续优化；不能检验，你就完全不知道它到底在帮你还是在坑你。

凡是结果模糊、好坏说不清的场景，先别上 AI，因为你根本没法管它。能被检验，是一个 AI 场景值得做的前提。

必问 2：数据和流程，准备好了吗？

接着问：这个场景要用的数据，齐不齐、干不干净、能不能调？这个场景的业务流程，理顺了没有、有没有标准？

我前面反复讲，大部分功夫在数据和流程上，流程不清、数据乱，AI 项目几乎一定失败。立项前如果这两样还没谱，那你立的不是一个 AI 项目，是一个先补数据流程的债、再上 AI 的项目，时间和预算都得按这个重新算。

别让供应商替你回避这个问题。数据流程没准备好就硬上，等于在沙子上盖楼。

必问 3：谁来养它？

这一问，能筛掉一大半注定要死的项目。

一个 AI 系统上线只是开始，它需要人持续地养，更新知识、收集反馈、改 badcase、答疑推广。立项时就要定死：这个系统上线后，具体哪个人负责养它？

如果答案是：大家一起看着、到时候再说，那就等于没人养。我跟客户聊到这一步，如果对方答不出具体的人，我会直接说：这个项目活不过半年。没有养护人的 AI 系统，会越用越废，最后静静死掉。

必问 4：出错了，谁兜底？

最后一问，关于责任。

AI 一定会出错，这不是会不会的问题，是迟早的问题。所以立项时就得想清楚：它出错了，谁来发现、谁来兜底、谁来负责？

AI 生成的内容，谁来审？自动流程跑错了，谁来兜？这个责任边界，必须在试试看阶段就划清楚，否则等真出了事，部门之间互相甩锅，最后一句“AI 不行”把整个项目埋了。能把“出错谁兜底”想清楚的项目，才是从玩票进入长期用的项目。

必问 5：这个项目，老板自己下不下场？

补一个问题，AI 项目不是 IT 一个部门能扛起来的事，它牵涉流程、数据、权限、考核、责任，这些只有一把手压得动。如果一个 AI 项目，从头到尾老板都不下场，全交给某个部门或某个供应商，这个项目大概率推不动。

我们服务客户有个判断：项目方案展示时，如果老板（或者真正能拍板的高管）连场都不愿意到，推说：太忙、让下面看着办，这本身就是个信号，这个项目在公司里的优先级不够高，启动了也容易半途而废。一把手不下场，AI 项目大概率会死。

真正实现了全组织转型的项目，背后的老板，无一例外都做到了把 AI 写进公司 OKR、和考核奖金挂钩的程度，光批钱、偶尔过问，是远远不够的。所以这一问，你别只问：老板支不支持，要往深里问一层：老板愿不愿意把这件事，变成全公司要一起背的目标？愿意，这项目才

有戏；不愿意，再好的方案也大概率烂尾。

第三张清单·一句话记法：结果能不能验、数据流程齐不齐、谁来养、出错谁兜底、老板下不下场，五个问题问不出具体答案，先别立项。

把三张清单合起来用

这三张清单，对应一个 AI 项目的三个决策时刻：

- **听方案时**，用第一张，拦住会忽悠你的供应商；
- **定 AI 干什么时**，用第二张，别把该人扛的责任甩给 AI；
- **拍板投钱前**，用第三张，确认这个项目真的准备好了。

你会发现，这三张清单，没有一条需要你懂技术。它们考的全是常识：有没有人负责、出错谁兜底、做成什么样算数、该不该让机器替人拍板。AI 项目的风险，绝大多数不藏在技术里，藏在这些常识里。而常识，恰恰是你这个当老板的最该把关、也最有资格把关的地方。

一句话

这章收束成一句话：**判断一个 AI 项目靠不靠谱，你不用懂模型，只要盯死三件事，供应商敢不敢划边界、该不该让 AI 拍板、这项目有没有人养、出错有没有人兜。**把这几句常识问到底，比你听懂任何技术名词都管用。

第 13 章 先培训、先诊断，还是先做智能体

老板问得最多的一个问题

我做企业 AI 落地，经常被问：我们想搞 AI，到底应该从哪开始？是先找你们做个培训，还是先上一套智能体？

这个问题背后，藏着老板的焦虑：钱不能乱花，第一步走错了，后面全错。培训花了几万块，员工听完没人用，等于白花；智能体上来就做大系统，做了半年用不起来，几十万打水漂。老板要的是一个不会踩空的第一步。

我通常会这样回答：没有一个所有企业都该先做的第一步，先培训、先诊断、还是先做智能体，取决于你公司现在卡在哪。卡的地方不一样，起步方式就不一样。选错起步方式，代价不只是慢一点，还会从一开始就在解决不存在的问题。

这一章我就讲清楚三件事：这三种起步方式分别适合谁、大厂和小企业为什么路径完全不同、培训到底能解决什么解决不了什么。看完你应该能判断出，你公司这第一步该迈向哪。

三种起步方式，分别治三种病

先培训、先诊断、先试点，对应的是企业三种完全不同的状态。你要先搞清楚自己是哪种状态，才知道该选哪条路。

状态一：全公司认知不齐，先培训

有一类企业，老板自己对 AI 很热，刷了很多视频、加了很多群，天天在公司群里转发“某某公司用 AI 提效了多少，解决了什么问题”。但你到他公司里一看，中层一脸茫然，员工连豆包、DeepSeek 都没认真用过，大部分人对 AI 的理解还停在“会聊天的搜索引擎”。

这种状态，老板和员工根本不在一个频道上。老板觉得 AI 无所不能，员工觉得 AI 跟自己没关系。这时候你直接上智能体，做出来也没人会用、没人想用，因为底下的人压根不知道这东西能帮自己干嘛。

这种企业，第一步该做的是培训。但注意，培训的目的目标不止是“教会一项技能”，更是把全公司对 AI 的认知拉到同一条及格线上。让老板知道 AI 的边界在哪、别神化；让员工知道 AI 能帮自己省哪些事、别恐惧。认知拉齐了，后面做项目才推得动。

我服务过的一家头部电商平台就是这个起点。它的财务团队有上百人。找我做培训时，核心诉求其实有两层：表面是提效，底下还有一层，大厂的危机感。财务这种高度标准化的工作，如果财务自己不进化、不学会用 AI，等业务部门先学会了，反过来就能把财务这套活给接管掉，所以财务部必须走在前面。

这种情况下，先做培训是对的。因为它要解决的首先是全员认知和危机意识，不是某一个具体场景的自动化。认知这一关不过，后面什么都做不动。

状态二：想做但说不清做什么，先诊断

第二类企业更常见。老板很想做 AI，也愿意花钱，但你问他：你具体想解决什么问题，他答不上来，或者答得空洞，我想提升效率，我想让公司更智能化，我想搭个知识库。

这种状态的特征是：有预算、有意愿，但没有靶子。他知道要做点什么，但不知道做什么最值。

这种企业，第一步既不该培训，也不该做智能体，该先做**诊断**，系统过一遍业务，找出哪些地方最值得上 AI。我见过太多老板，凭感觉拍了一个场景就让人去做，做完才发现这个场景一年也省不下几个钱，投入产出根本不划算。

为什么诊断这么重要？我举一个反面的例子。我接触过一个做女装选品的客户，老板想做一个 AI 自动选品的系统。一聊才发现，他是被几个 AI 博主的视频洗脑了，觉得别人都在做、自己不做就落后。但他公司真正的问题根本不在选品，选品靠的是买手多年的眼光和供应链关系，这恰恰是 AI 最难替代的；他公司真正高频、重复、耗人的活在客服和上新文案上。如果一上来不诊断、直接照他说的做选品系统，大概率是几十万砸进去做了个用不起来的东西。

诊断的价值，就是把“我以为该做的”换成“真正最该做的”。这一步怎么做，我在下一章（怎么拆业务）和第十五章（哪些场景值得做）里会讲。这里你只要记住：说不清做什么的时候，先别动手，先诊断。

状态三：有明确高频痛点，直接试点

第三类企业，老板能非常具体地说出问题：我客服每天回的问题里，有大量是重复的。我财务每个月做费用审核要耗掉不少人力。我们 ISO 审核表填一次要花很长时间。

这种状态，问题已经清清楚楚摆在那了，既不需要再花时间统一认知，也不需要大动干戈做全面诊断。该做的是直接挑这一个最痛的点，做一个小切口的试点，先跑通、先出结果。

我们服务过的一家精密制造企业，就是这种起点。它有家区域子公司，要应付客户的 ISO 审核，人工填一张审核表要数周。问题极其明确、极其具体。这种就不用绕，直奔着这个场景做一个能自动填审核表的智能体，把那数周压下来，就是实打实的成果。

判断自己是不是这种状态，有一个简单的标准：你能不能用一句话说清：哪个岗位、在哪个环节、每天/每周浪费多少人多少时间。说得出来，你就具备直接试点的条件；说不出来，你还在前两个状态里，别急着试点。

三种状态的对号入座

把上面三种状态整理成一张老板能直接对照的表：

你公司的状态	典型表现	第一步该做
认知不齐	老板热、员工懵，大家对 AI 理解天差地别	先培训（拉齐认知）
有意愿没靶子	想做、有预算，但说不清做什么最值	先诊断（找对场景）
有明确痛点	能一句话说清哪个环节最耗人	直接试点（小切口跑通）

这三种状态可以叠加，很多企业是混合的，比如认知不齐、同时又有一个明确痛点。这种就两条线一起走：一边给关键岗位做认知培训，一边挑那个明确痛点做试点。

但有一条铁律：不管你是哪种状态，都不要把做大系统当第一步，这是下一节要专门讲的。

大厂和小企业，起步路径完全不同

讲到这，有个问题绕不开。很多中小企业老板看 AI 落地的案例，看的都是大厂的，某某银行上了几百个大模型场景、某某互联网公司一年替代了几千万工时。看完热血沸腾，回去就想照着干，结果发现根本干不动。

这里我要专门给你提个醒：大厂和中小企业，做 AI 的起步路径完全不一样。别照搬大厂打法。你拿大厂的剧本演自己的戏，一定演砸。

大厂的打法：直接切复杂场景、搭体系

我在某大厂做风控的时候，做过一个东西叫策略部署工作台。当时的状况是：风控团队要处理大量策略流转，全靠邮件推进，邮件管理乱、容易漏，常常一条策略拖很久没人部署；历史策略散在各人邮箱里没法回溯；标题相似还经常回错邮件。

这是个又复杂、又涉及多团队协同的场景。我们的做法是直接搭体系：先做策略线上化，监听邮件、结构化存储、上状态机（待部署→待复核→待审批→待回复→已完结）、自动路由、全程留痕；在这个基础上再做 AI 化，自然语言邮件需求转结构化策略信息，识别出其中一部分邮件可以自动部署。最后效果是，超时遗漏基本消除，策略流转变得全量可管理。

注意，大厂为什么能这么干？因为它有数据、有专职团队、有现成的系统可以接入。我背后有一个专门的部署团队，有完整的 ERP 和后台，有几年沉淀下来的策略数据。在这种条件下，一上来就切复杂场景、搭体系，是合理的，它扛得住这个复杂度，也养得起这套系统。

小企业的打法：现成工具、轻量落地、单点跑通

中小企业完全是另一回事。人少、数据散、没有专职的 AI 团队，老板自己可能就是产品经理兼客服兼采购。这种条件下，你让它学大厂搭体系，纯属自杀。

中小企业的正确打法是：用现成的工具，轻量落地，单点先跑通，跑通了再复制。

我前面讲的某精密制造企业 ISO 审核就是。它没有去搭一个庞大的“制造业 AI 平台”，就奔着自动填审核表这一个具体的、脏累花钱的活去做。先把这一个点跑通，证明这套东西真能省时间、真能用，再说下一步。

这样做的原因是：中小企业的资源，只够你押注一个点。你把这一个点押对了、跑通了，有了结果，老板才有信心、有依据继续投。你要是一上来就铺开做大平台，资源摊得太薄，哪个点都做不深，最后大概率是个谁都不用的半成品。

这里再补一个打法，对想做 AI 的中小企业、或者大企业里想探索新场景的团队都有用，我把它叫：另起炉灶。与其在又老又复杂的现有系统、现有团队上硬改 AI，那种改造的内部阻力极大，不如单独拉出三五个人、一个轻装小团队，直接对着一个具体业务目标去跑，给他权限调资源、试工具，快速验证。失败了成本很低，成功了就是一条新业务线的雏形。在企业里做加法容易，做减法极难；老团队、老系统、老 SOP，有时候与其修修补补地往上贴 AI，不如另起一个轻装的小团队，从零跑通一个样板，再回头反哺大部队。这比指望一个庞大的旧组织自己长出 AI 能力，要快得多、也实得多。

一句话区分

维度	大厂（头部互联网平台这类）	中小企业
数据	多、有沉淀	散、不完整
团队	有专职 AI/技术团队	没有，老板身兼数职
系统	有现成系统可接入	系统散、甚至没有
起步打法	直接切复杂场景、搭体系	现成工具、轻量落地、单点跑通
最大风险	体系搭得动但用不起来	摊太薄，哪个都做不深

老板看自己该学谁，就看这张表里你公司更像哪一列。绝大多数中小企业，都该走右边这条路。你羡慕的那些大厂案例，是用大厂的资源喂出来的，不是你能直接复制的。

为什么不能一上来就做大系统

老板心里想的：既然要花钱做 AI，那就做个气派的、全面的，一步到位，省得以后再折腾。听上去很合理，实际上是项目的头号死法。

我接触过一个做养老健康的客户，一开始聊得很好，但越聊范围越大，又想做健康咨询、又想做营销获客、又想做内部管理，恨不得用一套系统把公司所有环节都 AI 化。这种需求一旦铺开，大概率做不成。因为它没有一个可以先跑通、先验证的小切口，所有东西都纠缠在一起，任何一个环节卡住，整个项目都推不动。

大而全的系统有三个必死的理由：

第一，没法验证。系统越大，跑通它的前提条件越多，数据要全、流程要顺、各部门要配合。这些条件凑齐之前，你看不到任何成果。而 AI 项目最怕的就是长时间看不到结果，老板的耐心和预算都是有限的，拖到看不到东西，项目自己就死了。

第二，越复杂越不稳。多个环节串起来，每个环节哪怕有九成准确，几步连乘下来，整体准确率会掉得很快。系统越大、环节越多，这种不稳定性越严重。一个不稳定的大系统，没人敢真用。

第三，资源全押在一个看不见底的窟窿里。前面提过，我在大厂时就踩过类似的坑，有一个季度推大模型探索，因为是探索型的、没有阶段性的可见成果，到季度末汇报时不被认可，明明投入了，但拿不出能验证的东西。这件事给我的思考是：任何项目，都要先做出能被看见、能被验证的阶段性成果，哪怕它很小。先搭一个能跑的小东西，后面的人才会有地基往上接。

所以正确的姿势永远是：小切口跑通，再往外扩。先做一个能在数周内看到结果的小场景，跑通了、老板看到价值了，再用这个成果换取下一阶段的资源和信任，这是企业 AI 唯一能稳着活下来的方式。

培训能解决什么，解决不了什么

最后单独讲培训，因为这是老板最容易误解的一项，也是最容易把钱花在错误预期上的一项。

很多老板把培训当成 AI 转型本身，我给全公司做了 AI 培训，我们就算转型了。这是个危险的误解。培训有它确切的作用，也有它确切的边界，你得分清楚。

培训能解决的：认知拉齐

培训真正能解决的，是把全公司的 AI 认知和基础使用能力，拉到一条及格线以上。

我们服务客户时，内部用一套从 LV1 到 LV10 的能力分级来看员工。大部分企业的员工，都还卡在 LV1 到 LV3，也就是只会跟 AI 聊聊天、问问问题，把它当个更好用的搜索引擎。我做培训的核心目标，是把这批人提到 LV5 到 LV7，能把 AI 真正用进自己的日常工作里，用 workflow 的方式批量解决问题，而不只是手动一问一答。

AI 使用水平 10 级

Lv	称号	能力标准
0	旁观者	知道 AI 但从没用过 · 全球 80% 在这
1	尝鲜者	帮我总结、帮我写邮件，你不会追问，不会补充背景，一般只用豆包或 deepseek，把 AI 当高级搜索引擎
2	对话者	会追问 + 补背景，知道提示词，开始意识到怎么问非常重要，工作中两三个场景应用，如写周报
3	驯化师	主动立规矩 + 调 Prompt，开始给上下文、拆分任务，AI 产生结果从随机变成大部分可用 超 70%
4	越境者	用 AI 做专业外的事，做过海报、做数据分析，为 GPT、Claude 付费，发现能力边界大幅扩张 超 90%
5	织网者	AI 嵌入 workflow、有固定用法、开始学会搭建 workflow 和智能体、开始设计自己和 AI 之间的协作方式
6	召唤师	从 ChatBot 跨到 Agent，开始尝试 ClaudeCode、skills，用 AI 做出自己的产品 超 97%，体验差最震撼
7	铸造师	设计 Agent，自己造 Skill、AI 反馈循环、复利曲线启动
8	造物主	开始研究 AI coding、CC、Codex，搭建自己工作台 超 99%
9	觉醒者	AI 成为思维方式的一部分，遇到问题第一反应怎么和 AI 一起做好 0.01%
10	一人军团	产出能力一个人对标几十人团队 · 系统能力 >> 个人能力 顶级 0.001%

来源 · 数字生命卡兹克

按我的实际经验，一套培训做下来，大概一个月后，百分之六七十的员工能把 AI 用进日常工作提效；继续陪跑一段时间后，这批人的工作效率能提升到原来的两三倍。这是培训实打实能拿到的结果，它把人的“会不会用”这个底盘抬起来了。

这件事极其重要，因为后面所有的 AI 项目，落地都要靠人。员工连基本的 AI 认知和使用都没有，你做再好的智能体，他也用不起来、不愿用，培训解决的就是这个底盘问题。

培训解决不了的：场景、流程、系统

但培训的边界也非常清楚。它解决不了三件事：

第一，解决不了场景诊断。培训能让员工会用 AI，但回答不了：我公司哪个业务最该上 AI 这个问题，那是诊断要干的事。

第二，解决不了流程改造。你公司那些又长又乱的业务流程，比如一个客诉要穿过客服、工单、分析师好几道人墙，培训改变不了它。流程改造是一项需要专门去拆、去重构的工程，不是培训能顺带解决的。

第三，解决不了系统落地。真正把 AI 嵌进业务、做成一个能稳定跑的智能体，这是个交付工程，涉及知识治理、流程嵌入、准确性把控、运营维护，培训给不了你这些。

培训是入口，只是入口。它把人拉到及格线，让公司有了能承接 AI 的人。但从人会用了到业务真的被 AI 改造了，中间还隔着诊断、拆业务、流程改造、系统落地这一整段路。把培训当成转型的全部，就好比以为招齐了会开车的司机，公司的物流体系就自动建好了，司机是必要的，但远远不是全部。

我服务过的某四大行之一也做培训，它是体制内的，出发点更纯粹，就是提效，从聊天式 AI 用法过渡到执行式 AI 用法。对照某头部电商平台那种带着大厂危机感的培训，你能看出：不同企业做培训的出发点不一样，但培训的边界是一样的，它都只负责把人拉到及格线，负责不了后面的落地。老板对培训抱有正确的预期，这笔钱才花得值。

老板该学到什么

这一章如果你是老板，带走这几条判断：

第一，起步方式取决于你卡在哪，不存在统一的第一步。认知不齐就先培训，有意愿没靶子就先诊断，有明确痛点就直接试点。先搞清楚自己是哪种状态，再选路径。选错状态，后面全是无用功。

第二，别照搬大厂打法。大厂直接切复杂场景、搭体系，是因为它有数据、有团队、有系统。你一个中小企业，该走的是“现成工具、轻量落地、单点跑通”的路。你羡慕的大厂案例，是用大厂资源喂出来的，复制不了。

第三，无论如何别一上来做大系统。大而全的系统没法验证、越复杂越不稳、资源全押进看不见底的窟窿。永远先做能在数周内看到结果的小切口，跑通了再扩。

第四，培训是入口只是入口。它能把人拉到及格线，解决不了场景诊断、流程改造、系统落地。对培训抱正确预期，这笔钱才不白花。

一句话

这一章你记住一句话就够了：**AI 转型的第一步，要先看你公司现在卡在哪，再看 AI 能干什么，卡在认知就先培训，卡在没靶子就先诊断，有明确痛点就小切口试点；唯独不要一上来就做大系统。**

第 14 章 如何拆解业务

大家都在讨论答案，没人对齐问题

我和同行朋友都反复见过一个现象。

一个 AI 项目启动，老板召集大家开会。会上热闹：技术的人说我们用多 Agent 架构，数据的人说先把数据中台打通，业务的人说我要一个能自动回客户的智能体。每个人都在给方案、给答案，会开了三场，听上去推进得很快。

但真正动起手来，问题全暴露了。技术按系统架构去拆，业务按使用场景去拆，老板心里想的其实是按部门考核去拆，三个人对着同一个需求，脑子里装的是三张完全不同的图。做到一半发现对不上，推倒重来。

这就是 AI 项目最大的隐形杀手：所有人都在讨论答案，却没有人先对齐问题。

选场景、上 AI 之前，有一道绕不过去的前置功课，叫拆业务。不是拆技术架构，是把你这摊业务本身拆开、摊平、看清楚，哪个环节高频、哪个环节重复、哪个环节耗人、哪个环节卡在某个老员工的脑子里。拆清楚了，你才看得见 AI 该嵌在哪。拆不清楚，后面所有的动作都是在猜。

这一章我们主要讲四件事：为什么必须先拆业务再上 AI、业务要拆到什么颗粒度、拆完怎么看、谁来拆。这是承接能力里任务定义这一环最核心的动作，做不好，后面再强的模型也使不上劲。

为什么不能直接上 AI，要先拆业务

老板最常有的冲动是：问题我都知道，直接做就完了，拆什么拆，浪费时间。

我太理解这种心情了，但这恰恰是返工成本最高的起点。启动前花数周把问题对齐、把业务拆清楚，看起来是慢，实际上能省下后面十倍的返工。问题没对齐就动手，等于是在错误的地基上盖楼，盖得越快，后面拆得越惨。

不拆业务的代价：解决了一个不存在的问题

我前面提过那个做女装选品的客户。他张口就要：AI 自动选品系统，问题在他看来清清楚楚。但如果当时不去拆他的业务、直接照做，会发生什么？

我们会吭哧吭哧做一个选品模型出来，然后发现：选品这事压根不该 AI 来扛，它依赖买手的眼光和供应链关系，是这家公司最核心、最难替代的能力。真正天天耗人、把员工累到没空干

别的活的，是客服和上新文案。

也就是说，不拆业务，你很可能花大价钱、解决了一个根本不该解决的问题，而真正该解决的问题一直晾在那。老板以为的痛点和真正的痛点，经常不是一回事。拆业务的第一个价值，就是把“老板以为该做的”和“业务里真正该做的”分开。

不对齐问题的代价：做出来的东西没人认

还有一种更隐蔽的代价。

你请一个外部团队做 AI，如果一开始不把什么时点、交什么成果、用什么标准验收对齐清楚，几乎一定会互相失望：你觉得对方做得不够，对方觉得你要得太多。对齐问题，不只是对齐做什么，还要对齐做到什么程度算成、什么时候交什么。这一步省了，后面全是扯皮。

