

025 / CAPABILITY LEDGER

CAPABILITY PORTFOLIO / TASK ·
EVIDENCE · AUTHORITY

AI营销人才 能力模型

把“会用工具”升级为可观察、可评测、可迁移并能对经营结果负责的人才能力操作系统。



能力由任务与证据定义

拆清输入、判断、输出与异常

用五级阶梯决定授权

稳定交付、异常处理与系统设计

十岗拥有不同专业纵深

共同底座不替代岗位判断

14/30/90 天落地

任务盘点、实战认证与团队重组

BOARD BRIEF 01

执行摘要

竹势 AI 营销智库出品

BOARD TALENT RULE / 董事会人才判断

能力不是工具使用记录，而是把经营任务转成可验证工作、处理异常并承担判断责任的稳定证据。

任务先于工具

证据先于等级

授权匹配风险

系统条件单独分账

企业正在快速增加人工智能工具、培训课程和使用要求，但多数组织尚未建立一套能够回答以下问题的能力体系：

哪些员工只是知道工具功能，哪些人能够稳定完成真实经营任务？

哪些岗位可以扩大人工智能授权，哪些岗位仍需严格限制自动执行？

一名内容人员完成大量初稿，是否代表其具备更高的营销能力？

一名投放人员节省了报表时间，是否意味着预算配置和实验判断得到改善？

当结果不佳时，问题来自个人能力、团队组合，还是数据、流程、权限与系统条件？

当技术快速变化时，企业应该培养什么稳定能力，避免培训投资随产品版本一起折旧？

这些问题不能靠“使用过多少工具”“参加过多少课程”“获得过多少证书”回答。工具使用记录能够证明接触，不等于证明能力；内容数量能够证明产出活动，不等于证明顾客反应、品牌质量或商业结果；节省工时能够证明局部效率变化，也不能直接证明收入、利润、客户价值或风险结果得到改善。

本报告提出一个基本判断：

AI营销人才的稀缺性，不在于熟练操作某个产品，而在于能否把模糊经营目标转化为可评测任务，调度人、模型、知识、数据与工作流，并对事实、品牌、客户、预算和结果承担判断责任。

这意味着企业需要从“工具中心的人才观”转向“任务—证据—责任中心的人才观”。

在工具中心的人才观中，企业通常先列出热门产品，再要求员工学习，最后用使用率、结业率、生成数量和工时变化证明项目成功。这种做法容易奖励三类低风险表演：第一，选择容易展示但经营影响有限的任务；第二，批量增加产出，却不承担事实、品牌和效果责任；第三，把复杂问题包装成个人操作技巧，回避流程、数据、权限和组织协同的真实障碍。

在任务—证据—责任中心的人才观中，企业先定义必须改善的经营任务，再拆解输入、动作、判断、输出、异常与责任，随后确定哪些能力属于个人，哪些必须由团队、流程、数据和系统共同提供。人才评测不再以“本人认为自己会”为依据，而是以真实任务、工作样本、过程记录、异常处置、复盘答辩和持续业务表现为证据。

本报告形成五套相互连接的原创管理框架。

第一套是 MARKET-8 能力组合：

任务建模（Model the task）：把模糊目标转化为边界清晰、输入明确、结果可评测的任务。

洞察与问题定义（Ask）：识别真正需要解决的顾客、品牌、增长或经营问题。

研究与证据（Research）：寻找、比较、验证并记录支持判断的证据。

知识与数据（Knowledge）：识别所需知识、数据、上下文、权限和质量条件。

系统编排（Engineer）：组合人员、模型、工具、规则、工作流和人工审核节点。

判断与质量（Test）：评测事实、逻辑、品牌、创意、客户体验、效果和风险。

执行与协同（Execute）：推动任务跨角色、跨系统、跨周期稳定完成。

责任与治理（Responsible control）：明确授权、审批、留痕、升级、纠错和最终责任。

第二套是 能力证据五级阶梯：

一级：知道；

二级：可在指导下完成；

三级：可独立稳定交付；

四级：可处理异常并持续优化；

五级：可设计系统并带动他人。

等级不是按工作年限、职位名称或考试成绩划分，而是按独立性、任务复杂度、异常处置能力、稳定性和影响范围划分。每一级必须对应可观察的工作样本。

第三套是 个人—团队—系统三层能力账本。该账本用于防止企业把数据缺失、流程断裂、权限不足、系统不稳定和决策迟缓错误归因于员工。个人能力可以被要求和评测，但个人不应为组织没有提供的数据、工具、授权和协同机制承担