所以拆业务这道前置功课，本质上是在做两件事：一是看清真正该解决的问题在哪，二是让所有相关的人，老板、业务、技术、外部团队，对着同一张图说话。这两件事做扎实，后面才省心。

这事不只我一个人这么看，有个同样做企业 AI 落地的同行，完整复盘过他给一家检测机构梳理需求的全过程，几个判断和我完全一致，我借来帮你看得更清楚：

第一，**客户嘴上说的需求，和他真正需要的，中间隔着好几层**。对方一开口是：我要做个能自动生成检测报告、还能审核排版的系统，但每家的流程都不一样，不把需求一条条捋清楚，连交付标准都是糊的。

第二，**AI 能跑流程，但不能替你发明规则**。报告该怎么分类、价格怎么算、哪些算合格，这些业务规则，得客户团队自己先梳理清楚，AI 只能在你给定的规则上跑，不会凭空替你定标准。这一条尤其要记住：很多老板以为上了 AI 就不用自己理规则了，正好反了。

第三，**别为了上 AI 而硬造一个多余的环节**。他那个客户本来想做个拍照录入的 App，但检测数据本来就在设备里、导个表就行，真正的问题是打通不是录入。

这三条全都指向同一件事：真正的功夫在动手做之前，把业务拆开、把规则理清、把真问题找准。这一步省了，后面用再好的 AI，都是在错的地基上跑。

业务拆到什么颗粒度

知道要拆了，接下来的问题是：拆到多细？

拆得太粗，看不出业务流程，我们公司就是销售、生产、财务三块，这种拆法等于没拆。拆得太细，陷进流程图的汪洋里出不来，几个月画不完图。拆的颗粒度，有一个明确的落点：拆到你能看出，哪个环节高频、重复、可以被 AI 替代或辅助为止。到这个程度就够了，不用再细。

沿着“岗位—流程—环节”往下拆

我拆业务走的是从岗位到流程、再到环节这条线。

第一层，先看岗位。列出这个业务里所有的岗位，问一句：哪几个岗位最耗人、人数最多、最贵？人最多、最贵的岗位，往往就是 AI 价值最大的地方。比如某头部电商平台的上百人财务团队，光是岗位规模就告诉你，这里值得深挖。

第二层，拆这个岗位每天/每周在干的流程。把这个岗位一天、一周的工作摊开，看他到底在干什么。比如风险分析师，他的流程是：接客诉工单→登录后台拉数据→查命中的策略→判断是误拦还是真风险→把结论写回工单。一条流程清清楚楚摆出来。

第三层，在流程里标出每个环节的三个属性。这是最关键的一步。对流程里的每一个环节，问三个问题：

频率高不高？这个环节是每天发生几百次，还是一个月才一次？

重复度高不高？每次做的事是不是大同小异，还是每次都不一样？

依赖经验吗？这个环节是按固定规则就能干，还是要靠老员工的手感判断？

标完这三个属性，哪个环节该上 AI，基本就跳出来了。

一个拆法演示

我们之前做的支付客诉那个项目，拆出来是这样的：风险分析师的核心流程，是处理客服转上来的客诉工单。我把一段时间的工单拉出来分析，沿着上面三个属性一标，发现：分析师每天处理的问题里，有一大半是重复的。同一类拦截、同一种异常，换个用户又来一遍，换一天还是这些。这个环节，频率极高、重复度极高，而且其中大部分的判断逻辑是可以拆清楚、写下来的。

这就是一个标准的 AI 切入点：高频、重复、判断逻辑可显性化。剩下那些不重复、信息不足、需要深度研判的，频率低、依赖经验重，就不该硬塞给 AI，该留给人。

拆到这个颗粒度，哪里该上 AI、哪里不该上就不用拍脑袋了，它是从业务里自己长出来的。

一个反向的判断：有些环节天生不该拆给 AI

拆的过程中，你也会拆出一批“看着能上 AI、其实不该上”的环节。判断标准是反过来的那三条：

责任边界不清的不拆给 AI。比如要对外做风险承诺、给客户拍板赔不赔，这种出了事要有人担责的，不能让 AI 直接决策。

强人际博弈的不拆给 AI。比如谈判、压价、安抚一个情绪激动的大客户，这些靠的是人和人之间的博弈，AI 顶多打个下手。

高风险终局决策的不拆给 AI。比如前面说的，把一笔真风险当误拦放行了，损失是真金白银。这种最后一锤子的高风险判断，留给人。

拆业务不只是找 AI 该干的活，同样重要的是划出 AI 不该碰的活。两边都拆清楚，你的图才完

整。

拆完怎么看、谁来拆

业务拆完、属性标完，你手里就有了一张业务地图：哪些环节高频重复可替代，哪些环节责任重博弈强不能碰，一目了然。这张图怎么用、由谁来画，是这一节要讲的。

拆完怎么看：对照“哪些值得做”

拆完不是目的，拆完是为了给下一步“选场景”提供原料。

你拿着这张业务地图，下一章（哪些场景值得做）里会告诉你怎么给每个高频环节标个价、看它成不成熟、该不该先做。但前提是，你得先有这张地图。没拆业务就去选场景，等于没看菜单就点菜，点出来的全是凭印象。

这里先记一个衔接点：拆业务负责摊开看见，选场景负责挑出该做的。这一章只管把业务摊开、摊清楚，挑哪个是下一章的事。

谁来拆：业务方必须在场，光靠技术拆不出真问题

拆业务这件事，绝对不能只交给技术的人去拆。技术人员对业务理解非常有限，他能画出系统架构图，但他不知道风险分析师扫一眼工单凭什么就判断出是误拦，不知道财务做费用审核时心里那套哪些必须卡死、哪些可以放的隐形标准。这些藏在业务老手脑子里的判断，才是拆业务最值钱的部分，而它只能从业务人嘴里挖出来。

但也不能只靠业务方，业务方对自己的活太熟，熟到说不清，你问他怎么判断的，他张口就是“凭经验啊”。他知道怎么干，但不知道哪些是可以被 AI 学走、哪些是 AI 学不了的。

所以拆业务的正确配置是：一个既懂业务、又懂 AI 边界的人，带着业务老手，一笔一笔把经验抠出来。我自己拆经验的笨办法是：找团队里最资深的人，拿他处理过的案例，一笔一笔过，他每判断一笔，我就在旁边追问，

“你第一眼看的是哪个字段？”“看到这个值，你为什么立刻觉得是这种情况？”“什么情况下你拿不准、必须自己深挖、不敢直接下结论？”

一开始他答得很笼统，凭经验，但追问得足够细、足够多，那层凭经验就裂开了，底下真正的判断结构露出来。把这些答案攒起来、归类，一个原来只可意会的判断过程，慢慢就变成了可以写下来的逻辑。

这个追问、显性化的过程，既是在拆业务，也是在为后面 AI 落地准备最关键的原料。它需要一个人同时懂业务那头在说什么和 AI 那头能用什么，才接得上。纯技术的人不懂业务，接不住；纯业务的人不懂 AI 边界，提炼不出来。这是拆业务最难、也最值钱的地方。

我也见过反例，印证这个配置有多重要。同行做过一个销售智能体，管理层觉得做得特别高级，结果一线导购根本不用，因为拆业务的时候，没让一线导购真正参与进来，做出来的东西

没踩在导购最忙、最需要帮忙的那个点上。判断一个业务拆得对不对，有个朴素的标准：做出来的东西，能不能在做事的人最忙的时候，给他一句真正用得上的帮助。答案藏在一线那里，不在会议室里。

老板该学到什么

这一章如果你是老板，带走这几条判断：

第一，上 AI 之前必须先拆业务，这道功课不能省。启动前花数周对齐问题、拆清业务，能省下后面十倍的返工。不拆业务直接做，大概率是花大钱解决了一个不该解决的问题。

第二，拆到能看出哪个环节高频、重复、可替代为止。沿着岗位→流程→环节往下拆，给每个环节标三个属性：频率、重复度、是否依赖经验。拆到这个颗粒度就够，不用画完整张流程图。

第三，拆业务也要拆出 AI 不该碰的活。责任边界不清的、强人际博弈的、高风险终局决策的，这三类要明确划出去，留给人。

第四，业务方必须在场，而且要有个懂业务又懂 AI 边界的人来主导。真问题藏在业务老手脑子里，光靠技术拆不出来；但业务老手又说不清自己的判断，需要有人一笔一笔追问、把它显性化。这是拆业务最难、最值钱的一步。

一句话

这一章你记住一句话就够了：**AI 落地最常见的死法，是所有人忙着讨论答案，没人先对齐问题，而对齐问题的唯一办法，是先把业务一笔一笔拆开，看清楚哪个环节真正高频、重复、值得 AI 来接。**

第 15 章 哪些场景值得做，哪些不值得做

选错场景，比不做还糟

上一章把业务拆开了，手里有了一张业务地图。这一章要回答的，是老板们最关心、也最容易选错的一个问题：这么多环节，到底哪个值得上 AI，哪个不值得？

这个问题选错的代价，比你想的大。不做 AI，顶多是没有享受到红利；选错了场景做，是真金白银地把钱、把团队的耐心、把老板的信心，全砸进一个不该做的地方。第一个场景选错，往往不只是这一个项目失败，是把全公司对 AI 的信任一起赔进去。

所以选场景不能拍脑袋，得有方法。下面是一套系统看自己企业的诊断方法，三步：第一步找场景（哪里有机会），第二步看成熟度（现在能不能落地），第三步排优先级（应该先做哪个）。三步走完，哪个该做、哪个该等、哪个压根别碰，清清楚楚。

这套方法的核心，是一个视角的转换。



第一步：哪里有 AI 机会，核心是给问题标价

找 AI 机会，绝大多数人的姿势是错的。他们从 AI 能干什么出发，AI 能写文案，那我做个文案

工具；AI 能做客服，那我做个客服机器人。这叫工具思维，它的毛病是：你会找出一大堆能用 AI 的地方，但其中大部分是鸡毛蒜皮，做了也省不下几个钱。

正确的姿势是反过来，从哪个问题最值钱出发。不是问哪里能用 AI，是问哪个问题如果解决了，对我的业务结果影响最大。

给每个问题标个价

一个具体动作：拿出你的业务地图，给上面每一个问题、每一个环节，标一个价。标价不是估算 AI 项目要花多少钱，是估算这个问题如果被解决，一年能给你带来多少钱，省下来的、多赚到的、避免损失的，都算进去。

我举几个标价的例子：

- **会议纪要。**很多老板第一个想做的 AI 场景就是自动整理会议纪要。你给它标个价：就算公司每天开会产生纪要的时间全省了，一年能值多少钱？可能几万块。这是个典型的能用 AI、但太便宜的问题。
- **销售成交率。**假设你有一支销售团队，人均每月成交 200 单。如果 AI 辅助能把人均提到 250 单，这一项标出来的价，可能是几百万、上千万。同样是上 AI，这个问题的价，是会议纪要的几百倍。
- **发货计划优化。**一家做实体货的公司，如果 AI 能优化发货计划、减少压货，省下来的现金流可能是几百万。

再看一组对比，你立刻就懂标价为什么能救命。同样是花钱上 AI，两个老板：

老板 A，花三十万做了个 AI 系统，去替代一个年薪十万的岗位，你算算，光回本就得起三年，头一年还净亏二十万。

老板 B，花十万做了个 AI，把原来四个接单员优化到一个人顶上，一年省下三个人的成本，不到十个月就回本，三年净省二十多万。

同样是上 AI，一个赔钱赚吆喝，一个真金白银地赚。区别不在 AI 用得好不好，在选的问题对不对。一个温馨提示：AI 转型的起点，不是去找最显眼的问题，而是要去找回报最快的问题。最显眼的问题（比如老板天天看到、最闹心的那个），未必回本最快；回本最快的，往往是那些“人多、重复、一优化就立省”的不起眼环节。

你把价标完，会发现一件残酷的事：大家抢着做的（会议纪要、文案这类），往往是最便宜的问题；而最值钱的问题，经常没人看见。因为最值钱的问题通常藏在业务深处，销售怎么成交、库存怎么周转、风险怎么控住，它们不像“会议纪要”那样显眼，但它们才是真正能改变业务结果的地方。

标价是为了“别把好钢用在便宜的问题上”

标价这个动作，本质上是帮老板做一道排序：别把能改变业务结果的 AI，浪费在太便宜的问题上。

一个 AI 项目从立项到落地、再到持续运营，认真做下来，投入不会小。所以你标价的时候，心里要有一杆秤，这个问题值不值得为它专门折腾一个 AI 项目？太小的问题（一年省不下几万的），用现成工具顺手解决就行，别专门立项；真正值得立项做定制的，得是那种解决了能明显改变一条业务线结果的问题。

标价的时候，还有一个快速判断器，一句话就能用：这个问题解决了，好处是看账单看得见的，还是只能看感觉？

这两类问题，优先级完全不同。看账单的问题，把四个接单员优化成一个、把一份报告的制作成本从八十块降到八块，它的价值直接写在账上，不用开会讨论，数字摆在那。这类问题，优先做，因为它的回报清清楚楚、没人赖得掉。看感觉的问题，比如 AI 让我们的方案显得更专业，员工用 AI 心情更好了，它可能真有价值，但这价值得靠主观感受去判断，会波动、会扯皮，老板今天觉得值、明天又觉得不值。这类问题，往后放。

有个朴素的用法：凡是你给一个问题标完价，发现还得靠一堆主观判断才能说清它到底值不值，就把它从优先做的清单里划掉。优先做那些算完账、数字硬邦邦摆在那的。这一刀切下去，能帮你避开一大半，听起来很美、做完却算不清账，的项目。

还要说清楚：给问题标价，不只是第四、第五层（组织协同、业务模式）那种大场景才用，它是贯穿所有层的筛选原则。哪怕你只是想做个培训，也得问一句：这培训值不值这个钱、能换回什么。标价是个习惯，不是个动作。养成了，你看任何 AI 机会，第一反应都是它值多少钱，而不是它酷不酷。

第二步：现在能不能落地，看成熟度

标完价、找到了那几个最值钱的问题，先别急着上。找到机会，不等于现在就能做。一个问题值一千万，但如果落地的条件根本不具备，你硬上，就是把一千万的诱惑变成一千万的坑。

第二步是给每个高价值的场景做一次成熟度体检，看它现在到底能不能落地。五个维度挨个过：

第一，数据有没有。AI 要做事，得有料喂它。这个场景需要的数据，你公司有没有、全不全、能不能拿到？比如做销售辅助，得有历史成交数据、客户画像、产品资料。数据是散在各人电脑里的几个 Excel，还是有完整可用的沉淀？数据没有，这个场景就先放下。

第二，流程清不清楚。这个环节的工作流程，能不能说清楚？谁在什么情况下、按什么标准做什么判断？上一章拆业务，拆的就是这个。流程是一团乱麻、连人都说不清，AI 更没法嵌进去。

第三，结果能不能验证。AI 给出一个结果，你有没有办法判断它对不对？能验证的，你才敢用、才能持续优化；不能验证的，AI 答得再流利，你也不知道它在不在胡说，不敢真用。

第四，有没有人负责。这个 AI 做出来之后，谁来用、谁来“养”它（持续维护、补知识、改 badcase）？一个没人负责的 AI 项目，上线就是死亡的开始。

第五，系统能不能接入。AI 要嵌进业务，得能接到你现有的系统里，ERP、CRM、客服后台。接不进去，它就只是个孤立的玩具，很难产生价值。

这五条，是一个场景能不能落地的体检表。五条都过，这个场景成熟，可以做；缺一两条，先补齐再做；缺一大半，这个场景现在还不成熟，标价再高也得等。

AI 适合什么、不适合什么

AI 适合做的：规则相对清晰、信息量大、需要做总结/分类/提取/生成的活。比如把一堆非结构化的资料整理成结构化信息，比如按既定规则做初步判断，比如生成内容初稿，比如从海量数据里提取关键特征。这些是 AI 的主场。

我做的几个项目都落在这条线上。某大厂那个规则模板推荐的项目，要解决的是“分析师重复造轮子、知识散在代码文档邮件里查不到”，这是个知识密集、需要提取和推荐的活，AI 适合。做下来，配置模板的时间明显缩短，推荐覆盖率也达到较高水平。某精密制造企业那个 ISO 审核也是，本质是从一堆历史资料里提取信息、自动回填审核表，AI 适合，数周的活能大幅压缩。

AI 不适合做的：责任边界不清的、强人际博弈的、高风险终局决策的。这三类我在上一章拆业务时讲过，这里再强调一遍，因为它是成熟度判断里最该警惕的。需要对外做承诺、需要谈判博弈、需要拍最后一锤子高风险的板，这些别指望 AI，它干不了，硬让它干就是埋雷。

（同行也讲过：有三类问题不该让 AI 直接回答，需要业务判断或风险承诺的、来源和责任不清的、知识体系本身没理顺的。说的是一回事：成熟度不到、或者本就不该 AI 拍板的，别硬塞给它。）

第三步：应该先做哪个，排优先级



走完前两步，你手里会剩下一批值钱、又成熟的场景。但你不能说这些都能做 AI 然后一起上，资源不够，一起上等于哪个都做不深。第三步，是把这批场景排个序，定出先做哪个、后做哪个、哪些先别碰。

排序的核心原则就一句：价值最高、阻力最低的，先做。

一个简单的分层：

一期先做：轻、快、能出成果的。第一个项目，目标不是做最值钱的那个，是先跑通、先出一个能被看见的结果。所以一期要挑那种价值还不错、但落地阻力小、能在数周到一两个月内看到成效的场景。它的作用是立威，让老板和团队亲眼看到 AI 这东西真能在我们公司跑起来、真能省钱，建立信心和经验。某精密制造企业那个 ISO 审核就是个好的一期场景：问题具体、范围可控、成果看得见。

二期再做：更值钱、但更复杂的。有了一期的成功和经验，再去啃那些价值更高、但落地更复杂、要动的数据和流程更多的场景。这时候团队有信心了、有方法了、老板也愿意投了，啃硬骨头的条件才成熟。

暂时别碰：成熟度不够的。那些标价很高、但数据没有、流程没理清、结果没法验证的，放进“等”的池子里。不是不做，是等条件补齐了再做。现在硬碰，只会赔钱赔信心。

只适合培训、不适合定制的：单点提效类。还有一类，问题真实，但其实不需要专门做个 AI 系统，比如员工各自写写文案、查查资料、整理整理信息。这类用现成工具加培训就能解决，别浪费一个定制项目的资源在上面。

这套“一期/二期/暂缓/培训”的分层，就是从一堆“都能做”里，挑出现在最该做的那一个。一个企业的 AI 转型，几乎都是从一期那个小而准的项目，一点点滚出来的。没有谁是一上来就铺开做对的。

一次诊断，胜过十次拍脑袋

把这三步连起来看：找场景（给问题标价）→ 看成熟度（五维体检）→ 排优先级（价值高阻力低的先做）。这就是一套系统诊断自己企业 AI 机会的方法。

这套方法是想让你明白一件事：企业最该做 AI 的地方，往往不是你最先想到的那个。最先想到的，通常是最显眼、最被博主念叨的（会议纪要、写文案、做客服），而它们往往不是最值钱的。最值钱的机会，藏在业务深处，要系统地标价、体检、排序，才挖得出来。

凭感觉拍一个场景就上，和系统过一遍业务再决定，这两者的差距，就是做一个用不起来的项目和做一个真能改变业务结果的项目的差距。一次像样的诊断，省下的是你后面在错误场景上反复折腾、反复赔钱的成本。

我接触的客户里，凡是第一个项目做成的，几乎都是先把业务系统过了一遍、想清楚了才动手

的；凡是上来就凭老板一句话直接做的，翻车的概率高得多。

老板该学到什么

这一章如果你是老板，带走这几条判断：

第一，找 AI 机会，核心是给问题标价，不是找哪里能用 AI。从哪个问题最值钱出发，给每个环节标个价。你会发现大家抢着做的往往最便宜，最值钱的问题经常没被看见。别把好钢用在便宜的问题上。

第二，找到机会不等于能做，先做多维成熟度体检。数据有没有、流程清不清、结果能不能验证、有没有人负责、系统能不能接入。五条都过才做，缺一大半就等。

第三，排优先级，价值高、阻力低的先做。一期挑轻快能出成果的立威，二期啃更值钱的硬骨头，成熟度不够的暂缓，单点提效的交给培训。

第四，你最该做 AI 的地方，往往不是你最先想到的那个。系统诊断一遍，比凭感觉上项目稳得多。第一个场景选对，后面的转型才滚得起来。

一句话

这一章你记住一句话就够了：**选场景不是找哪里能用 AI，是先给每个问题标个价，挑出最值钱又能落地的那个先做，你企业最该做 AI 的地方，往往不是你第一个想到的那个。**

第 16 章 为什么企业应该先从人机协同开始

最容易高估的一件事：全自动

场景选对了，该动手做了。这时候老板心里几乎都有一个默认的目标：做全自动。既然上了 AI，那就让它从头到尾自己干，人就别管了，这才叫先进、才叫值。

我理解这种期待，但这恰恰是企业 AI 落地翻车率最高的一个点。全自动看着最美，在企业场景里风险最大。这一章要给一个反直觉但极其重要的判断：别一上来追全自动，先从人机协同开始，让 AI 当副驾驶，人坐在驾驶位上。跑稳了、有数据了，再一点点把方向盘交出去。

这一章讲三件事：全自动为什么容易翻车、人机协同具体怎么落、为什么要先降级再放权。

全自动为什么容易翻车

你想象一个全自动的流程，分五步走，每一步交给 AI。假设 AI 每一步的准确率都很高，有九成。听上去很可靠了对吧？但你把五步连起来算： 0.9 的五次方，等于 0.59 。五步走完，整条流程的可靠性只剩下不到六成。（这个计算是假设每个环节独立串联，业务里如果有人复核或交叉校验，风险会被部分抵消。）

这就是错误复合效应。单看每一步，AI 都挺准；但全自动意味着把这些步骤串成一条没有人看守的链，误差一步步累积，最后的结果烂得超出你的想象。因为中间没有人，甚至你连它在哪一步错了都不知道，等发现的时候，错误已经走到终点、造成实际损失了。

我们之前做风控时，判错一笔的代价是真金白银的损失，把一笔风险交易当成误拦放行了，损失可能是大钱。所以在做客诉智能体那个项目时，从一开始就没追求全自动、没追求让 AI 解决 100% 的问题。我们目标是解决大部分高频重复的问题，少部分真正复杂、信息不足、需要深度研判的，智能体不硬答，直接转人工，转回给风险分析师。

为什么不贪那最后的 20%？因为那部分复杂问题恰恰是最容易判错、风险最高的部分。让 AI 硬扛这部分，省不了多少事，却把最大的风险敞口暴露出来了。把这部分留给人，既守住了风险底线，又把从大量重复劳动里解放出来。

人机协同具体怎么落

讲了半天人机协同，它落到地上到底长什么样？很简单：**AI 做初稿、初筛、辅助分析**，人负责判断和决策。AI 把累活、重复活、信息整理的活干掉，把一个半成品递到人手上；人在这个半成品的基础上做最后的判断、拍板。

给你看几个具体的分工方式，一看就懂：

AI 做初稿，人来定稿。比如写内容、写报告、写方案，AI 先生成一个草案，人在草案上修改、把关、定稿。这比人从零开始写快得多，又守住了质量。

AI 做初筛，人来复核。比如审核、风控，AI 先把海量的东西筛一遍，把明显没问题的、明显有问题的分开，把拿不准的挑出来给人；人只需要复核 AI 拿不准的那部分。人的精力从此只花在刀刃上。

AI 做辅助分析，人来决策。比如风险研判，AI 把相关数据拉齐、把特征提取出来、把可能的结论和依据摆在分析师面前；分析师基于这些数据，快速做出最终判断。我做的那套风险分析的项目，核心就是这个，AI 把分析师原来要花大量时间手动排查的活做掉，分析师反应时间从原来的多天压到一天内，但最后那一锤子，始终是分析师敲的。

上面三种分工有一个共同点：人始终在最后那个判断/决策的位置上，AI 在前面把活铺好。这就是人机协同的骨架。它不追求把人赶走，它追求把人放到最该由人来干的那个环节上。

有研究把这套人机各守其位的分工，按人介入多少可以分成三档，再看每一档的实际提效：

升级模式：AI 自主处理八成以上，人只复核它拿不准的例外，中位提效 71%。

人在环中模式：AI 和人一起干，但每个输出都要人过一遍，中位提效约 22%。

审批模式：每一步都要人点头同意，中位提效 30% 左右。

你盯着这个结果看：给 AI 多一点自主权、人只守住例外的那种模式，效果反而最好。这恰恰印证了这个分工思路，把大部分高频任务放手给 AI、把复杂例外留给人，不是保守，这是被数据验证过的最优区间。

但必须记住：不是监督越少越好。像医疗、金融这种错一次就出大事的场景，人工逐个审核是硬要求，有相关法规要求，该让人把关就得把关。所以先人机协同是按你这个场景的容错空间来定的，容错高、错了能补救的（像内容、像大部分客服），可以让 AI 多担一点；容错低、错一次就是大损失的（像钱、像命），人就得守得更紧。

有同行把这件事讲得很好，他叫先把 AI 降级，很多项目失败不是做不出来，而是把负担从一个地方转移到了另一个地方。降级不是降目标，是降低 AI 的决策权，给它找对位置。能用流程先解决的就别让 AI 扛，能用数据承担的就别让模型猜。

先降级，再逐步放权

老板听到人机协同，把 AI 降级，可能会担心：这是不是太保守了？是不是不够创新？我花钱上 AI，结果还得人盯着，那 AI 的价值打折了吧？

但降级不是不创新，降级是让项目活下来、可控。一个全自动但不敢信、不敢用的 AI，价值是零；一个人机协同、稳定能用的 AI，价值是实打实的。零和实打实之间，选哪个不用我说。

而且，人机协同不是终点，是起点。正确的路径是：先降级，跑稳；再凭数据，逐步放权。

具体怎么放？当一个 AI 在人机协同的模式下跑了一段时间，你会积累两样东西：一是真实的运行数据，它在哪些问题上判断准、哪些上容易错，二是一套发现和修正错误的机制，也就是 badcase 库。有了这两样，你就能看清楚：在哪些已经被验证得足够准的子环节上，可以放手让 AI 自动跑；哪些环节它还不靠谱，继续留着人。

放权是按子环节、凭数据，一块一块放的，不是一刀切地从全人工跳到全自动。我们做策略部署那个项目就是这个路子，先把整个流程线上化、让人在系统里高效流转（这是降级，AI 还没怎么介入），跑顺了、数据攒够了，才识别出其中一部分邮件需求是可以自动部署的，把这部分放给自动化，但仍然保留人工的 double check。没有一上来就说全部自动部署。这部分是从数据里分析出来的，不是拍脑袋定的。

这个节奏一定要拿住：今天的大部分草案加人工复核，比一步到位的 100% 全自动，好交付得多、也稳得多。先求稳、再求自动化的扩大，这是企业 AI 唯一不翻车的放权方式。

老板该学到什么

这一章如果你是老板，带走这几条判断：

第一，别一上来追全自动，先从人机协同开始。全自动有错误复合效应，每步九成准，五步连乘只剩不到六成。最危险的复杂问题留给人，AI 干高频重复的部分。Klarna 先替代后回调，就是活教材。

第二，人机协同就是 AI 做初稿/初筛/辅助、人做判断/决策。人始终站在最后那个拍板的位置上，AI 在前面把活铺好。

第三，先降级、再凭数据逐步放权。降级不是降目标，是让项目可控。放权要按子环节、凭运行数据一块一块放，不是一刀切跳到全自动。

一句话

这一章你记住一句话就够了：**全自动不是起点，而是结果；企业 AI 最稳的打法，是先让 AI 做副驾驶，人守住判断和责任，等数据证明它可靠了，再逐步把方向盘交出去。**

第 17 章 ROI 怎么算，才不容易被画饼

老板心里那杆秤：能挣回来多少

前面几章讲了从哪起步、怎么拆业务、选哪个场景、怎么稳着落地。但老板心里始终悬着最后一个问题，这个问题不解决，前面讲得再好他也不敢动手：这件事花了钱，到底能挣回来多少？

这是这本书一开头我就说的、老板脑子里仅有的四个问题之一。这一章，就专门算这笔账。

我先把一个最常见的误区点破：很多人算 AI 的 ROI，只会算一件事，省了几个人。上了 AI，裁掉两个人，一年省下三十万工资，项目花了二十万，所以赚了十万，账算到这就完了。

这么算，要么把账算窄了（很多价值没算进去，你会觉得不划算、不敢投），要么被供应商画饼忽悠了（对方拍胸脯说能省一半人，你没法验证）。这一章是一套既不把账算窄、也不被人画饼的算法：全口径地看 AI 的价值，用真实数字算一笔自己信得过的账，对那些算不出来的价值也有正确的处理方式。

只算省人，是算窄了

AI 的价值，绝不止替代人力这一项。如果你只盯着省了几个人，你会漏掉它一大半的价值。我把 AI 能创造的价值掰开，至少有六个口径：

第一，省人力。这是最直观的，把重复劳动自动化，省下人工。但请注意，这往往是 AI 价值里最小、最不重要的一块。

第二，提效率。同样的人，做事变快了。我们做的支付客诉那个项目，一个客诉问题从提出到拿到结论，原来要走完客服、工单、分析师一整条人墙链路，动辄要几小时，现在压到分钟级。我们没裁任何人，但每个分析师单位时间能处理的事翻了好几倍。这部分效率提升，你不能只算省了几个人，要算同样的人产出翻了几倍。

第三，降错误率。AI 把错误率压下来，省下的是返工成本和事故损失。我在某大厂做内容安全，低质内容占比明显下降，审核时间也大幅缩短。错误率降下来这件事，省的不只是人工，是后面因为错误引发的一连串麻烦。

第四，提响应速度。在很多业务里，快本身就是钱。做风险分析的项目，对一个新出现的攻击，反应时间从原来的多天压到一天内，在风控这行，响应越快，风险资金漏损就越容易被控制。在账号风险体系也一样，拦截效果明显提升，帮助团队更早识别和处置高风险账号。这些都不是省人能概括的，而是直接堵住了往外漏的钱。

第五，沉淀知识。这一项最容易被忽略，但对很多企业最值钱。一套 AI workflow，本质上是把原来锁在老员工脑子里的经验，变成了公司能留住、能复用的资产。老员工离职，能力不再随之消失；新人上手，有一个随时可问的老师傅。这相当于给你的企业上了一份经验不流失的保险。这份保险值多少钱？想想一个核心老员工突然离职给你带来的断档，你就知道了。

第六，撑增长。AI 还能直接帮你多挣钱，不只是省钱。我们做的跨境电商内容工厂，把单条内容成本大幅压低，同时日产量百倍提升，带来持续的自然流量。这里既有降本，更关键的是产能的解放带来了增长，原来一天做不了几条、铺不开量，现在能大规模铺量，直接撑起了业务增长。

所以你算 AI 这笔账，该往这些个方向多想一层：除了它能帮我省多少，更值钱的问题是它能不能让我做成原来做不成的事。

你把这六个口径都摆出来，再回头看只算省人那种算法，就知道它有多窄了。AI 真正的价值大头，往往在效率、错误率、响应速度、知识沉淀和增长上，而不在省下的那几个工资上。一个有意思的旁证：德勤的调研发现，那些 ROI 领先的企业，最看重的是营收增长（占 49%），而不是降本，会算账的人，都把眼睛放在了增长上，不只盯着省钱。

怎么算一笔老板信的账

口径搞清楚了，具体怎么算成一笔账？一个简单但够用的框架：投入 vs 多维收益。

投入：算全，别只算开发费

投入这边，别只算供应商报的那个开发价。一个 AI 项目的投入，至少包括三块：

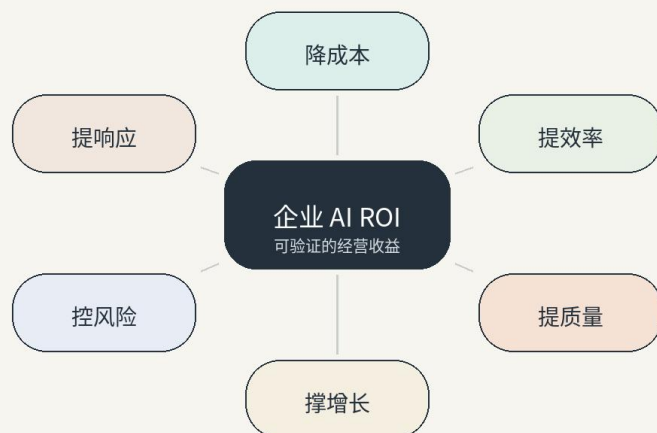
- **建设投入：**开发/采购的费用，这是显性的。
- **运营投入：**上线之后持续的维护、知识库更新、badcase 修正、模型/接口的调用费用。这部分经常被低估，但它是长期的。
- **人的投入：**业务方配合拆业务、参与测试、学习使用的时间成本。

把这三块加起来，才是真实的投入。只算第一块，你会低估成本，后面追加预算时措手不及。

收益：按六个口径，能量化的就量化

ROI 六类收益图

ROI 不能只算省人，要把经营价值算全



收益这边，拿前面那六个口径，挨个问：这个项目在这六项上，各能带来多少？能量化的，落到具体数字。我用支付客诉那个项目演示一下怎么落数字：

提效率： 处理时长从完整流程压缩到小时级，覆盖了大量高频客诉问题，准确率达到较高水平。

省人力 + 沉淀知识： 分析师从大量重复问题里被解放出来，他们最贵的工时，重新花在了研究新欺诈手法、对抗黑产这些真正只有他们能干的事上。同时，他们脑子里的判断经验被沉淀成了公司资产。

光这一个项目，在提效率、省人力、沉淀知识三个口径上都有产出。如果只算省了几个人，你根本看不到它真正的分量。

数字要能验证，不能听供应商一张嘴

这是不被画饼的核心，任何收益数字，都要问一句：这个数，谁来验证、怎么验证？

供应商跟你说上了我这套能提效 50%，你要问：提效 50% 怎么定义？用什么基线对比？上线后用什么数据来证明？如果对方答不上来，或者说上了就知道了，这就是画饼。

真实的项目，收益数字一定是能对照基线、能被数据验证的。我前面给的那些数字，完整流程压缩到小时级、拦截效果明显提升、内容质量显著改善，每一个都有上线前的基线、有上线后的实测，经得起对照。一个不敢给基线、不敢谈验证方式的收益承诺，默认它是画的饼。

行业回报水位

老板可能会问：别人做 AI，回报到底是多少？给你几个可以参考的外部数据，帮你心里有个谱，也帮你识别哪些承诺明显离谱。

平均回报：微软联合 IDC 2024 年的一项调研测算，企业每投入 1 美元在生成式 AI 上，平均回报约 3.7 美元，头部企业能到 10.3 美元。

真实的骨感面：麦肯锡 2025 年的报告同时指出，只有约 39% 的企业报告了 AI 对利润（EBIT）的实际影响，而且多数影响还不到 5%。另一份 2025 年企业生成式 AI 研究提出，约 95% 的试点没有产生可衡量的损益回报。这个口径更适合作为风险提醒，别把它读成所有项目都会失败。

把这两组数据放在一起看，可以看出：AI 确实能带来可观回报，但很多企业还没拿到。那些没拿到回报的项目，问题往往不在模型，在前面几章讲的那些：场景选错、流程没拆、贪大求全、上线没人养。

这对老板算 ROI 的启示是：别用头部 10 倍那种数字来给自己画饼，也别把“95%”当成命运。选对场景、稳着落地，才能拿到健康回报；拍脑袋上、贪大求全，项目很容易卡在试点。回报高低，取决于你的落地能力，不取决于 AI 本身。

难量化的价值，别硬算 ROI

有些东西，你明知道它有价值，但实在算不出一个精确的 ROI 数字。比如：

决策提前：AI 让你比对手早一天看到风险、早一周看到趋势，这个早值多少钱？

运营透明：原来一团黑箱的业务流程，上了 AI 之后全程可追溯、可管理，这份看得见值多少钱？

经验不流失：前面说的那份老员工离职能力不断档的保险，值多少钱？

这些价值是真实的，但你硬要算出一个精确的 ROI，会算得很勉强、说服力还差。对这类难量化的价值，别硬算，换一种方式让老板自己掂量，用风险铺垫和损失回想。

所以对决策提前、运营透明、经验沉淀这类价值，正确的算法是问老板：你过去因为“看不见”“反应慢”“人一走经验就没了”，吃过哪些亏？把这些吃过的亏摆出来，老板自己心里就有数了，这套系统不是为了多赚一笔算得清的钱，是为了堵住那些算不清、但确实有漏的窟窿。这种风险铺垫的算法，比硬凑一个 ROI 数字，对老板更有说服力，也更诚实。

算账时最该警惕的两个陷阱

提醒老板算 ROI 时最容易掉进去的两个陷阱：

陷阱一：用建设型成果冒充经营型成果。供应商交给你的，常常是建设型成果，我搭了什么系统、上了什么能力。但你真正要的是经营型成果，它给我的经营结果带来了什么。这两者中间有巨大的鸿沟。一个项目能力都搭好了，不等于经营结果改善了。算 ROI，一定要算到经营结果那一层，别被我们能力都建好了糊弄过去。

陷阱二：不对齐时点就谈回报。这是我们和老板都要负责的，AI 项目的回报是分阶段释放的，一期可能只是把基础搭好、看到小成果，真正的大回报在二期、三期。如果一开始不把什么时点、看什么回报对齐清楚，老板用终局的预期去看一期的成果，一定会失望、一定会觉得被画饼了。算 ROI 要连着时间轴一起算，这个阶段该看到什么、那个阶段该看到什么，事先说清楚。这样既不会互相失望，也能及时发现项目是不是真的在往回报走。

行业里有个对照数据：BCG 2025 年的调研发现，78% 的高管期望 18 个月内见到回报，但只有 23% 的中层觉得能交付。这个巨大的预期差，就是时点没对齐的后果。

老板该学到什么

这一章如果你是老板，带走这几条判断：

第一，只算省人是算窄了。AI 的价值有六个口径：省人力、提效率、降错误率、提响应速度、沉淀知识、撑增长。大头往往在后五项，不在省下的那几个工资上。

第二，算账要投入算全、收益能验证。投入要算建设+运营+人的投入三块；收益每个数字都要问：谁验证、怎么验证，不敢给基线、不敢谈验证方式的承诺，默认是画饼。

第三，难量化的价值别硬算 ROI，用损失回想。决策提前、运营透明、经验不流失这类，问自己：过去因为没有它吃过哪些亏，比硬凑数字更有说服力。

第四，警惕两个陷阱：别拿建设成果冒充经营成果，别不对齐时点就谈回报。ROI 要算到经营结果那一层，要连着时间轴一起算。

第五，回报高低取决于你的落地能力，不取决于 AI。大量试点没拿到回报，问题往往都在落地，选错场景、流程没拆、贪大求全。方法对了，健康回报完全拿得到。

一句话

这一章你记住一句话就够了：**AI 的 ROI 不能只算省了几个人，也不能只停留在供应商嘴里提效 50% 那种虚数，它是一笔投入算全、收益能验证、连着时间轴算的真账；算不清的价值，就去回想你过去因为没有它吃过的亏。**

第 18 章 智能体到底是什么

一个被你反复听到、却没人给你说清楚的话

这两年，你大概在无数个场合听过智能体这个词。供应商来谈合作，开口就是我们帮您搭一套智能体；行业论坛上，人人都在讲智能体元年、智能体重构一切；你手下的人也来跟你说，老板，我们也得上智能体了。

你点点头，但心里其实没底：这东西到底是什么？和我之前用的那个能聊天的 AI，有什么区别？

这一章你记住这一句就够了：智能体本身是个壳。它能不能真做事，不取决于这个壳，取决于壳背后那一大堆你看不见的东西搭得好不好。

很多老板买了智能体用不起来，不是智能体这个概念有问题，是壳背后那一大半的活没人干。这一章就是带你看清楚：壳是什么、壳背后是什么、为什么这事跟你这个拍板的人有关系。

智能体不是聊天机器人

你打开手机，问 AI 帮我写段祝酒词，这道菜怎么做，它答你一段话，这是聊天机器人（ChatBot）。你问一句，它答一句；你不问，它不动；它答完就完了，不会自己去查东西、不会自己分几步干、也不记得你上礼拜问过什么。

它就像一个口才很好的顾问，坐在那儿等你提问。你问什么，他凭脑子里现成的知识答什么。但他不会起身去帮你跑腿、不会去翻你公司的档案柜、更不会自己判断这事得分三步办。

智能体的定义很简单，你不用记术语，记这四件事它能干：

第一，它能听懂一个任务，不只是听懂一句话。你跟聊天 AI 说帮我分析这笔交易，它给你一段分析文字。你跟一个智能体说：帮我处理这个客诉，它知道这背后是一连串的话：先去查这笔交易为什么被拦、再判断是误拦还是真风险、再决定该不该放行、最后给出一个能用的结论。它理解的是任务，不是一句话。

第二，它有记忆。它记得这个任务做到哪一步了、前面查到了什么、你之前交代过什么规矩。不是每次都从零开始。

第三，它会自己调工具。这是智能体和聊天机器人最大的区别。聊天机器人只会跟你说话，智能体会动手，它能去查数据库、能调一个查征信的接口、能打开你公司的某个系统拉一份数据。它不只是说，它能做。

第四，它会自己规划步骤。一个复杂任务，它知道先干啥、后干啥、这一步的结果决定下一步

怎么走。不是你手把手喂一步它走一步，是你给它一个目标，它自己拆成几步去完成。

把这四件事合起来，就是我说的智能体：一个能理解任务、有记忆、会调工具、会自己规划步骤的工作单元。你可以把它想成一个新来的员工，不是一个只会回答问题的实习生，是一个你交代一件事、他能自己分几步把这件事办完、办的过程中会去查资料、会用各种系统、办不了还知道来找你的员工。

这个区别，对你这个老板太重要了。因为：聊天机器人，提升的是个人效率，你手下某个人写东西快了点。智能体改变的是流程，一整条原来要好几个人接力干的活，可能被它接管掉一大段。前者是给个人配把好用的工具，后者是动你的业务流程。这俩的分量，差着量级。

ChatBot 相当于大脑，智能思考，智能体相当于在大脑的基础上加了手脚，既可以思考，也可以帮你干活。

Chatbot 和 Agent 的区别

Chatbot 负责回答；Agent 负责完成



一个真做出来的智能体长什么样

光讲定义太空，给你看一个我真做过的东西，你就明白一个能做事的智能体是什么样子。

我在某互联网大厂做支付风控时，搭过一套东西，内部叫“风险 Co-Pilot”。它要解决的问题是这样的：这家公司的支付业务，每天有大量用户因为各种原因触发风控，交易被拦、账户异常、提现失败。其中相当一部分会变成客诉，涌到风险分析师那里。

我们当时拉过数据：分析师每天要处理大量这样的客诉，相当一部分工作时间都耗在排查这些 case 上。

排查一个 **case** 有多麻烦？分析师得先看这笔交易触发了哪条风控、再去关联各种数据（用户画像、设备、行为、历史关联等）、再判断这到底是误拦还是真有风险、判断完还得想该怎么处理、要不要出一条新策略去堵这个漏洞。一套流程下来，链条很长，而且高度依赖这个分析师的个人经验，同一个 **case**，老分析师和新分析师，可能得出不一样的结论。

我们做的，不是搭一个能聊天的 AI 让分析师去问。我们搭的是一套能自己把这套排查流程跑下来的智能体系统。它干的活是：接到一个风险线索 → 自己去关联各方数据、给这笔交易打上标签、生成一份分析报告 → 辅助分析师做判断 → 甚至能推荐一条应对策略、自动把策略生成出来 → 推到旁路引擎上去验证 → 再生成一份结果报告。

这就是前面说的那四件事全用上了：它理解的是排查这个风险这个任务（不是答一句话）、它记得这个 **case** 查到哪了（记忆）、它自己去调各种数据接口（工具）、它自己按关联→分析→推荐→验证这个顺序往下走（规划）。

这套东西做出来之后的效果：

客诉量：明显下降。很多原来要分析师全手动做的事，系统自己处理掉，分析师只需要做确认动作。

反应速度：明显变快。原来一个新出现的攻击，分析师要花不少时间研判；现在能更快给出处理结论。

风险资金漏损明显降低。这些数字背后，是分析师被从大量重复排查里解放出来了，他们终于能把时间花在真正难的、需要人脑的攻防对抗上。

这个例子想说明：一个真能做事的智能体，是嵌在你业务流程里、能自己跑完一段流程、能拿出可量化结果的东西。它不能只停在聊天窗口。这个判断，是你后面验货所有智能体项目的基准。

智能体只是个壳，干活靠这些依赖

为什么我说智能体是个壳？

因为智能体这层会理解任务、会规划、会调工具的能力，现在已经不稀缺了。任何一家像样的供应商，都能给你搭出这个壳。真正决定这个智能体能不能在你公司做事、干得好不好的，是壳背后这几样东西。

第一样：Prompt（指令）

就是你怎么跟这个智能体交代活。但企业里的 Prompt，不是你随口跟 AI 聊天那种。它是一套写得死死的规矩，告诉智能体：遇到 A 类情况怎么办、遇到 B 类怎么办、什么时候你必须停下来转人工、什么时候你绝对不能自己瞎猜。

Prompt 是把你公司的业务规矩翻译给 AI 听的那道工序，这道活的质量，直接决定智能体答得对不对。同一个模型，规则写得清楚和写得含糊，出来的结果天差地别。这不是技术活，这是业务活，得有懂你业务的人，把那些藏在老员工脑子里的判断标准，一条条写清楚。第 7 章我

讲的把老师傅的手感翻译成能写下来的逻辑，落到智能体上，就是这个 Prompt。（或者通过工作流来解决）

第二样：知识库

智能体光有通用知识不够，它得懂你公司的事，你的产品、你的政策、你的历史案例、你这行的 know-how。这些东西，装进一个它能随时查的库，就是知识库。

这是企业智能体项目里最容易翻车、也最值钱的一块，我专门用第 20 章讲它。这里你先记一句：文档丢进去不等于有了知识库。以前文档是给人看的，得经过一道清洗、整理、划边界的脏活，AI 才用得上。我做过的项目里，这道脏活的工作量，经常是整个项目的大头。

第三样：工具

就是智能体能动手调用的那些东西，查数据的接口、你公司的某个系统、一个能算账的程序。智能体之所以能做而不只是说，靠的就是这些工具。

你公司的系统接口干不干净、数据通不通，直接决定智能体能不能动手。很多企业卡在这，智能体这个壳搭得挺漂亮，但你公司的系统老旧、接口对不上、数据散在十几个地方各说各话，智能体想动手发现没处下手。这就是为什么我一直说，AI 进不进得来，大半取决于你 AI 之外的家底。

第四样：MCP

这个词你听供应商提过。我用大白话给你讲：MCP 就是一个统一的插座标准。

以前智能体想连你的某个系统，得专门定制一根线、接一个口，连十个系统就得接十种不同的口，又贵又乱。MCP 出来之后，相当于行业约定了一种统一的插座规格，你的系统按这个规格留个口，任何智能体来了都能直接插上用。

MCP 对你的意义，是降低了把 AI 接进你现有系统的成本和门槛。这是个好东西，它让智能体接入你公司各种系统这件事，从每个都得定制变成标准化对接。但你别被供应商拿这个词唬住，MCP 解决的是怎么接的问题，解决不了你的系统本身乱不乱、数据本身脏不脏的问题。插座标准统一了，但你家电路是乱的，照样用不起来。

第五样：Skill（技能）

你可以把 Skill 理解成：把一套做某件事的标准打法，打包成智能体随时能调用的一个招式。

比如怎么处理一笔退款客诉，这套流程是固定的：先核身份、再查订单、再判断符不符合退款条件、再走流程。把这套打法打包成一个 Skill，以后智能体遇到退款客诉，直接调这个招式，按这套标准打法走，不用每次现想。

Skill 这个东西，本质上是在帮你把公司里那些标准打法沉淀下来、变成可复用的资产。这跟我这本书反复讲的主线是一回事，AI 转型最值钱的事，是把锁在人脑子里的经验，变成公司能留住、能复用的东西。Skill 就是这件事在技术层面的一个落点。一个老板真正该上心的是：公司里那些值钱的标准打法，有没有被一条条沉淀下来，而不是有没有用上 Skill 这个名词。

这五样过完，你回头看那句结论就懂了：智能体这个壳，只是把这五样东西组织起来、让它们

协同做事的一个外框。壳人人会搭，但 Prompt 写得准不准或 workflow 设计合不合理、知识库治没治理、工具接没接通、标准打法沉没沉淀下来，这一大半的活，才是决定成败的地方，也才是真正费功夫、费钱、费人的地方。

AI 越智能、越自主，它对外围那套工程基础的要求反而越高。这听着反直觉，AI 不是越强越省事吗？恰恰相反。一个只会聊天的 AI，你给它一句话、它回一句，对你的系统、数据、流程没什么要求；但一个要自己调工具、自己跑多步流程的智能体，它每动一步，都要求你的数据是通的、接口是稳的、规矩是清的，任何一处不到位，它越自主，翻车翻得越彻底。

有一份研究表明，真正用了智能体 ic（自主智能体）的只占两成，但这两成的中位提效高达 71%，远超高自动化的 40%。为什么这么高的提效，却只有两成企业用上？就是因为壳好搭，壳背后那套工程基础，大多数企业还没搭起来。报告里提到：Agentic 不只是换了个新的聊天界面，它是重新定义了人和机器在工作流里的分工。这正是我反复强调的，别盯着那个壳，盯着壳背后你这家公司的家底。

不是一个大智能体，是一群各干各的小智能体

很多老板以为，智能体就是做一个聪明的大 AI，什么都让它干。一上来就想：我搞一个无所不能的智能大脑，把公司的事全交给它。

这个想法，恰恰是企业智能体项目最常死的地方。

我前面讲的那套风险 Co-Pilot，不是一个大智能体，是 7 个各司其职的小智能体在协作。怎么分工的？

有一个主驱动的智能体，它的活就是分类调度，线索来了，它先判断这是哪一类风险，然后派给对应的下游去处理，就像一个调度员。下面挂着好几个专门的分析智能体，每个只专精一类风险的深度分析。还有一个专门负责从海量日志里提取关键特征的智能体，它的活就这一件，但干得专、干得稳。

为什么要这么拆？因为让一个智能体啥都干，它就啥都干不精、还不稳。一个负责调度、几个负责专项分析、一个负责取数据，各管一段，每段都简单、可控、好检验。组合起来，才是一套既能干复杂活、又稳得住的系统。

我在大厂做过多个智能体，有专门答客服问题的、有替代人值夜班做安全运营的、有做智能告警的、有做监控预警的。做多了我越来越确信一件事：好的智能体系统，几乎都是一堆简单的小智能体各守一摊、拼起来的，不是一个号称无所不能的大家伙。

为什么这事跟老板有关系

讲了这么多，你可能会想：这些不都是技术细节吗？我一个老板，搞清楚智能体是个壳、背后有五样东西，有什么用？

用处很大，我给你三条，都是你拍板时能直接用上的判断。

第一，你不会再被我们做了个智能体这句话糊弄。

以后供应商跟你说我们帮您搭好了智能体，你心里有数了：壳人人会搭，得问的是壳背后那一大半，你给它配的规矩（Prompt）清不清楚？知识库治理了没有？接我系统的工具通不通？这些活做了多少？问得出这几个问题，你就从听不懂、只能信的位置，挪到了能验货的位置。

第二，你会明白买了智能体用不起来，问题大概率出在哪。

我见过不少企业，智能体买回来用不起来，老板第一反应是这 AI 不行、这供应商是骗子。其实十有八九，问题常常出在壳背后那一大半没人干，规矩没写清楚、知识库是一堆没整理的烂文档、要接的系统接口根本不通。智能体这个壳是供应商交付的，但壳背后那一大半，很多得靠你公司内部配合着一起干。你不知道这一层，就会一直怪错地方、也补不到正点上。

第三，你会对花多少钱有个正确预期。

如果你以为智能体就是个壳、搭起来就能用，你会觉得这东西不该花多少钱、做起来不该花多少时间。然后真做起来，发现钱和时间都远超预期，你就开始怀疑是不是被坑了。真相是：壳很便宜，壳背后那一大半才是大头。知道这一点，你报预算、定周期的时候，心里才有底，也才不会在该投入的地方抠门、最后把项目做夹生。

老板该学到什么

这一章，如果你是拍板的人，我希望你带走三条判断。

第一，智能体是个壳，价值在壳背后。

会理解任务、会规划、会调工具，这个壳现在人人会搭，不稀缺。决定它能不能在你公司做事的，是 Prompt、知识库、工具、标准打法这些壳背后的东西。你掏的钱，大头该花在这一大半上；你验货，也该盯着这一大半。

第二，真能做事的智能体，是嵌在流程里、能跑完一段活、能拿出数字的，不能只停在聊天窗口。

判断一个智能体真不真，就看它能不能接管你一段业务流程、能不能给出可量化的结果。能聊天但接不进流程的，是玩具，不是生产力。

第三，好系统是一群简单小智能体拼出来的，不是一个无所不能的大脑。

谁上来就给你画一个什么都能干的智能大脑，你要多一个心眼。好的智能体系统，几乎都是各司其职的简单部件拼出来的，这背后的道理，下一章细讲。

一句话

这章收束成一句话：**智能体是个壳，壳人人会搭；能不能在你公司真做事，看的是壳背后那一**

大半，规矩写没写清、知识理没理顺、系统接没接通。老板不用懂技术，但要知道：钱和功夫，大头都在你看不见的那一半。

第 19 章 如何判断一个智能体项目是真落地，还是 Demo

一场让你点头的演示，和你掏完钱之后的沉默

先看一个很常见的场景。

供应商来给你演示一个智能体。大屏幕上，他敲进去一句话，智能体哗哗哗自己分了好几步，调了数据、出了图表、给了结论，最后还自动生成了一份报告。整个过程行云流水，看着智能。会议室里所有人都在点头，你也觉得这东西真神，就要它了。

然后你掏了钱，三个月后，这套东西在你公司里基本没人用。要么是答得不对、没人敢信；要么是接不进你的系统、跑不起来；要么是上线没多久就出了个错，把人吓回去了，从此束之高阁。

演示时那么惊艳，落地后这么沉默，这中间隔着的，就是 Demo 和真落地的鸿沟。这道鸿沟，深得超出绝大多数老板的想象。

这一章，就是给你一双能看穿这道鸿沟的眼睛。让你在演示现场，就能大致判断出：这家给我演的，是一个真能在我公司做事的东西，还是一个只为了让我点头掏钱的漂亮 Demo。

Gartner 在 2025 年有个判断：到 2027 年底，超过四成的 Agentic AI（也就是智能体类）项目会被取消。他们还点破了一个现象，叫智能体 washing，数千家号称自己在做智能体的厂商里，真正名副其实的，大约只有一百三十家。（来源：Gartner 2025）

翻译成成人话：市面上十家说自己做智能体的，可能九家是在贴标签、做 Demo。你不会判断，大概率会踩坑。

还有个来自做企业 AI 一线的观察：从一个能演示的 demo，到一个真能在生产环境稳定跑的产品，中间差的不是 20% 的工作量，是 200%，也就是说，你看到那个惊艳的 demo，顶多算走完了三分之一的路，后面还有两倍于此的活没干。AI 让做出一个能演的 demo 变得很容易，但做成一个真能用的产品，门槛高很多。这就是为什么那么多项目演示时人人叫好、落地时全员沉默。

Demo 和真落地，到底差在哪

很多老板分不清这俩，是因为它们在演示现场看起来一模一样，都是输入一句话，智能体自己跑一通，出个漂亮结果。

但它们考验的东西，根本不在一个层面。**Demo** 考验的是能力，真落地考验的是能不能在你这儿活下来。

一场 **Demo**，只需要证明一件事：这个智能体能做这件事。它在一个干净、理想、供应商精心准备好的环境里，跑通一遍给你看。数据是挑好的、问题是设计过的、流程是顺的。它只要跑通这一遍，演示就成功了。

真落地要过的关，多得多。我列给你看：

第一关：数据。**Demo** 用的是供应商准备好的漂亮数据。你公司的数据呢？散在十几个系统里、格式各不相同、还有一堆历史脏数据。智能体进了你这儿，面对的是这摊烂账，不是演示用的样板间。

第二关：权限。智能体要做事，得能访问你公司的系统、调你的数据。这就牵扯到一连串问题：它有没有权限？哪些数据能给它看、哪些不能？这些在 **Demo** 里全不存在，**Demo** 是在供应商自己的地盘上跑的。

第三关：流程。**Demo** 是智能体自己跑一段。真落地，这个智能体得嵌进你现有的业务流程里，它的上游是谁、下游是谁、它的结果交给谁用、谁来接它的活。它不是孤零零跑一遍，它得是你流程链条上的一环。

第四关：准确率。**Demo** 给你看的那一遍，是对的。但真上线，它一天要处理成百上千个 case，这里面它能对多少？错的那些怎么办？一个偶尔错一次的智能体，在 **Demo** 里看不出来，在生产里是要命的。

第五关：责任。它万一答错了、判错了、把一笔该拦的放过去了，谁负责？这个问题 **Demo** 永远不会涉及，但它是你长期用下去绕不开的坎。（这一关太重要，我专门留了第 22 章讲。）

第六关：运营。上线只是入口。知识会过期、业务会变、会冒出新的错例。谁来持续地喂它、修它、养它？没人养的智能体，用不了多久就废了。

Demo 只要过能力这一关，真落地要过这六关。一个项目敢不敢、能不能带你过这六关，就是它真不真的分水岭。一上来只给你秀能力、绝口不提后面这六关的，基本可以判定：它停在 **Demo**。

判断真落地的几个标准

从 Demo 到生产的六道关

跨不过这六关，就还不能叫“真落地”



Demo 证明“能跑”；生产证明“能长期稳定地跑”。

六关讲完，你可能觉得太多、记不住。可以压缩成四个判断题。下次看智能体项目，无论是供应商的方案还是你内部团队的汇报，拿这四问去套。

第一问：它的结果，能不能被检验？

这是最重要的一条，放在第一位。

你想一件事：为什么写代码、做翻译这类 AI 用得最顺？因为它的对错能被立刻验证，代码跑一下就知道对不对。而有些活天生难，比如让 AI 写一段战略分析，好不好没有客观标准，全凭人感觉。

一个场景能不能让智能体真落地，很大程度上取决于：它干出来的结果，有没有办法被检验对错。能检验，你就能知道它什么时候在犯错、就能持续纠正它、就敢逐步放权。不能检验，你永远不知道它是真懂还是在一本正经地胡说，你也就永远不敢真信它。

我们做风险 Co-Pilot 时，为什么敢上？因为它的每一个判断，都能拿的风险结果去验证，它说这笔是误拦，后续有没有真出事，是查得到的。它说要上一条策略，推到旁路引擎上小流量一跑，有效没效，数据说话。整套系统是“感知→分析→推荐→验证→上线”一个能验证的闭环，不是它出个结论就完事。

所以你在现场就该问一句：这个智能体给的结果，我们拿什么去判断它对不对？答得上来、答得具体的，是真在做事；答得含糊、说它很智能您放心的，要警惕。

第二问：流程的主权，在不在你手里？

真落地的智能体，是嵌在你流程里的一环，而且关键的那几步，主动权得攥在你手里。

什么意思？意思是它不会从头到尾全自动跑完、让人在旁边干看着；它会在关键节点，尤其是高风险、要担责的那几步，它停下来，把判断交给给人，人点了头它再往下走。

我们做的那套系统，设计上就是这样：智能体把分析、把策略都准备好，但最后要不要上这条策略，是人确认、一键配置、再小流量上线。关键的方向盘，始终在手里。这不是不信任

AI，这是企业里能长期用下去的姿势，出了事有人能拦得住、能兜得住。

你该问的是：这套流程里，哪几步是智能体自己定的、哪几步是人能拦下来的？一个号称全自动、全程不用人管的智能体，听着很先进，但在企业场景里，往往是最危险的（为什么，我下一段专门讲）。

第三问：出错了，能不能兜得住？

它一定会出错，没有不出错的 AI。所以关键不是它会不会错，是它错了之后，这个系统接不接得住。

接得住，意味着：它拿不准的时候，知道停下来转人工，而不是硬着头皮瞎答；它答错了，有机制能发现、能纠正、能避免下次再错；一个错误，不会顺着流程滚成一场灾难。

我们做这些项目，有一个原则：知识库没覆盖到的、它拿不准的，直接转人工，这是设计。我宁可让它在复杂难题上老老实实说：我不会、转人工，也绝不让它在这些地方逞能瞎答。第 7 章我讲过这个分工道理，放到智能体上是一样的：一个能稳定处理大部分高频任务、并且老实把复杂问题交还给人的智能体，远比一个号称什么都能干、但你不信它的智能体有价值。

你该问的是：它遇到搞不定的情况，会怎么办？是转人工，还是硬答？答会转人工、有兜底的，是真懂企业；答它什么都能处理的，是没在生产环境里挨过打。

第四问：这套东西，能不能复用、能不能留下来？

一个真落地的项目，做完会沉淀下东西，一套整理好的知识库、一套写清楚的规矩、一套能复用的标准打法。下次做类似的事，这些能接着用，而不是每次都从零开始。

一个 Demo 性质的项目，做完什么都不剩。演示完、款一付，那套东西换个场景就用不了了，因为它本来就是为了演示那一遍临时搭的。

你该问的是：这个项目做完，我公司能留下什么可以反复用的资产？这一问，把一锤子买卖和能长出能力的项目区分开了。

越简单，往往越靠谱

面对智能体项目，越是给你画得复杂、画得宏大的，你越要警惕；越简单、越说得清的，往往越靠谱。

这话听着别扭，直觉上，不是越复杂越厉害、越值钱吗？恰恰相反。我用做项目的体会告诉你为什么。

第一，复杂本身就是成本，而且是你后面要一直还的成本。

复杂度不是能力的证明，是成本。一套架构复杂的多智能体系统，不代表做它的团队本事大，只代表你以后要为它的难维护、难排查一直付账。它做的时候贵、调的时候难、出了问题排查起来更要命，因为环节太多，你都不知道是哪一环出的错。而且它越复杂，越不稳，一个环节抖一下，整条链就可能崩。复杂度不是免费的，你今天图它看着厉害，明天就要为它的不稳定

和难维护一直买单。

我前面讲的风险 Co-Pilot，虽然有 7 个智能体，但每一个都简单、专一、只干一件事，一个调度、几个专项分析、一个取数据。它的复杂，是很多个简单部件拼起来，不是一个绕来绕去谁也看不懂的大家伙。这两种复杂，天壤之别。

第二，真正难的判断力，是忍住不把系统搞复杂。

做技术的人，有个天然的冲动：把系统越做越花、越做越炫。但企业里真正值钱的，是反过来的能力，能把一件事用最简单的办法稳稳办成，忍住不去堆那些用不上的复杂功能。

我见过把复杂的多智能体架构，拆掉、换成简单的几步串行，结果准确率反而更高、更稳、更好维护的情况。减法比加法难，也比加法值钱。一个供应商，如果上来就给你堆一堆花活、画一个庞大复杂的架构图，大概率他想秀的是技术、是要卖个高价，而不是真为你的稳定落地着想。

第三，简单的东西，老板才看得懂、才管得住。

一套你完全看不懂的复杂系统，出了问题你只能听供应商解释，完全被牵着走。一套简单、说得清的系统，你大致知道它在干什么、哪一步在干啥，你才有判断、有主动权。看不懂的复杂，是把主动权交了出去；看得懂的简单，是把缰绳攥在自己手里。

所以，下次看方案，如果一个供应商三句话给你讲清楚了他要做什么、怎么验证、人在哪几步介入，另一个供应商画了一张你看半天看不懂的宏大架构图、满嘴新词，别被后者唬住。在企业落地这件事上，前者赢面大得多。

一个反过来的样子：真落地长什么样

讲了这么多警惕什么，我们正过来看一眼，一个真能落地的智能体项目，在你眼前会是什么样子。

它会从一个小而具体的场景切入，不贪大。比如就解决客服处理某一类支付客诉，这一件事，而不是上来要重构整个客户服务体系。

它会坦诚地跟你谈那六关，会主动问你数据在哪、乱不乱，会跟你确认要接哪些系统、权限怎么给，会告诉你它打算处理多少、剩下多少转人工，会跟你商量出错了怎么兜底、上线后谁来运营。一个肯跟你谈这些麻烦事的供应商，比一个只给你秀漂亮 Demo 的，靠谱得多。

它会给你一个能验证的结果和可量化的目标，目标要写成这一类客诉，处理时间从 X 降到 Y、覆盖率达到 Z、准确率不低于多少，达不到不算交付，不能只写会变得很智能。

它做完会留下能复用的东西，一套知识库、一套规矩、一套打法，你下次能接着用。

你拿这个样子，去比对眼前的项目，八九不离十就能判断它真不真。它越接近这个样子，越靠谱；它越像那个输入一句话、哗哗跑一通、智能的演示，越要小心。

老板该学到什么

这一章，如果你是拍板的人，我希望你带走三条判断。

第一，演示惊艳不代表能落地，中间隔着六关。

Demo 只要过能力一关，真落地要过数据、权限、流程、准确率、责任、运营六关。看演示的时候提醒自己：它给我秀的是能力，但能不能在我这儿活下来，得看那六关。一上来只秀能力、不谈这六关的，基本是 Demo。

第二，四个问题当场就能验货，结果能不能检验、流程主权在不在你手、出错能不能兜、做完能不能复用。

记不住六关没关系，记住这四问。下次无论看供应商方案还是听内部汇报，挨个套一遍。答得具体的是真做事，答得含糊、只说很智能您放心的，要警惕。

第三，越复杂越要警惕，越简单越靠谱。

复杂度本身就是成本，是你后面要一直还的账。真正的本事是用最简单的办法把事办成，而不是堆一个谁也看不懂的庞大架构。谁上来给你画宏大复杂的架构图，你多留个心眼；谁三句话讲清楚要干啥、怎么验证、人在哪介入，这才是冲着落地来的。

一句话

这章收束成一句话：**Demo 比的是能不能演成功，真落地比的是能不能在你公司活下来；验货就盯四件事，结果能不能检验、方向盘在不在你手、出错能不能兜底、做完能不能留下能复用的东西。记住，越简单的，往往越靠谱。**

第 20 章 知识库 / RAG 项目怎么判断靠不靠谱

一个老板最容易上当的项目类型

在所有 AI 项目里，知识库是最容易让老板上当的一种。

为什么？因为它听起来太简单、太美好了。供应商跟你说：把您公司的文档、制度、手册全传进来，以后员工有问题直接问 AI，它就能从这堆资料里给您找出答案。

一听太合理了，我公司这么多年攒的文档、规章、案例，全在那儿躺着没人看，要是 AI 能把它们用起来、随问随答，那不就是个永不下班、什么都懂的老员工吗？于是你拍板：做！把文档一传，等着用。

然后，十有八九会失望，员工问它，它答得驴唇不对马嘴；问 A 它答 B；有时候还一本正经地编一个根本不存在的政策。用了几次，大家就不信它了，这个知识库慢慢就没人碰了。

这一章主要讲知识库这东西为什么这么容易翻车、一个靠谱的知识库项目长什么样、你作为老板该问哪几个问题，才能在掏钱之前就大致判断出它行不行。

一个核心判断：知识库的难点，从来不是 AI 不够聪明，是你那堆文档根本不是给 AI 准备的。这件事的本质，是个管理问题、是个脏活累活的问题，不是技术问题。

为什么文档传进去，AI 还是答不对

我们服务过一家制造业企业，是个大集团下面的一个质量体系部门。他们有个头疼的活：应付客户的 ISO 审核。大客户隔段时间就要来审一次，要求他们填一大堆审核表，证明自己的质量管理、流程规范都达标。这事原来全靠人工填，一次要耗两三周，上下游十几个部门参与。

他们初期想得很好：我们公司这么多年的质量文档、流程记录、历史审核资料堆成山，能不能搭个知识库，让 AI 自动把这些审核表填了？

我们最后是用一套多智能体系统帮他们做自动填表的，但这个项目里最难的、占了最多功夫的，都花在了清洗数据。

他们那堆文档有扫描进电脑的图片、有格式五花八门的表格、有十几年前的 PDF、有谁随手画的流程图、还有大量信息根本没写在文档里、只存在某个老员工的脑子里。这哪是什么知识库，这是一个堆了十几年、谁也没整理过的杂物间。

你把这么一个杂物间原封不动丢给 AI，让它从里面找答案，它当然找不对。它看到的是一堆它

根本读不懂、对不上、还互相矛盾的东西。这就好比让你一个新员工，从一个乱糟糟、没编号、文件还缺页的档案室里，准确无误地翻出某份资料，即使他在聪明也没办法在这屋子根本没法找。

我们真正花大力气把这个杂物间，一点一点清洗、整理、归类成一个 AI 能看懂、能用的知识库：把图片里的信息提出来、把乱七八糟的表格统一格式、把过期的和有效的分开、把藏在人脑子里的补上、把互相矛盾的捋顺。这道脏活累活干完之后，后面搭 RAG、写智能体、让它自动填表，反倒都不算难了。

这就是知识库的难，九成难在把你那堆“给人看的、乱的”文档，洗成“给 AI 用的、整齐的”知识。

海外有 AI 落地调研也有类似结论，他们做知识库项目，做过的最重要的一件事，就是花大量时间去打磨怎么把知识切分、组织好；模型本身，反而不是关键。在他们研究的项目里，42% 的情况下，用哪个模型其实完全可以互换，换个模型，结果都差不多，因为决定成败的是模型周围那套：知识怎么洗、怎么切、怎么组织的功夫。所以下次有人跟你说：我们用了最先进的大模型，知识库肯定准，你心里要清楚：模型不是重点，你那堆知识洗没洗、理没理，才是。

为什么会这样？因为你公司现在所有的文档，都是给人看的，不是给 AI 准备的。人看文档，能连蒙带猜、能凭经验补全、能去问同事；AI 不会，它只认整齐、清楚、有边界的东西。你这十几年的文档积累，本来就没按将来要给 AI 用的标准来攒。所以几乎每个企业上知识库，第一步撞上的都是这堵墙，文档一大堆，但没一份是 AI 能直接用的。

RAG 为什么会翻车

知道了难点在哪，我们再来看几个知识库翻车后能看到的现象，它们一出现，就说明知识库大概率是有问题的。

现象一：答非所问。你问它：我们公司差旅报销标准是多少，它给你扯了一段关于差旅审批流程的话。问的是 A，答的是 B。这通常说明：知识没整理好，AI 没能从那堆乱资料里找到真正对应的那一条。

现象二：答案过期了。它信誓旦旦告诉你一个政策，结果那是三年前的老规定，早改了。这说明：知识库里新旧版本混在一起，没人把过期的清出去，AI 分不清哪个是现行的。

现象三：它答了它本不该答的。你问一个它知识库根本就没覆盖的问题，它编了一个听起来煞有介事的答案给你。这是最危险的现象，它在一本正经地胡说，而你还信了。说明这个项目没有设计不知道就说不知道的机制，放任 AI 瞎编。

现象四：它答得模棱两可、不敢给准话。问它一个该有明确答案的问题，它给你一堆可能、也许、建议您咨询相关部门。说明知识本身就是乱的、矛盾的，AI 也捋不清，只好和稀泥。

这四个现象，你在试用任何一个知识库时都该盯着看。只要反复出现其中任何一个，这个项目的底子，也就是知识治理那一层，大概率没做好。这时候该追问的是你们的知识到底清洗、整理了没有，而不是先怪 AI 笨。

验货：你该问供应商的几个问题

现在到了最实用的部分，你不懂技术，但你要在掏钱之前，有能力把一个知识库项目问出真假。给你一串问题，你照着问，问到对方答不上来或者支支吾吾的，就要当心。

第一问：知识有没有经过清洗和整理？还是直接把文档丢进去？

如果对方跟你说很简单，您把文档发我们，传进去就能用，警惕。这话等于说他们要把你那个杂物间原封不动塞给 AI。真正靠谱的供应商，一定会跟你谈知识治理这件事，会告诉你你的文档需要清洗、需要整理、这部分是工作量的大头。一个肯跟你谈脏活的供应商，比一个把事情说得轻轻松松的，靠谱得多。

第二问：谁来维护它？知识过期了、变了，谁负责更新？

知识库不是搭完就一劳永逸的，政策会变、产品会更新、会冒出新情况。这些变化谁来同步进知识库？这事要是没人管，知识库用不了多久就开始答过期答案，慢慢就废了。你要问清楚：维护这件事，是供应商持续负责，还是交给你公司哪个人、哪个岗位？有没有说定？没有明确“谁来养”的知识库项目，基本可以判死缓，上线那天就是它开始烂掉的那天。

第五问：怎么知道它答得对不对？有没有一套题去考它？

一个靠谱的项目，会有一套考卷，拿一批有标准答案的问题，定期去考这个知识库，看它答对多少，这叫评测。没有评测的知识库，你根本不知道它到底准不准，全凭用的人感觉好像还行或者好像不太行。你要问：你们怎么衡量这个知识库答得准不准？有没有一套题、多久考一次？答得上来的，是真把准确性当回事；答不上来的，是做完就不管了。

把这五问连起来，清洗了没、有没有边界、不知道会不会瞎编、谁来维护、怎么验证准确性，你就有了一把验货的尺子。

一个靠谱知识库的样子

反过来，一个靠谱的知识库项目长什么样。你拿它去对照，心里就有数了。

我们做的那套营销风控的规则模板推荐系统，可以当个参照。那个场景的痛点是：分析师配置策略时，总在重复造轮子，因为有用的模板和知识，散落在代码、文档、邮件里，没人能一下子查到。

我们怎么把它做成一个靠谱的知识库的？

它的知识是被结构化整理过的。我们没有把代码和文档原封不动丢进去，而是把每一个模板的关键信息，它叫什么、有哪些参数、依赖什么、输入输出是什么，一条条提取出来、整理清楚。而且每一条信息，都绑定了它的出处。这样 AI 给你一个答案时，能告诉你这答案是从哪儿来的，你能核对。

它有明确的边界和兜底。你问它，它先理解你到底想要什么、缺哪些关键信息，然后才推荐。如果你给的信息不够、它判断不了，它不会硬编一个，而是明确告诉你还缺这几样信息，让你补上。这就是不知道就说不知道的设计。

它的好坏是被量化考核的。我们给它定了硬指标：推荐的前三个里命中正确答案的比例要达到多少、生成的配置一次通过的比例要达到多少、出结果要多快。有了这些数字，它准不准，不靠感觉，靠考卷。

这套东西做出来，原来需要小时级的配置工作，压到了分钟级；同类模板的重复新建，每月少了三成。

一个靠谱的知识库，就是这个样子，知识被整理过、有出处、有边界、不瞎编、被量化考核。你拿这五条去比对眼前的项目，差得越远，越要小心。

老板该学到什么

这一章，如果你是拍板的人，我希望你带走三条判断。

第一，知识库的难，九成在脏活，不在 AI。

文档传进去 AI 就能用，这是最坑老板的一句话。真相是，你那堆给人看的、乱的文档，要花大力气清洗整理成 AI 能用的知识，这道脏活才是工作量的大头。谁把这事说得轻松，谁不靠谱。知识库本质是知识治理，是管理问题，不是技术问题。

第二，认得四个翻车现象，答非所问、答案过期、瞎编乱答、模棱两可。

这四个现象但凡反复出现，就说明知识治理这一层没做好。这时候别怪 AI 笨，该追的是知识到底洗没洗、理没理。

第三，五个问题验货，清洗了没、有没有边界、不知道会不会瞎编、谁来维护、怎么验证准确性。

掏钱之前，拿这五问去考供应商。其中最关键的两条：它不知道的时候必须会说我不知道、转人工，绝不能瞎编；以及，必须说清楚谁来养它。这两条答不利索的，直接出局。

一句话

这章收束成一句话：**知识库答不对，十有八九是你那堆给人看的文档没被洗成给 AI 用的知识；验货就问五件事，洗没洗、有没有边界、不知道会不会瞎编、谁来养、怎么考。一个会说我不知道的知识库，比一个什么都敢答的，可信一百倍。**

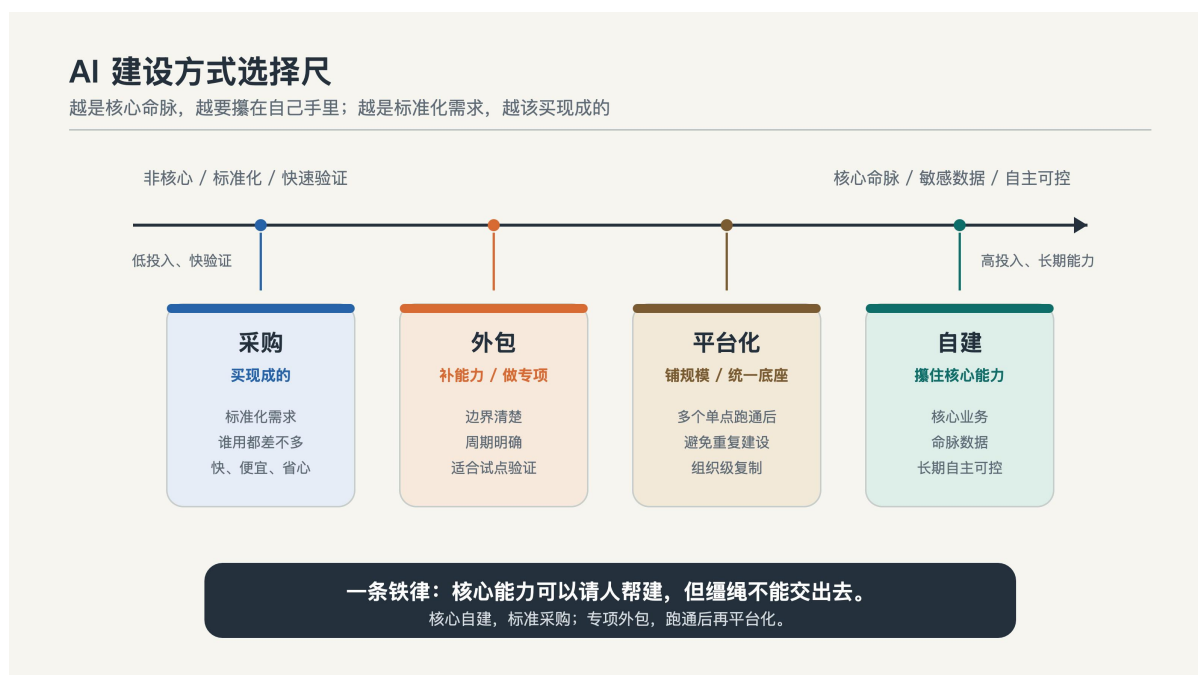
第 21 章 自建、采购、外包、平台化怎么选

同一句“我要做 AI”，四条完全不同的路

老板们想清楚了要做 AI、也选好了场景，那下一个问题就是：这东西，到底是我自己组队来做，还是买现成的，还是包给外面的公司，还是上一套平台让大家都能用？

这四种路子分别是自建、采购、外包、平台化。每一种，投入不一样、风险不一样、你能攥在手里的主动权也不一样。选错了，轻则多花冤枉钱，重则把自己的命脉交到别人手里。

这四种没有哪种是绝对好的，只有“适不适合你这个事”。一个基础判断：越是你的核心命脉、越要长期攥在自己手里的东西，越往自建那头走；越是边边角角、标准化、谁做都一样的东西，越往采购那头走。中间地带，用外包补人手、用平台铺规模。这条线怎么划，决定你这笔钱花得值不值、稳不稳。



四条路，分别适合什么

下面来展开说说四条路分别是什么，该什么时候用，不需要记名称，记每一种背后的适用条件就行。

第一种：采购，买现成的

就像你买办公软件一样，市面上有成熟的 AI 产品，你买个账号、开通服务，拿来就用。

什么时候用它：你要解决的，是一个标准化的、大家都差不多的需求。比如让员工用上一个通用的 AI 助手提高写东西的效率、比如用一个现成的会议记录工具自动出纪要。这类需求，你公司和别家公司没什么本质区别，市面上早有人做好了，你没必要自己造轮子。

采购的好处：快、便宜、省心，今天买明天用。**坏处：**它是给所有人做的，不会专门为你公司的特殊情况量身改；你公司那些独特的、藏在业务深处的需求，标准产品多半满足不了。

我的判断：凡是谁用都一样的需求，优先采购，别自己折腾。把自己的人力和预算，省下来留给那些真正只有你才需要、市面上没有的事。

第二种：自建，自己组队做

你拉起自己的团队（或者带着外部伙伴），从你公司的业务出发，定制一套属于你自己的东西。

什么时候用它：这件事是你的核心业务、涉及你的命脉数据、或者你必须长期自主可控。这种事，标准产品满足不了，外包又不放心，只能自己掌握。

我服务的那家区域特色连锁餐饮，做的就是自建。他们要做的，买个现成的点餐工具解决不了这件事。他们真正要做的是以智能体为枢纽，把菜品、供应链、会员、门店经营，甚至独有的健康餐饮养生知识整个打通重构。这是他们整个企业数字化的命根子，数据是他们的、流程是他们的、那套养生知识库更是他们独一份的护城河。这种东西，你能买吗？买不到。你敢完全包给外人吗？不敢。只能自己主导着建，一层一层长出来。

自建的好处：贴合你的业务、数据和能力攥在自己手里、能持续长成你的组织能力。**坏处：**慢、贵、对你自己的组织能力要求高，得有懂行的人主导。

我的判断：核心业务、命脉数据、要长期自主可控的，往自建走。这笔投入贵，但你买的是主动权。

第三种：外包，包给外面的公司做

你把一个明确的活，交给专业的外部团队，他们做完交付给你。

什么时候用它：你有一个边界清楚、周期明确的活，自己又没那么多人手或专业能力，适合找外部团队短期把它干出来。比如你想先做一个试点项目验证一下、比如某个专项你内部没人会做。

外包的好处：借到专业的人、不用自己长期养团队、能快速把一个项目推出来。**坏处：**做完之后，这套东西到底能不能在你手里持续用、持续改，是个大问号。外部团队交付完就走了，后续要改、要维护、出了问题要排查，你能不能接得住？这是外包最容易踩的坑。

我的判断：外包适合“边界清楚、周期明确”的专项，尤其是早期试点。但一定要约定清楚：做完之后的文档、知识留不留得下来，你能不能维护。否则就是一锤子买卖。

第四种：平台化，铺一套底座，让全公司都能用

它搭的是一套统一底座，让公司各个部门都能在上面长出自己的 AI 应用，而不是单个工具。

什么时候用它：你已经跑通了好几个单点、想把 AI 能力铺到整个组织、不想让每个部门各搞各的、各自为战。这是组织级的事，通常是企业 AI 走到比较深的阶段才该考虑。

平台化的好处：统一管理、统一标准、避免重复建设、能规模化复制。**坏处：**投入大、周期长，而且一上来就上平台，是企业最常犯的贪大求全的错。平台是用来铺规模的，前提是你已经有规模可铺，你得先跑通了几个场景，知道大家都在用什么、需要什么，再搭平台去承接。没跑通就先搭平台，十有八九搭出一个没人用的空壳。

我的判断：平台化是水到渠成的事，不是一步到位的事。先用前三种把单点跑通、跑出价值，再考虑用平台把它铺开。顺序反了，大概率是个昂贵的摆设。

怎么根据自己的情况选

四条路讲完，你可能还是会问：那我这个具体的事，到底走哪条？可以用三把尺子，挨个量一遍，答案基本就出来了。

第一把尺子：这是不是你的核心业务？

你先问自己：这件事，是我赖以生存、别人抢不走的命脉，还是一个谁都有、不出彩也死不了的辅助环节？

是核心命脉，往自建走，攥在自己手里。是辅助环节、谁用都一样，能采购就采购，别浪费力气。

举个例子：那家餐饮企业，把健康餐饮养生推荐+全链路经营数据打通当成核心命脉，所以自建；但他们公司日常写写通知、整理整理材料这类活，完全可以买个通用 AI 助手解决，犯不着自己造。同一家公司，不同的事，走不同的路，别一刀切。

第二把尺子：你要不要、能不能自主可控？

第二问：这件事，你需不需要完全攥在自己手里？数据敏不敏感？将来要不要频繁地改、自己说了算？

需要自主可控、数据敏感、要频繁迭代的，往自建那头走，哪怕慢点贵点。你能容忍交给外人、数据也不那么敏感的，可以采购或外包。

这里有个提醒，跟数据安全相关（下一章会专门展开）：有些数据，涉及你公司的商业机密、客户隐私、核心算法，这种东西能不能往外送、能不能放在别人的服务器上，你心里要有杆秤。这类强敏感、强核心的，往往逼着你只能走自建、还得是私有部署。

第三把尺子：你的预算和周期，经得起哪种？

第三问，现实一点：你这次准备投多少钱、能等多久？

钱少、要快、想先验证一下值不值得，别上来就自建、更别上来就平台化，先采购个现成的、或者外包一个小试点跑跑看。等跑出价值、看清方向了，再决定要不要加码自建。

钱够、愿意为长期能力投入、这事又确实是核心，那就别图便宜，踏踏实实自建。在核心的事上图便宜，是最贵的。

这三把尺子量下来，你会发现一个规律：多数企业的正确打法，不是非此即彼，是组合着来，核心的事自建、标准的事采购、专项的事外包、跑通了再平台化。一家公司里，这四种路同时存在，是正常的、健康的。谁要是跟你说：您整个公司就该全用我这一种，这话本身就值得怀疑。

把这个组合着来再往深推一层，你别把做 AI 想成一个整体的、要一次性决定自建还是采购的大块头。其实一个 AI 应用，是分好多层的，底下是算力、模型，中间是数据、知识库、工具接口，上面是贴着你业务的应用逻辑。这每一层，你都可以单独决定是“租别人的、配现成的、还是自己造”。比如：最底层的模型，你大可不必自建（自己训模型是巨头的事），租用、调用现成的就好；但最上面那层贴着你核心业务的逻辑、和你那套独有的知识库，就该牢牢自建、攥在自己手里。很多企业的毛病，是把这八九层混成了一个“是或不是”的决定，要么幻想全自建，要么干脆全外包。真正会算账的老板，是一层一层看：哪层是通用的就租、就买，哪层是你的命脉就自己攥住。把这笔账拆到每一层去算，你既不会在通用的地方浪费钱造轮子，也不会命脉的地方把家底交出去。

别把核心能力，完全交到别人手里

讲完怎么选，我要专门拎出一条最容易被忽视、但代价最大的提醒。

在 AI 这种变化飞快的领域，千万别把你的核心能力，整个打包外包出去，自己一点不掌握。

为什么？我给你两个理由。

第一，你会失去主动权。

如果你公司最核心的那套 AI 系统，从架构到数据到怎么运转，全在一家外部公司手里，你自己人一窍不通，那么这家公司就拿捏住了你。他要涨价、他要拖延、他哪天不想干了、甚至他被别人收购了，你都只能干瞪眼。你以为你买的是服务，其实你把脖子伸过去了。尤其当这套系统是你的核心命脉时，这种依赖是致命的。

第二，AI 变得太快，你必须自己跟得上。

AI 这个领域，几个月就一个样。模型在变、能力在变、做法在变。如果你的核心 AI 能力完全靠外人，那么这个领域每往前走一步，你都得等外人来告诉你、帮你跟上。你对 AI 的理解和判断力，会一直停留在听别人说的水平，永远长不出自己的肌肉。而在这个时代，一个企业自己有没有 AI 的判断力，差距会被越拉越大。

核心的事也可以找外部帮忙。关键在于：主导权和核心资产是不是还在你手里。

还是那家餐饮企业，他们做最核心的自建，也带着外部的专业团队一起，但主导的是他们自己，数据是他们的、那套独有的知识库是他们的、整个架构是围着他们的业务长出来的。外部团队是来帮他们做事、补能力的，不是来当这套系统的主人的。这就是健康的姿势：借外部的

力，但缰绳在自己手里。

这种借外部的力、但主导权在自己的模式，其实正在成为企业 AI 落地的一条主流答案。我观察到，连 OpenAI、Anthropic 这些最顶尖的 AI 公司，给大企业做落地时，采用的方式通常更重：派工程师驻到客户现场，而不是简单地卖个产品或包个项目，和客户一起，把混乱的流程一点点拆开、再用 AI 接起来，业内把这种角色叫 FDE（前线部署工程师）。它更像四条路之外的一种组合形态：共建 + 驻场 + 持续陪着迭代，区别于传统采购和传统外包。

为什么这种模式正在兴起？因为 AI 落地最难的那部分，把你那套独有的业务、数据、规矩接进 AI，买个标准品解决不了，甩给外人做完也很难扎根，它必须有懂 AI 的人，蹲在你的业务现场，和你一起长出来。AI 时代的核心项目，更适合“我们一起建，但建的过程我全程在场、最后东西归我”。这恰恰是借力但缰绳在自己手里最落地的一种形态。

老板该学到什么

这一章，如果你是拍板的人，我希望你带走三条判断。

第一，先分清核心和非核心，再决定走哪条路。

核心命脉、敏感数据、要长期自主可控的，往自建走，贵也值，买的是主动权。辅助的、标准化的、谁用都一样的，能采购就采购，别浪费力气造轮子。一家公司里四种路同时存在，是正常的，别一刀切。

第二，平台化是水到渠成，不是一步到位。

平台是用来铺规模的，前提是你已经有规模可铺。先用采购、外包、自建把单点跑通、跑出价值，再用平台铺开。没跑通就先搭平台，大概率是个没人用的昂贵摆设。

第三，核心能力可以借力，但缰绳得攥在自己手里。

别把命脉系统整个外包出去、自己一窍不通，那是把脖子伸给别人。在 AI 这个快变的领域，你必须自己长出判断力。请外部团队帮你建可以，但主导权、核心数据、关键资产，必须在你这儿。借力可以，交命不行。

一句话

这章收束成一句话：先问这事是不是你的命脉，是命脉就自建、攥在自己手里；是边角料就采购、别造轮子；专项缺人手用外包、但要约好东西留得下；跑通了再平台化、别一步到位。一条铁律：核心能力可以请人帮建，但缰绳不能交到别人手里。

第 22 章 安全、准确性和责任边界怎么把关

前面几章，讲的都是怎么把 AI 做出来、怎么验货、怎么选路子。但你心里最深的那个顾虑，我们还没解。

那个顾虑就两个字：出事。

AI 把一笔不该放的钱放过去了，谁负责？它给客户答了一个错误的承诺，这事算谁的？它把公司的机密数据，顺着哪个口子漏出去了，怎么办？它自动跑的一个流程，中间错了一步，等发现的时候已经酿成大祸，这个锅谁背？

只要这两字没想清楚，你就永远只敢让 AI 试试看，不敢真把它放进核心业务里长期用。而企业 AI 的价值，恰恰要在长期用力才能兑现。

所以这一章，是你从试试看迈进放心用的最后一道门槛。核心三件事：数据安全怎么守、准确性怎么保、出了事谁负责。讲完，你心里那块石头能落地。

这张核心记住一句话：AI 的安全和准确，从来不是模型自带的，是你设计出来的、管出来的。你要是指望买一个又安全又准确的 AI 回来就万事大吉，那一定会出事。真正靠谱的做法，是在它周围，搭一圈人来把关、用机制来兜底。

一、数据安全：哪些能往外送，哪些打死不能

数据安全，是三件事里最不能含糊的，因为它一旦出事，往往是不可逆的、是要命的。

第一刀：先把你的数据分个级。

不是所有数据都一样金贵。你得先在心里把数据切成几档：

最敏感的： 客户隐私（身份证、手机号、交易记录）、公司核心机密（配方、核心算法、未公开的财务）、涉及合规红线的（法律明确管着的那些）。

比较敏感的： 内部经营数据、还没公开的业务信息。

不太敏感的： 公开的产品资料、对外的宣传材料、行业通识。

分完级，规矩就清楚了：越敏感的数据，越要往自己手里收；越不敏感的，越可以放开用外部服务。

第二刀：决定哪些能上云、哪些必须留在本地。

你用一个外部的 AI 服务（比如某个云上的大模型），本质上是把数据发到别人的服务器上去处理。这对不敏感的数据没问题，方便又便宜。但对最敏感的那一档，客户隐私、核心机密，我的判断很明确：这类数据，要么别往外送，要么就走私有部署，放在你自己能完全控制的环境里跑。

私有部署，简单说就是把 AI 装在你自己的机房或者你专属的服务器上，数据从头到尾不出你的门。它贵、它麻烦，但对那些漏了就完蛋的数据，这个钱不能省。第 21 章我讲过，有些强敏感、强核心的事，逼着你只能走自建、还得私有部署，根子就在这儿，是数据安全的底线逼出来的。

第三刀：就算在内部用，也得管住谁能看什么。

不是说数据留在公司内部就安全了。你还得管权限：这个 AI、这个员工，能看哪些数据、不能看哪些。该脱敏的脱敏（比如把身份证号、手机号该挡的挡掉），该限权的限权。

我们当时做账号风险处置，有一次批量处置了一批高风险账号。结果出事了：里面误伤了一批 MCN 机构和大 V，他们突然登录不了，投诉炸了锅。复盘下来，问题出在几个环节：该保护的名单只同步了一份、不全；识别的颗粒度有问题；再加上用户反馈的流程太慢，问题发生后，反馈隔了一段时间才到我们这儿，黄花菜都凉了。

这件事之后我们定了一条规则：任何严重处置动作，上线前必须跟相关业务方把信息同步到位，还得留足够长的反馈时间，让出问题的人能及时把声音传回来。这条规则，放到今天任何一个会自动做事的 AI 系统上，都适用，你让 AI 自动去动用户、动数据、动业务，就必须想好：它动错了，谁能第一时间发现、第一时间喊停。

关于数据安全，还有两个来自一线、也被研究印证的判断，帮你把这件事看得更立体。

第一，安全不该只被看成 AI 的绊脚石，它也可能变成你的护城河。很多老板一听数据安全就头疼，觉得是拖慢项目的负担。但在有些研究报告的案例里，安全要求从来没有真的“杀死”过一个项目；相反，那些被逼着把数据保护做扎实的公司，后来反而能接下竞争对手不敢碰的敏感生意，比如处理客户的金融数据、医疗记录。报告里有个银行的例子：它原来是任何数据不出防火墙，上 AI 后不得不改造，做法更细：数据出门前先脱敏，用假的姓名和金额替换真实的，云端只负责理解意图，返回时再把真实信息拼回来。这套机制一旦建好，后续所有要碰敏感数据的 AI 应用，都能复用它。所以安全投入是前重后轻的：第一次建起来费劲，但一旦建成，就成了你的资产、你的门槛，别人因为没这套东西，根本接不了你能接的活。这也顺带告诉你：私有部署之外，脱敏 + 用假数据替换 + 云端只处理意图也是一条又安全、又能用上云端强模型的路。

第二，你不管，员工也在偷偷用 AI，而且更危险。你是不是觉得我们公司还没正式上 AI，所以没有数据安全问题？恰恰相反，行业调研显示，七八成在工作中用 AI 的员工，用的是公司没批准的工具；其中过半的人，承认往这些工具里输入过公司的敏感信息。这叫“影子 AI”，你不给他们配安全可控的工具，他们就自己拿着公司机密，去喂外面的免费 AI，这才是最大的窟窿。有家半导体公司做内部安全排查，发现员工私下在用的 AI 工具，竟然有上千种。所以数据安全这件事，不是等我们正式上 AI 再说，是你现在就该管，与其堵，不如赶紧给员工配上安全、好用的官方工具，把他们从外面的野路子上引回来。

数据安全这一块，你作为老板要问的问题很简单：最敏感的数据，在这个 AI 项目里，会被发到哪儿去？谁能看到？漏了怎么办？这几个问题供应商和你自己的团队答不清楚，这事就先别上。

二、准确性：它一定会错，关键是错了接不接得住

首先，没有不出错的 AI。任何跟你拍胸脯说：我们的 AI 百分百准确、绝不出错的，都是在骗你。

所以，保准确性的核心思路是：承认它会错，然后设计一套机制，让它错得起、兜得住、纠得回来。准确性是这么“管”出来的，不是模型天生就有的。

把这套机制，拆成你能听懂、也能去要求的几条。

第一条：按“错了有多严重”来定准确率的标准。

不是所有场景都需要一样高的准确率。一个错误，后果越严重，要求就得越高。

我们做风控时，有一套很清楚的分级：动作越狠，准确率门槛越高。直接拦截一笔交易（后果重，误拦了用户骂街），准确率必须达到很高水平才敢上；而像加个验证码、提示一下这种轻动作（就算误判了，用户多验证一下也没大碍），门槛就可以低一些。

这个道理你直接能用：你公司哪个 AI 应用要是动作很重、错了代价大，就得用最高的准确率标准卡它，甚至卡完还得加人工复核；哪个动作很轻、错了无所谓，标准可以松，放手让 AI 多干。一刀切地要求所有 AI 都“必须百分百准”，既做不到，也是浪费。

第二条：让它在拿不准的时候，老实说我不知道。

这条我在前面几章反复讲，这里再提一遍，因为它是准确性的命门。

一个会说我不知道、转人工的 AI，比一个什么都敢答的 AI，安全一百倍。我们做的所有系统都有一条基本原则：知识库没覆盖的、它判断不了的，直接转人工，禁止它自己瞎猜、瞎补、瞎答。它说我不会，不是失败，是它知道边界、是设计得好的表现。

验收时一定要确认这一点：它遇到没把握的事，会停下来交给给人，还是会硬编一个答案？后者是埋雷。

第三条：给它的每个答案，留个出处。

让 AI 给结论的同时，告诉你这结论是从哪条知识、哪份资料来的。这样人能核对、能追溯，不是它凭空说一句你就得信。有出处的答案，你能验；没出处的答案，你只能赌。

第四条：用一套考卷定期考它，还要专门收集它的错题。

这是准确性能不能长期保持的关键，也是最容易被偷工减料的一环。

考卷（评测集）：准备一批有标准答案的问题，定期拿去考这个 AI，看它答对多少。没有这套考卷，你根本不知道它准不准，全靠用的人感觉。

错题本（badcase 机制）： 把它答错的、出问题的案例，专门收集起来、分析为什么错、然后去改进。一个有“错题本”的系统，会越用越准；一个没有的，会一直在同一个坑里反复栽。

你问供应商和团队一句：你们拿什么考它准不准？它答错的怎么收集、怎么改？答得上来的，是真把准确性当回事。

第五条：重要的事，人必须在关键环节把关。

这就是我反复讲的人机协同，AI 干初筛、初判、出初稿，人在关键的、要担责的那一步把关、拍板。准确性不光靠 AI 自己，更靠“AI + 人”这套组合。

前面讲人机协同时提过 Klarna 的案例，这里换个角度看同一件事：它的问题不只是“AI 能不能替代客服”，更是责任边界怎么设计。AI 可以接住高频、标准、低风险的问题，但一旦进入复杂投诉、敏感沟通、赔付判断，就必须有人接管。

这个反复，恰恰证明了那条基本原则：AI 不是一刀切地全替代人，它干它擅长的大批量、标准化的部分，把复杂的、高风险的、容易出错的那部分，稳稳地交还给人。准确性，正是在这个分工里守住的。

三、责任边界：出了事，得有人能背、有地方能兜

最后一件，也是你最关心的：出了事，谁负责？

这件事，绝大多数企业在上 AI 的时候，压根没提前想。等真出了事，才发现：AI 是个工具，它不会背锅，锅最后还是落在人头上，但落在谁头上、怎么落，事先没说清，就成了一笔糊涂账，部门之间互相推诿，最后伤的是公司。

所以责任边界，必须在项目上线之前就白纸黑字定清楚，不能等出了事再扯皮。

可以先思考这三个问题：

第一：AI 生成的内容，谁来审、谁签字？

AI 写的对外文案、给客户的答复、生成的报告，在它发出去、用出去之前，谁负责看一眼、谁负责签这个字？特别是对外的、有承诺性质的内容，必须有个人把最后一道关、并且为这道关负责。不能是 AI 写的，发出去就发出去了，那等于没人负责。

第二：自动流程出了错，谁兜底、谁能喊停？

如果你上了一个会自动做事的 AI 流程，你必须事先想清楚：它跑错了，谁来兜这个损失？以及更重要的，谁有权、有能力在它闯祸的时候，第一时间把它摁停？

前面讲的那个误伤高风险账号的事故，核心教训之一就是喊停太慢，问题发生后，反馈隔了一段时间才到。一个自动流程，必须配一个能快速发现问题、快速喊停的机制和负责人。没有这个，自动化跑得越欢，出事时滚得越大。

第三：这个 AI 系统，谁是它的“主人”？

每一个上线的 AI 应用，都得有一个明确的主人，一个具体的人或岗位，负责养它、管它、为

它的表现负责。没有主人的 AI 系统，是无主之物，出了事没人认领，平时也没人维护，慢慢就烂掉了。我们在前面讲知识库时说过没人养的知识库上线那天就开始烂，放到任何 AI 系统上都成立。谁是主人这个问题，上线前必须有答案。

把这三件事，变成你的一张上线检查表



安全、准确性与责任边界

数据安全、准确性、责任边界，讲完了。我帮你把它压成一张能直接用的检查表。任何一个 AI 项目要从“试试看”转成正式长期用之前，你拿这几问挨个过一遍，过不了的，先别上。

数据：最敏感的数据会被发到哪儿？谁能看？漏了怎么办？最敏感的是不是该走私有部署？

准确性：这个应用动作重不重？准确率标准卡到位了吗？它拿不准会不会老实说不知道？答案有没有出处？有没有考卷和错题本？关键环节有没有人把关？

责任：AI 出的东西谁审、谁签字？自动流程错了谁兜、谁能喊停？这个系统谁是主人？

老板该学到什么

这一章，如果你是拍板的人，我希望你带走三条判断。

第一，安全和准确不是买来的，是设计和管出来的。

别指望买个又安全又准的 AI 就万事大吉。真正靠谱的，是在 AI 周围搭一圈机制：数据分级、

权限管控、准确率分级、拒答兜底、出处可查、评测和错题本、人在关键处把关。这一圈机制，才是安全和准确的来源。

第二，最敏感的数据，要么别往外送，要么私有部署。

数据先分级，越敏感越往自己手里收。客户隐私、核心机密这一档，漏了就是不可逆的灾难，该花的私有部署的钱不能省。这是底线，不是成本问题。

第三，责任边界必须上线前定清，别等出事扯皮。

每个 AI 系统上线前必须回答三件事：它出的东西谁审谁签、自动流程错了谁兜谁喊停、这个系统谁是主人。这三件事没答案，这个项目就还没准备好长期用。

一句话

这章收束成一句话：**AI 的安全和准确，来自你在它周围搭的那一圈机制；而出了事谁负责，必须在线前就白纸黑字定清，谁审、谁兜、谁是主人。守住这三道，你才敢把 AI 从试试看放进核心业务长期用。**

第 23 章 谁牵头，各角色怎么分工

很多企业会犯同一个错：把 AI 转型这件事，一股脑交给 IT 部门。

老板的逻辑是：AI 是技术，IT 是管技术的，那就 IT 牵头。一纸任命下去，IT 负责人接了这个活，然后呢？然后大概率，半年后你会发现公司上了几个聊天机器人、装了几个 AI 插件，数字报上来挺好看，但你那些最耗钱、最依赖人的核心业务，一点没动。

这一章我们需要明确：AI 转型需要老板、业务、IT、HR 几方一起上，而且只有老板能压得动。谁牵头、谁干什么、怎么不让大家各干各的，这三个问题，你作为老板必须自己想明白，想不明白，这事就是散的。

为什么不能只交给 IT

把 AI 交给 IT，错在哪？IT 懂技术，但不懂业务；业务懂场景，但不懂技术。AI 落地恰恰卡在这两者中间。

我们在前面的章节里反复讲过一件事，AI 落地最难、最值钱的活，是把一个岗位上老师傅脑子里只可意会的判断，挖出来、写下来，变成 AI 能学的东西。我在某互联网大厂做支付客诉智能体的时候，这一步占了整个项目最大的工作量：找最资深的风险分析师，一笔一笔地过他处理过的案子，追问他你第一眼看哪个字段，看到这个值你为什么立刻判断是误拦。

这件事，IT 干得了吗？干不了。IT 工程师再聪明，他不知道一笔交易为什么被风控拦截，不知道这行的判断标准，不知道哪些经验是该沉淀的。这些只有业务方知道。

那全交给业务方行不行？也不行。业务方知道“对”的标准是什么，但他不知道怎么把这套标准变成 AI 能用的东西，不知道哪些活适合 AI、哪些 AI 干不了，不知道一个 Agent 落地需要打通哪些系统、接哪些数据。

所以，AI 落地这件事，天生就坐落在业务和技术的交界处。它需要一个既懂业务、又懂 AI 边界的角色把两边接起来。这个角色，大部分企业内部是缺的。你光把任务丢给 IT，等于让一个只懂半边的人去干一件需要两边都懂的事，从结构上他就做不成。

这也是为什么这么多企业的 AI 项目最后只剩下一堆小工具，因为只有“做个聊天插件”这种纯技术的活，IT 一个部门能独立完成。一旦涉及改流程、动业务，IT 就推不动了，他没那个权限，也没那个业务理解。于是大家都退回到自己能独立搞定的那一小块，加起来就是公司上了一堆边角料 AI、核心业务纹丝不动。

还有一个更深层次的原因，一家公司里，信息和决策是顺着层级一层层“翻译”下去的，老板的战略，翻译成中层的计划，再翻译成一线的执行；一线的情况，又一层层翻译、汇总、上报回

去。这条“翻译链”上，每一道转手都有损耗、有延迟、有走样。而 AI 能创造的大价值之一，恰恰是打掉这条翻译链上的损耗，让信息和决策更直接地流动。但问题来了：这条翻译链，恰恰是中层、是 IT 这套现有体系在维护的。你让维护翻译链的人，去主导一件消灭翻译链的事，这在利益上是拧着的，他凭什么去革自己的命？所以 AI 转型这种要打破现有信息流、权力流的事，只能由站在翻译链顶端、不靠这条链吃饭的人，也就是一把手，来推。交给 IT、交给中层，等于让既得利益者去拆自己的台，推不动是必然的。

各角色分别做什么

AI 转型是个多角色协同的事，我把它拆成五个角色，每个角色干一件谁也替代不了的事。你可以对照着看自己公司这几个位置上有没有人、做没做到位。

老板：定方向、配资源、压组织

这是我在“一把手工程”那一章专门讲过的三件事，只有你能干。定方向，AI 转型先打哪个业务、要什么结果；配资源，给钱、给人、给时间，而且要给到位，不能一边喊重视一边一分钱不掏；压组织，当业务和技术吵起来、当某个部门不配合、当中层为了怕担责而拖延的时候，只有你能拍板压下去。

这三件事，任何下属都替你干不了。你指望一个 IT 总监或者一个项目经理去“压组织”，他压不动，因为他没那个位置。AI 转型里一切跨部门的僵局，最后都得回到你这儿解。

业务方：定场景、出标准、验结果

业务方是 AI 落地真正的甲方，虽然他不掏钱。他要回答三个问题：我这个部门哪个环节最耗人、最重复、最值得上 AI（定场景）；这件事干得对的标准是什么、老师傅是怎么判断的（出标准）；做出来的东西到底好不好用、能不能解决我的问题（验结果）。

IT：做支撑、保安全、管接入

IT 的活也很关键，但它的位置是支撑，不是主导。系统怎么接、数据怎么打通、权限怎么控、哪些数据能上云哪些必须本地、出了安全问题怎么兜，这些是 IT 的主场。一个 Agent 要能进 CRM、ERP、OA 里读信息、触发流程，这些管道的活只有 IT 能干。

把 IT 摆在支撑位，不是降低它，是让它干它最擅长、也最该干的事。它不该被推到前台去定业务场景，那不是它的活。

HR：推机制、抓培训、调激励

很多老板做 AI 转型完全忘了 HR，这是个大漏。员工愿不愿意用 AI、用了之后有没有好处、不用会不会落后，这些全是机制问题，是 HR 的主场。培训怎么组织、火种用户怎么选怎么奖、AI 提效之后绩效怎么调、岗位职责怎么重新定义，这些不解决，工具做得再好也没人持续用（这是下一章要专门讲的）。

数字化/转型负责人：统筹方法论、对齐目标、管节奏

如果你公司有专门的数字化团队或者转型办公室，他们的活是穿针引线，把上面四个角色串起

来，统一方法论、对齐各方目标、管整体节奏。注意，这个角色是协调者，不是执行者，更不是替老板做决策的人。

我把这五个角色整理成一张表，你可以直接拿去对照：

角色	核心职责	一句话
老板	定方向、配资源、压组织	跨部门僵局的最终拍板人
业务方	定场景、出标准、验结果	AI 落地真正的甲方，不参与必败
IT	做支撑、保安全、管接入	管管道，不主导业务
HR	推机制、抓培训、调激励	让员工愿意用、持续用
数字化负责人	统筹方法论、对齐目标、管节奏	穿针引线，不越位决策

这张表你可以对照着数一遍：这五个位置，你公司现在分别是谁？有没有空位？大部分企业一数就会发现：业务方那一栏经常是空的（没人真正下场），数字化负责人那一栏也经常是空的（没人统筹），最后五个角色的活全压在 IT 身上。压不动，是必然的。

怎么避免各自为战

角色分清楚了，还有一个更隐蔽的坑：就算每个角色都有人、都在做事，他们还是可能各干各的，拼不成一个整体。

为什么会这样？

很多企业上 AI，半年下来做了几十个 AI 小工具，客服部做了个问答机器人、市场部做了个文案助手、行政做了个会议纪要工具。每个工具单独看都挺好，调用次数报上来很好看。但你回头一看公司的核心业务流程：一个都没改。

为什么？因为改流程是个又难又得罪人的活。要动别人的工作方式、要协调好几个部门、做砸了还要担责。而做个小工具呢？简单、出活快、数字好看、谁也不得罪。

组织的激励机制，奖励的恰恰是后者。年底汇报，你说：我牵头改造了订单履约流程，卡了两个月还没完全跑通，听上去像在找麻烦；你说：我部门上线了 8 个 AI 工具，日均调用 3000 次，听上去像个先进典型。于是每个部门负责人都做了一个对自己最理性的选择：挑软柿子捏，做小工具，躲开硬骨头。每个人都理性，加起来就是整个公司的非理性，半年过去，工具一堆，流程没动，转型原地踏步。

这个局，部门负责人自己解不开。因为让他去碰流程、去协调别的部门，既超出他的权限，又不符合他的利益。这个局只有一把手能解。

怎么解？三件事：

第一，设一个跨部门的统筹机制，别让 AI 转型散落在各部门。可以是一个专项小组，可以是一个挂在你直接领导下的转型办公室，关键是要有一个超越单个部门的角色，来管哪些该做、先做哪个、谁配合谁。前面那张表里的数字化负责人就该坐在这个位置。这个机制直接对你汇报，这样它才压得动各个部门。

第二，目标要对齐到业务结果，不是对齐到用了多少 AI。这一点极其关键。如果你考核各部门的指标是上了几个 AI 工具，调用了多少次，那你就是在亲手鼓励大家刷漂亮数字、躲硬骨头。你要考核的是业务结果，这个流程的处理时长降了多少、这个岗位的人效提了多少、这块的成本省了多少。指标定在哪，大家的力气就往哪使。你想让大家去啃流程这块硬骨头，就得让啃硬骨头的人得到承认。

第三，你自己得在场。不是挂个名当组长，是真的在关键节点出现，定方向的时候你在、各部门吵起来的时候你拍板、有人想躲硬骨头的时候你点破。有数据说明，一把手持续参与的 AI 项目，成功率远高于中途撒手的。原因很简单：跨部门的协同，本质是利益的重新分配，而利益分配这件事，公司里只有你一个人有最终的话语权。

老板该学到什么

这一章，我希望你带走三条判断。

第一，AI 转型不是 IT 的活，是一把手领着多个角色一起干的活。把它全交给 IT，等于让一个只懂半边的人去干一件两边都得懂的事。结果一定是只剩下一堆小工具，核心业务一动不动。

第二，五个角色，各守其位，业务方千万别缺位。老板定方向配资源、业务方定场景出标准、IT 做支撑、HR 推机制、数字化负责人统筹。这五个位置你公司缺了谁，补上谁。最容易缺、又最致命的，是真正下场的业务方。

第三，别让大家各刷各的漂亮数字。用业务结果而不是用了多少 AI 来定目标，设一个直接对你汇报的统筹机制，并且你自己要在场。否则每个人都会理性地挑软柿子捏，公司层面就是集体躲开了真正该改的东西。

一句话

这一章你记住一句：**AI 转型卡住的时候，十次有九次卡的不是技术，是组织没把角色摆对，没人压得动跨部门的事。**

而这件事最难的地方在于，你越是熟悉自己的公司，越容易默认现有的部门划分够用了，直接把 AI 塞进现有的组织架构里，塞给那个看起来最该管的部门。但 AI 转型恰恰是个会打破现有部门边界的事。角色怎么重新摆、机制怎么压得住、目标怎么从用了多少 AI 扭到换来什么业务结果，这些问题如果不在一开始说清楚，后面再多工具都会散成一地。

第 24 章 怎么让员工持续用起来

公司花了几万块请人做了一场 AI 培训，讲师讲得热闹，员工听得点头，现场气氛很好，大家都觉得开了眼。培训结束，老板觉得这事算成了。

两周以后你再去看，没人用了。该用 Excel 的还在用 Excel，该手写会议纪要的还在手写，那些培训里教的工具，装都没装。钱花了，认知好像提了一点，但工作方式一点没变。

这个场景，我们在服务客户的过程中见过无数遍。这是一个老大难的问题：做出来，和用起来，是两件完全不同的事。你能把工具做出来，能把员工召集起来培训，但你很难让他们持续地、自发地把 AI 用进每天的工作里。

这一章就讲这件事：怎么让 AI 在你的企业里活下去，而不是上线即死。我还会把一个相关的问题并进来讲，AI 时代，你的企业到底需要什么样的人。

为什么工具上线就死

工具没人用，老板第一反应往往是员工抵触新东西。这个判断基本是错的，员工不用，几乎都是有具体原因的，而且这些原因大多怪不到员工头上。

我把工具上线就死的原因拆成五条，你对照看看自己中了几条：

第一，入口不顺。一个 AI 工具，如果员工要专门打开另一个网页、再登录一次、再切换界面才能用，那它基本就死定了。人是有惰性的，工具但凡比原来的工作方式多两步，大家就退回老路。真正能活下来的 AI，都是嵌在员工本来就在用的系统里，在他天天打开的 OA 里、CRM 里、飞书钉钉里，顺手就能用。入口的摩擦力，直接决定生死。

第二，场景不真。很多 AI 工具是为了上 AI 而上 AI，解决的是个伪需求。员工试一次，发现这件事没解决他真正头疼的事，就再也不碰了。AI 必须切在员工真正觉得烦、累、重复的那个点上，他才有动力用。

第三，知识质量差，答得不靠谱。员工对一个 AI 工具的信任，极其脆弱，一次不靠谱的体验，就足以让他永久放弃。他问了一个问题，AI 答得驴唇不对马嘴，或者一本正经地胡说，他心里就给这工具判了死刑：不行，不能信。之后你再怎么优化，他都不会回来了。所以宁可让 AI 在拿不准的时候老老实实说这个我答不了，请找人工，也不要让它硬答、答错。

第四，习惯没变。工具是新的，但人的工作习惯是旧的。没有人推着、没有新的 workflow 强制，员工会自然地滑回原来的做法。改变习惯是反人性的，光给个工具不够。

第五，没人推、没人管。工具上线那天最热闹，之后就放任自流，没人运营、没人答疑、没人盯使用情况。AI 工具不是上线就完事的家电，它需要持续地养。没人养，它就慢慢死了。

有一项调研发现，大约 85% 的员工没法把 AI 培训里学到的东西真正用到自己的岗位上（Docebo 的调研）。注意，不是没培训，是培训了也用不起来。这恰恰说明：认知和落地之间，隔着一条很宽的沟。培训解决的是认知，跨过那条沟，得靠机制。

让员工持续用起来的一套机制

知道了为什么死，就知道怎么让它活。我把让员工持续用 AI 的打法，整理成一套可落地的机制。这套东西不复杂，但每一条都得真做，少做一条，效果就打折。

第一，培养火种用户

这是整套机制里最重要的一条，你不要指望全员一夜之间都会用 AI、爱用 AI，那不现实。你要做的是先找到一小撮火种，那些本来就对 AI 有兴趣、学得快、愿意折腾的员工，把他们重点培养起来，让他们先用爽、用出成果。

火种用户的价值有两层，第一层，他们是活的样板。同事看到身边一个跟自己干一样活的人，用 AI 把原来要两小时的事半小时干完了，这个冲击力，比老板开十次动员会、请十个讲师都管用。人是会跟身边人比的，最有效的推广，是他能干我也能干的同时带动，不是自上而下的号召。第二层，他们是种子教练。火种用户能就近帮同事答疑、分享自己的用法，把 AI 的使用从一个人扩散到一群人。

我们服务客户时，无论是给某头部电商平台的财务部做、还是给某四大行之一的多个部门做，核心动作都是一样的：不是把所有人拉到同一个水平，是先把一批人提上去，让他们成为带动其他人的火种。

员工用 AI 的能力是分层的，从最低的只会跟 AI 一问一答，到能搭 workflow 批量解决问题（这套能力分层，我在第 5 章讲培训目标时拆过，这里不重复）。火种，就是你要优先往上送几个台阶的那批人：先把他们从只会聊天提到能把 AI 用进具体工作、能提效那一层，一个月下来不少的活能用 AI 提效，再往上继续陪跑一段时间效率翻倍也看得见。不追求所有人都到顶，先把火种顶上去。

第二，把内部案例攒起来、传出去

外面讲一百个别人家的 AI 案例，不如你自己公司一个案例有说服力。你公司哪个员工用 AI 解决了一个具体问题、省了多少时间、效果怎么样，把这些攒成一个个内部案例，在公司里传播。员工看到的是我们公司、我同事、我熟悉的业务，代入感和可信度完全不同。

第三，把使用挂到激励上

这一点得讲究方式，直接考核：你这个月用了多少次 AI，会逼出一堆刷数据的假动作（我在上一章讲组织激励时专门提过这个坑）。正确的挂法，是挂到结果和示范上，谁用 AI 实打实提了效、出了成果，给承认、给奖励；谁成了火种、带动了一片人，给荣誉、给资源。让用得好的人尝到甜头、有面子，比罚不用的人有效得多。

但这里还有一个更隐蔽、也更致命的激励问题，你必须看见：很多员工不用 AI，不是懒，是

怕。一个客服用 AI 把工单处理速度提了三倍，本来是好事。可如果他的考核还是每天处理多少单，那他提速之后，要么只是更早地闲下来、显得无事可做，要么老板看他活轻了，开始琢磨是不是可以少要几个人。你要是这个客服，你还敢拼命用 AI 提效吗？当用 AI 把自己变高效的结果，是自己可能被优化掉，员工最理性的选择就是：偷偷藏着不用，或者用了也不声张。

所以光鼓励没用，你得从两头同时动：一头，把考核从干了多少量改成创造了多少价值，让提了效的人腾出手去做更值钱的事，而不是更早下班等着被裁；另一头，要明明白白告诉员工，AI 替代的是你本来要招的人，不是你现在已有的人；它让你现有的人，顶得上原来好几个人的活。把这层窗户纸捅破，员工心里那块石头落了地，才会真正放开手去用。

第四，留一个顺手的答疑入口

员工用 AI 的过程中一定会卡壳，不会问、问不出想要的、遇到报错。如果这时候他找不到人问，卡两次就放弃了。你要留一个低门槛的答疑入口，可以是火种用户兼任，可以是一个内部群，关键是让员工卡住了马上有人能问。

第五，定个简单的使用规范

哪些数据不能往 AI 里喂、AI 生成的东西谁来把关、什么场景能用什么场景不能用，这些得有个简单清楚的规矩。不是为了限制，是为了让员工敢用、用得安心，也防止出事。规范要简单到一张纸能讲完，太复杂没人看。

第六，定期复盘，持续养

每隔一段时间回头看一次：哪些工具在被用、哪些没人碰、员工卡在哪、哪个场景效果最好。把死掉的工具砍掉，把活的工具优化，把好的用法推广。AI 工具是要持续运营的，不是上线就能撒手的。这件事一停，工具就开始往下滑。

这六条，核心就一句话：让 AI 在企业里活下去，靠的不是一次性的培训，是一套持续运营的机制。培训是入口，机制才是让它活下去的氧气。

AI 时代，企业要什么样的人

讲完怎么让员工用起来，得往深处再走一层：AI 进来之后，你的企业到底需要什么样的人？这个问题，直接关系到你该重点培养谁、该把谁放到火种的位置上。

我的判断是：AI 时代，人的价值正在从会做事转向会指挥 AI 做事。具体说，有四种能力会变得越来越值钱：

会提问。AI 的输出质量，很大程度上取决于你怎么问它。同一个 AI，会问的人能从它那里榨出十倍的价值，不会问的人觉得它也就那样。会把一个模糊的需求，拆成清晰、具体、AI 能执行的指令，这是新的核心能力。

会判断。AI 会给你一个答案，但这个答案对不对、能不能用、哪里有问题，得人来判断。AI 越强，会判断的人越值钱，因为只有他能识别 AI 什么时候在胡说。

会拆任务。一件复杂的事，怎么拆成 AI 能一步步完成的小任务，哪些交给 AI、哪些自己来，

这种任务编排的能力，是把 AI 从玩具变成生产力的关键。

会跟 AI 协作。知道 AI 的边界在哪、什么活该让它干、什么活不能信它，把 AI 当成一个有能力但也会犯错的协作者来用，而不是当成一个无所不能的神，也不是当成一个不靠谱的玩具。

这四种人，就是你最该重点培养、最该放到火种位置上的人。他们不一定是技术最强的，但他们是 AI 时代的高产出者。

温馨提示：别让你的员工把“判断”也外包给 AI。

AI 能帮人做事，这是好事。但有一个危险的滑坡：执行交给 AI 没问题，可一旦连判断也交出去，员工不再自己思考对不对、好不好，AI 说什么就是什么，长期下来，团队会慢慢失去判断 AI 输出对错的能力。越依赖，越不会；越不会，越只能依赖。等哪天 AI 出了个看似合理实则错误的结论，整个团队没有一个人能识别出来，那就出大事了。

有个同行把这个现象叫“能力萎缩陷阱”，你把一项判断越多地让渡给 AI，你自己在这件事上的能力就越退化，直到有一天，你既离不开它、又再也没能力检验它到底对不对。这是个慢性的、不容易察觉的过程，等你发现的时候，往往已经晚了。所以这件事得提前防，不能等出了事才想起来。

AI 可以接管“手”的工作，但“脑”的工作，尤其是判断这一环，必须留在人这里。你培养员工用 AI，不是培养他们变成 AI 的传声筒，是培养他们成为能驾驭 AI 的人。这两者的区别，长期看就是你这个组织有没有竞争力的区别。

所以你在推 AI 的时候，要同时往里加一条文化：用 AI，但别停止思考。鼓励员工质疑 AI 的输出、核对 AI 的来源、对重要的结论保持自己的判断。这条看起来跟多用 AI 矛盾，其实是一体两面，只有保持判断力的人，才能真正用好 AI。

老板该学到什么

这一章，三条判断带走：

第一，做出来不等于用起来，工具上线就死是常态。入口不顺、场景不真、答得不靠谱、习惯没变、没人运营，逐条排查，你的工具大概率中了好几条。

第二，让员工持续用，靠的是机制不是培训。火种用户是核心打法，先把一小撮人提上去，让他们用出成果、带动一片。再配上内部案例、结果激励、答疑入口、使用规范、定期复盘。培训是入口，机制是氧气。

第三，AI 时代，会提问、会判断、会拆任务、会协作的人最值钱；但千万别让团队把判断也外包给 AI。手可以交给 AI，脑必须留在人这里。一个停止思考、全盘信 AI 的团队，是没有未来的。

一句话

这一章你记住一句：**AI 工具最大的敌人，不是它不够聪明，是它上线之后没人养、没人用，做起来只是开始，让它在组织里活下去，才是真功夫。**

而让它活下去这件事，最难的不在工具本身，在于你得改变一群人的工作习惯、还得在组织里点起几个能带动别人的火种。这件事急不得，也很难靠几场培训一蹴而就。老板要盯住的，不只是工具有没有上线，还要看机制有没有跟上、人有没有被带起来、团队的判断力有没有被保留下来。

第 25 章 不同部门和行业，应该从哪些 AI 场景切入

到这一章，你心里大概已经有了一个很具体的问题：道理我都懂了，那我这个公司，到底先从哪个部门、哪个场景下手？

这是个好问题，但它有个陷阱。很多老板希望我直接给一张表，你是制造业，那你就做这几个；你是餐饮，那你就做那几个。如果我真这么给你，你照着做，大概率会做歪。

因为先做哪个，没有标准答案。它取决于你公司的具体情况，哪个部门的活最耗人、哪块的数据最干净、哪个负责人最愿意配合。同样是做财务 AI，一家财务流程标准化程度高的公司，和一家全靠老会计手工对账的公司，该切的点完全不同。

所以这一章我不给你一张照着抄的清单。我给你一套判断逻辑，教你自己看，哪个部门、哪个场景该先切。至于各主要职能部门具体能切哪些场景，我整理成了一份查阅用的速查表，放在附录里，你需要的时候去翻就行。这一章，我只谈判断。

哪些部门最容易出成果

AI 最容易出成果的地方，是那些“高频、重复、有数据、结果能验证”的活。四个条件占得越多，越该先切。我们拆开讲这四个条件：

高频，这件事每天大量发生。频次越高，AI 一旦做好，省下的总量越大。一天处理三笔的活，AI 做得再好也省不出多少；一天处理三千笔的活，AI 提效一点点，乘以三千就是巨大的量。

重复，这件事的处理方式高度雷同，不是每次都要重新动脑子。重复度高，意味着规律可循，AI 学得会。

有数据，这件事有现成的、相对干净的数据或文档可依托。AI 不是凭空变出答案的，它得有东西可学、可查。

结果能验证，做得对不对，能比较容易地判断出来。这一条很关键：一个活如果“好不好”没法判断，AI 做出来你也不敢用、不知道该不该信。

拿这四个条件去套你公司的各个部门，有几个部门通常会亮起来：

财务，几乎是首选。财务的活天然符合这四条：费用报销审核、发票核验、对账、合同条款检查，全是高频、重复、有规则、结果能对。一张报销单合不合规，有明确标准，AI 判断完人能复核。所以财务是我见过 AI 落地成功率最高的部门之一。行业数据也印证这一点：超过六成的财务部门已经在用 AI（Gartner 2025 的调研），它是最早被 AI 渗透的职能之一。

我服务过一家电商公司，专门给他们的财务部做 AI 提效。为什么是财务？除了财务的活适合 AI，还有一个更深的原因，这家公司的老板意识到，财务这个部门如果不主动用 AI 进化，是有被时代淘汰的风险的。

客服，典型的高频重复场景。用户的问题大量重复，同样的几十类问题，换个人来问又是一遍，这种活 AI 接得很好。有研究显示，客服用上 AI 之后，平均每小时能多解决 14% 的问题，新手客服的提升甚至能到 34%（NBER 的研究）。我在某互联网大厂做的支付客诉智能体，把客服一个原来要走完整流程的问题压到小时级，就是这个逻辑。

内容相关的岗位，AI 能重构整条生产线。市场、运营、电商内容这类，要持续产出大量文案、图片、视频。AI 在这里的价值不只是帮你写一条文案，而是能把热点采集→内容生成→审核→发布整条流水线重构掉。我们帮一家做跨境电商带货的客户搭过内容工厂，把单条内容的生产成本从原来的几十块美金降到了个位数，日产量翻了几十倍。

这三个部门，财务、客服、内容，是大多数企业最该优先考虑的切入点。因为它们活够重复、数据够现成、结果够能验证。AI 落地最稳的入口，从来是这种不起眼的地方，不是那些听上去很炫的大场景。

哪些部门看着热闹，其实难落地

反过来，有些部门听上去很适合上 AI，老板也最容易心动，但实际一做就翻车。

责任太重、错了代价大的活，先别让 AI 直接拍板。比如涉及法律风险的决策、涉及大额资金的判断、涉及人身安全的环节。这些活 AI 可以碰，但不能独立拍板，它可以辅助分析、可以做初筛，最后那一下必须人来。

强人际博弈的活，AI 干不了。比如谈判、复杂的客户关系维护、需要察言观色随机应变的事。这些依赖人与人之间的微妙判断，不是规则能描述的，AI 在这儿使不上劲。数据极度分散、责任边界不清的活，先别碰。如果一件事的相关信息散落在十几个人的脑子里、几十个 Excel 里、谁也说不清归谁管，那它现在还不到上 AI 的时候，你得先把数据和责任理顺，AI 才有立足点。强行上，只会把混乱放大。判断一个场景该不该让 AI 直接答，有个简单的提问框架。遇到下面三类问题，AI 只能辅助，不能拍板：

需要业务判断或要做出风险承诺的（比如这单能不能放款）；

信息来源不清、出了错没人能负责的；

这块的知识体系本身还没理顺、连人都说不清标准的。

这三类，你一旦设成让 AI 直接回答，员工问一次发现不靠谱，就再也不信这个 AI 了。一次糟糕的体验，就足以毁掉一个工具。与其让 AI 在不该它管的地方逞强，不如老老实实划清它的边界。

钱该先砸哪个：给问题标个价

讲到这，你可能发现一件事：符合高频、重复、有数据、结果能验证的场景，在你公司可能不止一个。财务能做，客服能做，内容也能做，那先做哪个？

这里不要重新发明一套判断方法，回到第 15 章讲过的“问题标价法”就够了：不是看哪里能用 AI，而是看哪个问题最值钱、同时现在最能落地。

你要问的是这几个问题：

这个问题一年能省多少钱、少损失多少钱，或者多赚多少钱？

它是不是高频发生，而不是偶尔才出现一次？

数据、流程、负责人是不是已经具备，还是做之前还要补一大堆地基？

解决它以后，结果能不能被看见、被验证、被老板和团队认可？

所以，不同部门切入，也不是谁声音大就先做谁，也不是哪个部门看起来最时髦就先做哪个。先按价值和成熟度排一遍：价值高、成熟度也高的，先做；价值高但成熟度不够的，先补条件；价值低但容易做的，用现成工具或培训解决，别专门立项。

一句话：部门只是入口，标价才是排序方法。第 15 章讲的是完整方法，这里你只要记住，选部门、选场景，最后都要落回“值多少钱、能不能落地”这两个问题。

不同行业，特性不一样

部门的判断逻辑讲完了，再说行业。不同行业做 AI，切入点的特性差别很大：

制造业：重在流程和质检。AI 的着力点常在质量检测、设备巡检、工艺流程优化、合规审核这些环节。这个行业数据相对结构化，但历史数据往往脏乱（图纸、报告、表格什么格式都有），清洗是大头。

餐饮连锁：重在门店标准化。难点是怎么把总部的经营经验、菜品推荐逻辑，标准化地复制到每一家店、每一个流动性很大的服务员身上。AI 在这里能把依赖人的经验变成系统稳定输出。

电商/跨境电商：重在内容运营和数据分析。内容生产、选品、投放优化是主战场，AI 能重构内容生产线、能从海量数据里找规律。

金融：重在安全合规。这个行业 AI 渗透率最高，但每一步都被安全和监管框着。数据能不能出本地、AI 的决策能不能审计、责任怎么界定，是比 AI 聪不聪明更前置的问题。

政务/体制内：重在材料流程和合规。大量的活是文件起草、材料审核、流程办理，AI 提效空间大；但保密要求高、采购流程长、对“稳”的要求远高于对“新”的要求，落地节奏完全不同。

最后补一组数据，帮你对不同行业的 AI 落地速度有个谱。有报告统计了各行业企业 AI 采纳的增长速度，差异很大：科技行业最快（同比约 11 倍），其次是医疗（约 8 倍）、制造（约 7 倍）；而金融、专业服务这些，虽然增速没那么炸裂，但绝对体量最大、用得最深。值得注意

的是医疗，它是监管最重的行业之一，却跑在最前面。这恰恰说明一件事：强监管行业也能上 AI。只要把合规这道坎迈过去，后面的落地反而又快又深。所以你别拿我们这行管得严当不做 AI 的借口，管得严，反而意味着你把安全合规这道护城河建起来之后，对手更难追上你（这一点我在讲安全那章已经说过）。

特意把行业特性放在最后、而且只用一段话带过，是想强调：行业特性决定的是切入点的色彩，但不改变前面那套底层逻辑。不管你是哪个行业，该问的还是那几个问题，哪个环节高频重复有数据？哪个问题最值钱？哪个能落地、阻力最低？行业知识帮你更快找到候选场景，但最终拍板靠的还是这套判断，不是靠我们行业别人都在做什么。

别盯着同行在做什么照搬。我见过很多老板，听说同行上了某个 AI 系统，自己也急着上一个一模一样的，结果完全不适配自己的情况。同行的做法是参考，你自己的问题标价单才是答案。

老板该学到什么

这一章，三条判断带走：

第一，先切高频、重复、有数据、结果能验证的活，财务、客服、内容通常是首选。越重复的地方，AI 落地越稳。别一上来就盯着那些听上去很炫的大场景。

第二，责任重、博弈强、数据散的活先别碰，尤其别让 AI 在该人拍板的地方拍板。需要业务判断、来源责任不清、知识没理顺，这三类问题，AI 只能辅助。一次不靠谱的体验，就足以毁掉员工对工具的信任。

第三，选先做哪个，靠的是给问题标价，不是看 AI 能干什么。从哪个问题最值钱出发，在值钱和能落地之间找那个交叉点。最值钱的问题，往往是你一直没注意到的那一个。

一句话

这一章你记住一句：**别问 AI 能用在哪儿，问我哪个问题最值钱，前者让你做出一堆鸡毛蒜皮的小工具，后者让你把 AI 砸在真正能改变经营结果的地方。**

而把公司里每个问题标准地标上价、再判断哪个能落地、该先动哪个，这件事说起来简单，真做起来，需要有人系统地把你整个业务过一遍。你企业最该做 AI 的地方，往往不是你坐在办公室里最先想到的那一个。凭感觉上项目，和系统诊断一遍再决定，差的可能就是这个项目最后是赚是赔。这一遍诊断，最好有个见过足够多企业、知道该怎么给问题标价的人，陪你一起过。

结语 AI 转型的终点，是组织能力

这本书翻到这里，如果你只记住一句话，我希望是这一句：AI 的价值从来不会自动发生，工具越像一个能做事的人，你的企业越需要先具备让它持续产生价值的能力。

回看全书：有一条暗线，贯穿始终

前面六篇，看上去讲的是不同的事：认知、案例、坑、怎么开始、怎么判断、组织。但它们底下藏着同一条线，我管它叫承接能力，企业把 AI 接进业务、并让它持续安全可控产生价值的能力。

把这条暗线给你抽出来：

你得说得清，AI 该干哪件事。哪个岗位、哪个流程、哪个环节适合 AI 介入，要它减什么成本、提什么质量。说不清，AI 项目就变成一场大型许愿，所有人都在讨论答案，没人对齐问题，这是任务定义能力。

你得把知识理清楚。企业的问题从来不是数据太少，是数据太乱，知识锁在老员工的脑子里、散在群聊和旧文档里。没整理就接 AI，只会把混乱放大，这是上下文整理能力。

你得让 AI 能进到系统里做事。真实的工作发生在 CRM、ERP、OA、邮件里。AI 要能进去读信息、触发流程，才产生效率，否则它就是个聊天框，这是工具连接能力。

你得有人养它。AI 早期一定会出错，关键是有没有人盯着、把每一个错误变成下一轮的改进。没人养，它就永远停在玩具阶段，这是反馈迭代能力。

你得说得清谁负责。AI 生成的回复谁审、自动流程出错谁兜底、敏感数据怎么管，这些决定了你的 AI 能不能从试试看进到长期用。这是治理与责任能力。这五种能力，横跨了这本书的好几篇。

它是整本书的骨架。你回头会发现：讲知识库那章在补第二种，讲智能体那章在补第三种，讲坑和经验那几章在补第四种，讲安全责任那章在补第五种。每一篇，都在帮你长出一种承接能力。

这也解释了一个让老板困惑的现象：为什么买了同样的工具、用了同样的模型，有的企业做出了成果，有的企业砸了一堆钱什么都没留下？差别不在工具，工具是一样的、是买得到的。差别在承接能力，一个有承接能力的组织，AI 进来就能生根；没有承接能力的组织，AI 进来就是水过地皮湿。

外面有一组数据：全球已经有近九成的组织在用 AI，但真正把它用出规模化价值的，只有大约三分之一（麦肯锡的调研）。还有研究发现，企业 AI 落地的挑战，七成来自人和流程（BCG

的判断)。很多企业已经开始用 AI，但只有一部分真正用出了价值，中间的差距，差的就是承接能力。

终点是组织能力，不是工具

所以，AI 转型的终点到底是什么？

不止是人人会用 AI 工具，你公司每个人都装了 AI 助手、每天都在用，这很好，但这只是起点，不是终点。

前面几章反复讲的那条线，到这里可以收束成一句话：AI 转型的终点，不是员工各自变快，而是 AI 进入你的数据、流程、组织和经营方式，成为这家公司运行方式的一部分。

这条路不好走

承接能力，说得清任务、理得清知识、连得上系统、有人养、有人负责，没有一样是企业自己能轻松补齐的。它们都是慢功夫、笨功夫、磨人的功夫。整本书的那些判断，听上去都不复杂，但每一条真要落到你自己的公司里，都要趟过一堆的坑。

而且这件事有个难处：你越熟悉自己的业务，越容易看不清哪些是该沉淀的经验、哪些环节最值得动 AI。这不是因为你不够聪明，恰恰是因为你太熟了，熟到那些问题在你眼里都成了本来就这样。当局者看不清的东西，需要一个在外面趟过很多企业、又懂 AI 边界的人，陪你坐下来，把你的业务重新过一遍。

这本书，没办法给你一套照着抄就能成的公式，AI 转型没有这种公式。别问要不要上 AI，可以先问够不够格用好 AI；**流程不清，AI 一定失败；别贪 100%，守住 80/20；最值钱的活，是把老师傅脑子里的经验挖出来；AI 转型的终点，是组织能力。**

如果合上这本书，你脑子里能留下这几句话，然后在做下一个 AI 决策的时候，它们会自己跳出来提醒你，那这本书就值了。这条路不好走，但它是对的路。你不需要一口气把所有能力补齐，先从一个场景开始，把任务、知识、工具、反馈和责任一层层接住。

AI 不缺想象力。缺的，是你的企业接住它的那双手。把这双手练出来，就是这本书最后想说清楚的事。

附录 企业各岗位 AI 应用场景速查表

先说清楚这张表怎么用。

它是一面镜子，你对着自己公司的组织架构扫一遍，在每个岗位、每个环节停一下，问一句：这件事，AI 是不是也能在我们这儿干？凡是你看到后觉得这个我们也天天在做的，就是最值得先看一眼的地方。

每一条只有三栏：场景 / 喂什么进去、出什么来 / 替你省了什么。看不懂技术没关系，你只需要判断一件事，这个活，在我公司值不值得让 AI 接。

下面十一个岗位，基本覆盖了一家公司从前台到后台的日常。

一、财务 / 会计

财务是 AI 落地看得见效果最快的地方：活高频、规则清、数据现成，做出来老板当场就信。我不是财务出身，为了接住几场财务培训，把财务这条线上的 AI 场景从原始单据一路梳理到经营决策，下面按财务工作流的顺序排了一遍。其中有的我在培训现场真演示过，有的重度依赖内部数据和系统接口、目前还是评估方向；具体哪些能落地，要结合你公司财务的实际来判断。

场景	喂什么 → 出什么	替你省了什么
发票识别与查验	一摞发票图片 → 结构化数据（号码、金额、税号、明细逐条）+ 真伪与合规校验	逐张手抄手查变批量自动出表
自动生成记账凭证	一句业务描述（谁付款、货发没发）→ AI 理解语义判断借贷科目、生成分录	不靠死板的字段映射规则，按业务含义记账
报销合规审核	报销单 + 公司制度文档 → 逐项合规判断（超标、缺票、违规），把握不足自动转人工	把对制度、核标准从逐单人工变自动初审
杂乱流水整理成表	格式混乱的报销/流水文本 → 统一规范表格（日期归一、缺失补全）	脏数据秒变规整，可直接导出带筛选的 Excel
银行流水 / 多流对账	银行流水 + 业务收支（或实物、资金、发票、订单多套数据）→ 自动匹配 + 差异清单	流水来了自动逐笔比对，免人工对账

企业 AI 转型白皮书

多主体报表合并对账	几张子公司报表 → 合并表 + 差异标红 + 可视化报告	跨子公司对账自动找差异点，几分钟搞定核对
交易流水清洗	几千行含重复、含退款的交易 CSV → 去重表 + 异常标记 + 虚增金额量化	几千行秒级去重，算清重复带来的金额偏差
杂乱文件夹核对	混着合同、订单、发票的文件夹 → 多张结构化表 + 数据对不上的点	散落多格式文件一次结构化，连不一致都点出来
月度费用环比分析	各月部门费用表 → 带预警颜色的 Excel（环比超 30% 标红）+ 文字结论	自动揪出某部门差旅 +87% 这类异常科目
深度经营分析	公司背景 + 经营目标 + 数据 → 多维度经营分析（收入/费用/利润，对比行业）	不只是算数，像懂业务的人那样分析
经营数据出汇报 PPT	一张经营数据表 → 商务风汇报 PPT（指标卡 + 趋势图 + 风险建议）	月度/季度汇报不用手搓，直接拿去开会
Excel 自动标色	数据表 + 一句话规则（超 120% 标红）→ 标好色的成品表	不用学条件格式公式，说人话就出成品
预算分解与执行监控	年度目标 + 历史数据 → 各部门预算分解 + 超支预警与原因归因	预算从拍脑袋分配变数据支撑，超支自动报警
经营情景模拟	一个外部变化假设（如原料涨价）→ 对利润、现金流的影响测算	如果.....会怎样几秒出测算，不用手搭模型
现金流预测与资金预警	银行流水 + 收支数据 → 未来 7 / 30 / 60 天现金流预测 + 资金预警	资金紧张提前看到，不等付不出钱才发现
税务申报与合规辅助	业务数据 + 税务问题 → 申报辅助 + 合规口径建议	把业务的灵活处理拉回合规轨道
应收账款监控 + 自动催收	应收明细 → 超期分级标记（30 / 60 / 90 天）+ 到期自动提醒负责人	回款滞后自动揪出并催，不靠人脑记账期
应付款拆付款单	一张应付款清单 → 按供应商拆成 N 张格式统一的付款通知单 PDF	省逐个手填，字段齐全可直接走签批
多份合同条款	几份合同 PDF → 关键条款横向对比表 + 风险	自动揪违约责任 20 倍不

对比	逐条分级	对等这类坑，不漏
制度问答	财务/报销制度文档 + 一句话问题 → 引用原文条款的精确答复	制度散在多份文档，秒级定位第十八条并解释
财务数据脱敏处理	含敏感信息的财务数据 → 用编号、区间替代后的脱敏版本	用上 AI 又不泄露真实金额、客户和账户

二、行政 / 法务 / 合同

没有专职法务的公司，这一栏尤其值钱，相当于给行政配了个初审律师。

场景	喂什么 → 出什么	替你省了什么
合同审阅 + 改写	对方发来的合同 + 我方背景 → 风险清单 + 改成保护己方的版本 + 逐条建议	没法务也能按己方立场把合同审到能签
PDF 合同转 Word	扫描/PDF 版合同 → 可编辑 Word（条款、表格、签字栏还原）	PDF 转可编辑 + 结构化提取一步到位
文档精准修改	把这份文档某处改成 X → 就地精准改到段落/元素级	不用打开手动翻找，说一句就改好
散材料整合进档案	客户/各方发来的零散 Word 材料 → 提取整合进统一项目档案	散落材料统一归档，格式一致可检索
材料接受众重写	一份正确但没用的材料 + 目标读者（如不懂行的高管） → 看得懂的版本	把专业材料翻译成对方能接住的话
多份方案对比	几份方案/产品 PDF → 各自定位、能力 + 差异分析	快速吃透多份文档并找出协同与差异点

三、市场 / 品牌 / 内容

内容岗是 AI 提效倍数最高的地方，但也最容易出 AI 味，关键是把真实素材和人的判断喂进去。

企业 AI 转型白皮书

场景	喂什么 → 出什么	替你省了什么
分人群批量写长文	目标人群 + 标题 + 素材库 → 一批成稿（钩子 + 论点 + 收藏点）	一次产出整批，风格、字数、命名全自动统一
写稿前素材调研	人群定义 + 资料源 → 调研资料 + 写作框架两份文件	把找素材、定角度标准化，框架做一次反复用
选题与标题打磨	一个主题/观点 → 十几个多角度标题 + 传播力打分	不用憋标题，一次产出角度齐全的选题库
长文转平台风格	一篇原文 + 目标平台 → 平台化改写稿（换标题、加钩子、拆短段）	技术稿一键变小红书/公众号风，省人工重写
文章转图文卡片	Markdown 文章 → 一批 3:4 图文卡片 PNG（封面 + 页码 + 尾页）	文章一键变可直发的图，几十篇一次跑完
爆款公式拆解	别人的爆款文 → 为什么火分析 + 可复用框架	把零散爆款经验压成能复用的写作公式
竞品博主逆向拆解	竞品博主的一批文章 → 写作套路 + 选题清单	摸清对手打法，产出自己的可写选题
内容批量把关	一批待发文章 → 是否合规/是否编造的校验 + 改写	批量内容质检，避免对外发布失真
文章抓取 + 总结	一个或多个文章链接 → 正文 + 核心观点结构化总结	几千字长文几秒吃透，顺手判断对业务有没有用
某账号全量归档	一个公众号/博主 → 全部历史文章 CSV + 归档文档	一个号的内容全搬下来做素材底料
海量文章打标分类	上百篇文章 + 分类标准 → 每篇打好标签的结构化数据	人读不过来的批量分类，几分钟一致口径搞定
长报告速读	几万字行业报告 → 结构化摘要 + 一句核心判断	几万字压成几屏，快速判断要不要深读
视频拆解成复	一条想模仿的视频 → 逐镜头分镜 + 可喂给生	想做个类似的直接落成可

刻指令	成模型的指令	执行指令
-----	--------	------

四、销售 / 商务

这一栏的核心不是 AI 替你卖，是 AI 帮你把每次沟通的信息榨干、把每个客户摸透。

场景	喂什么 → 出什么	替你省了什么
客户录音转纪要	一段沟通录音 → 逐字稿 + 结构化纪要（痛点、预算、风险、待办）	几小时录音几分钟出要点，关键信号被显性化
需求翻译成方案	客户模糊的口头需求 → 功能分层清单 + 架构 + 落地场景	模糊需求变可确认方案，让客户先认方案再谈价
多档报价设计	已拆解的功能模块 + 同行情报 → 基础/进阶/全功能三档 + 话术	报价有理有据，三档让客户自选
见客户前尽调	公司名 / 创始人 → 核心事实 + 画像 + 几个懂行切入点	会前用证据武装到牙齿，进门能对话
行业痛点判断	目标行业 → 痛点 Top10 + 每条 AI 能不能解的判断	想清楚该推哪些场景、避开哪些，不聊偏
同行案例弹药库	对标品牌 + 场景 → 带名带数的案例卡片	给客户有威慑力的同行论据，靠案例建信任
沟通前提问清单	一个新客户的背景 → 该问对方哪些信息的结构化清单	把销售沟通标准化，不漏问关键信息
拜访后复盘校准	拜访后的语音反馈 → 复盘 + 纠正之前判断 + 下次三件事	每轮见客户形成反馈→校准→下次闭环
历史案例精准映射	一个新需求 → 历史项目里最匹配的案例 + 可讲金句	一句话打掉你做过类似的吗这种质疑
内部案例脱敏成材料	内部真实项目文档 → 脱敏后可发给潜客的对外案例	不能公开的项目脱敏后变可传播，不丢说服力

五、客服 / 客户成功

客服和客户成功适合先做副驾驶：AI 帮人分类、检索、起草、总结，人负责判断和兜底。

场景	喂什么 → 出什么	替你省了什么
常见问题自动回复	产品手册 + 售后政策 + 用户问题 → 标准答复草稿	高频问题先由 AI 接住，减少重复解释
客服对话副驾驶	当前聊天记录 + 客户画像 → 下一句回复建议 + 风险提示	新客服也能按成熟话术处理问题
客诉自动分类	一批客诉文本 → 按原因、严重程度、责任方打标	不靠人工翻聊天记录，快速看清问题分布
工单摘要与流转	工单详情 + 历史处理记录 → 摘要 + 建议转派部门	工单从谁看谁理解变成标准化流转
售后政策问答	售后制度 + 用户诉求 → 可引用条款的处理建议	避免客服凭感觉承诺、口径前后不一
退款 / 赔付初审	订单、聊天记录、规则 → 初步判断 + 需要人工确认的点	把规则清楚的初筛交给 AI，人处理灰区
客户情绪预警	对话文本 + 历史互动 → 情绪等级 + 升级提醒	投诉升级前提前发现，减少被动救火
客服质检与话术改进	一批客服录音 / 文本 → 质检评分 + 可改进话术	抽检变批量质检，培训材料从真实问题里长出来

六、风控 / 审核 / 合规运营

风控、审核、合规类场景适合先做初筛 + 解释 + 预警，高风险动作继续由人兜底。

场景	喂什么 → 出什么	替你省了什么
风险案例初筛	业务单据 / 交易 / 申请材料 → 风险等级 + 命中原因	人先看高风险，低风险自动归档

审核材料一致性校验	多份申请材料 → 信息不一致清单	证件、合同、表单之间的矛盾自动揪出来
异常交易排查	交易明细 + 历史规则 → 异常线索 + 排查路径	从海量流水里先筛出值得看的点
策略命中解释	风控规则 + 命中记录 → 为什么拦截 / 放行的解释	黑盒规则变成业务能看懂的话
客诉风控复盘	客诉文本 + 处理结果 → 风险归因 + 规则改进建议	把投诉倒逼成规则优化，而不是只处理个案
审核话术生成	风险原因 + 客户背景 → 合规、克制的沟通话术	减少一线随口解释带来的二次风险
Badcase 归因	一批误判 / 漏判案例 → 共性原因 + 改进清单	风控不只看结果，还能追到规则和流程问题
风险日报自动生成	当日风险数据 + 重点案例 → 日报 + 预警事项	管理层不用等人工汇总，天天看到风险变化

七、战略 / 调研 / 情报

老板和战略岗最常问的行业到底怎么样、对手在干嘛、这事能不能做，AI 都能给到带来源的答案。

场景	喂什么 → 出什么	替你省了什么
工具/方案选型调研	一个方向（如某类业务系统）→ 分类候选清单 + 选型建议	省逐个试错，一次拿到可决策的清单
行业/公司深度调研	一个主题 / 公司名 → 结构化画像 + 带来源的数据	见客户前、做决策前快速建认知，数据可溯源
动态核实 + 竞品对比	一条传言 / 一张截图 → 核实后的判断 + 竞品对比	把朋友圈级传言核实成有据可查的判断
政策匹配度分	一份政府扶持政策原文 → 符不符合 + 能拿多	一大份政策秒变我能拿多

析	少 + 先申哪个	少钱的行动清单
网站取数成报告	政府/资讯网站 + 主题 → 结构化调研报告	逛官网找信息升级成可存档的调研报告
大页面全量抓取	反爬/被截断的大网页 → 完整内容入库	把抓不全的大页面完整落地成可用素材
付费/登录内容提取	已订阅的付费专栏/会员内容 → 完整正文归档	付费、需登录的内容也能完整搬运学习
社媒声音抓取	几篇相关笔记 → 正文 + 全部评论 + 场景分类清单	把用户真实吐槽批量结构化，连折叠评论不漏
业务流程答疑	一个不熟的业务流程 → 每环节核心问题 + 常见坑	快速建立流程认知，为对接/沟通打底
数据找规律	一批内容/经营数据 → 高互动规律 + 可复用框架	把什么会火/什么有效从拍脑袋变成数据支撑
工具生态全景	一个工具品类 → 全景地图 + 官方 vs 社区区分	一次摸清某类工具生态，选型不走弯路

八、项目 / 交付（适用于做 IT、做方案、做乙方的团队）

场景	喂什么 → 出什么	替你省了什么
AI 落地方案设计	业务场景 → 可评审的技术方案（输入、输出、运行逻辑、边界）	做个智能体具体到可评审、可写进任务书
可复用底座设计	架构想法 → 分层架构 + 技术选型 + 运行链路	确立可复制交付模式，降低每次定制成本
竞品方案拆解	竞品产品手册/方案 → 功能本质 + 客户真实意图 + 报价锚点	客户发个手册问能不能做翻译成商业判断
需求可行性评估	客户需求说明书 → 能不能做 + 风险点 + 报价区间	几分钟变成接单决策依据，提前点破预期坑

流程图/泳道图 绘制	系统逻辑 / 项目排期 → Mermaid 流程图、泳道图、甘特图	方案可视化，直接贴进文档或 PPT
架构图绘制	一段架构描述 → 架构图 + 嵌入交付文档	纯文字方案瞬间补上可视化，专业度上一档
多轮沟通整理成文档	多份沟通记录 → 一份项目总览文档（规划、分工、报价、边界）	散在多处的项目信息汇成一份可对外的总览
按部门拆落地场景	客户的业务流程 → 按部门拆的 AI 场景清单 + 各自 ROI	抽象的 AI 转型变成客户看得懂的具体清单

九、人力 / 培训

无论你是做内训还是对外讲课，这一栏把备课、讲、复盘整条链路都覆盖了。

场景	喂什么 → 出什么	替你省了什么
课程该不该做 / 分层	定位 + 客户诉求 → 要不要做的判断 + 课程分档逻辑	把模糊的业务决策拆成可执行的产品结构
大纲设计 + 迭代	客户背景 + 对接录音 → 模块化大纲，随反馈快速调整	一次性方案变可复用模块库，跨客户改改就用
课件 PPT 批量制作	课程框架 → 商务风 PPT（逐页 + 配图占位）	把内容直接落成可用 PPT，省逐页手搓
演示 Demo 评估	现有 Demo 列表 + 受众 → 保留/替换/新增建议	备课阶段就把超时、敏感、翻车的 Demo 筛掉
逐字稿 + 时间评估	PPT 全文 → 每页口播稿 + 每模块时长	讲师不用临场组织语言，提前发现时间不够
考核题/互动设计	课程范围 → 选择判断题 + 可现场用的互动问答	测验题自动生成不超纲，解决现场冷场
现场录音复盘	培训现场的几小时录音 → 带原话佐证的复盘报告	把一场培训拆成可量化、可执行的改进项

交付物批量生成	大纲 + 课件 → 交付清单 + 按内容生成的成品文件	从想清单到出成品一步到位
培训素材池建设	知识库目录 → 案例池 + 方法论池 + 工具对比池	散在各处的素材系统化成可直接进课的弹药

十、知识管理 / 资料沉淀

这一栏是承接能力里上下文整理能力的直接落地，把散在各处、锁在人脑里的知识，变成组织能调用的资产。

场景	喂什么 → 出什么	替你省了什么
整库灌入做预热	整套公司资料 → AI 加载后秒懂全部背景	省去每次沟通重复交代背景
定向检索 + 聚合	一个目录 + 一批关键词 → 按主题聚合的报告	散在十几个文件里的信息一次性聚合
大型素材库摸家底	上千上万篇素材 → 内容地图（主题分布、质量、趋势）	上万篇几分钟变一张可读地图，看清家底
在线知识库迁本地	在线知识库链接 + 筛选规则 → 按层级镜像到本地	成百上千篇结构化搬运，省逐篇复制
素材归类入库	一批文章 + 主题划分 → 建好分类子页 + 填好表	散文章自动归类、提摘要、批量入库
文档双向搬运	文档链接 → 归档到本地 / 写回在线 / 修好乱表格	在线与本地知识库双向搬运，格式自适应
素材资产化	若干文件夹 → 哪些值得变成案例/方法论的评估清单	把散落记录系统性体检，隐性经验显性化
知识库体量盘点	一个知识库目录 → 体量统计 + 精简/重构建议	摸清家底，维护臃肿知识库时减负
资料秒级定位	一个模糊关键词 → 精确文件位置 + 它是干嘛	在庞大资料库里秒级找到

	的	那个东西在哪
--	---	--------

十一、通用办公自动化（全员适用）

这一栏不挑岗位，是所有人每天都在重复、最该交给 AI 的活。

场景	喂什么 → 出什么	替你省了什么
录音转文字 + 总结	一个会议/沟通录音 → 逐字稿 + 结构化要点	几小时录音几分钟出全文加要点
会议视频 + 转录下载	一个会议/培训视频链接 → 视频 + 完整逐字稿	没下载权限也能完整拿到视频和文字稿
长视频/播客沉淀	一个音视频链接 → 音频 + 逐字稿 + 可检索要点	一小时内容一键变本地存档加要点
每日工作晚报	当天的工作痕迹 → 自动生成的结构化复盘晚报	全程零人工，天天跑、可存档
邮件汇总推送	近 7 天邮件 → 过滤营销后的汇总 + 重点提示	每天自动看有没有要紧事，不用翻邮箱
定时任务搭建	一个每天自动做某事的需求 → 到点自动跑	把重复的体力活一次性配掉
会话/资料归档	对话/资料 → 按项目归档 + 全量备份	不丢历史，换电脑也能恢复
收藏素材批量下载	一批收藏/链接 → 批量下到本地素材库	逛收藏一键沉淀成可用素材库
聊天截图解读	群里/私聊的截图 → 对方意图 + 下一步建议	碎片化聊天信息变成清晰局势判断
长文转流程图	一篇讲流程的文章 → 可视化流程图	文字描述一键变看得清的图

最后说一句

这张表里一百多个场景，你不会全用上，也不该全用上。它的真正的价值，是让你在翻完这本书、合上它之前，心里有一张我们公司的 AI 地图，知道哪个岗位、哪个环节，是你下一步最该动手的地方。

地图我给你画到这儿了。往哪儿走第一步，这个判断，得你自己来。

参考来源

来源	正文使用位置	正文口径	状态
Stanford Digital Economy Lab, 《企业 AI 落地手册：51 个成功案例》 (2026 年 4 月)	前言、第 1、2、5、7、8、11、13、16、18、20、22、25、结语	51 个成功企业 AI 项目；技术可跑，难点在技术之外；数据、流程、组织、变革管理是主要挑战	已核原文，正文可用
McKinsey, The state of AI in 2025	第 1、5、13、17、19、结语	约 88% 组织已常态化使用 AI；约三分之一开始规模化；39% 报告企业层面 EBIT 影响；高绩效组织更重视工作流重构	已核官方页，正文可用
BCG, Artificial Intelligence capability page ; BCG, The Leader's Guide to Transforming with AI	第 1、6、11、结语	10-20-70 原则：算法 10%、技术和数据 20%、人和流程 70%	已核官方页，正文可用
RAND, The Root Causes of Failure for Artificial Intelligence Projects	第 1、第 10	AI 项目失败往往源于业务目标、领导认知、数据、基础设施、项目管理等问题，不宜写成“纯技术失败”	已核官方页，正文可用
Cloudera × Harvard Business Review Analytic Services	第 2、第 3	仅 7% 企业认为数据完全 AI 就绪；数据准备和治理是企业 AI 卡点	已核官方页，正文可用
Docebo, AI training relevance article	第 2、第 24	85% 学员认为 AI 培训没有充分帮助他们在岗位里理解或使用 AI	已核官方页，正文可用

企业 AI 转型白皮书

NBER, Generative AI at Work	第 15、第 25	生成式 AI 用于客服后，平均生产率提升，低经验员工提升更明显	已核论文页，正文可用
Klarna 官方新闻稿；CX Dive, 2025-05-09	第 10、第 16、第 22	官方新闻稿支撑 2024 年处理量、人工等效、重复咨询和响应时间；CX Dive 根据 Klarna 发言说明 2025 年重新招募人工客服，并保留人工服务入口	已核官方新闻页，正文可用
Gartner, Agentic AI projects cancellation forecast	第 19	超过 40% Agentic AI 项目可能在 2027 年底前取消，原因包括成本、业务价值不清、风险控制不足	已核官方新闻页，正文可用
Stanford HAI, 2025 AI Index Report	第 1	AI 商业使用率提升；生成式 AI 相关投资和采用加速	已核官方页，正文可用
The GenAI Divide: State of AI in Business 2025 公开预印本；MIT Media Lab NANDA 主页	第 1、第 10、第 17	公开预印本提出约 95% 的企业生成式 AI 试点没有产生可衡量的损益回报；核心讨论是 workflow 集成、学习和组织落地	已核官方页，正文可用
Microsoft / IDC, The Business Opportunity of AI	第 17	企业生成式 AI 投入的平均回报，以及头部企业更高的 ROI 水位	已核公开 PDF，正文可用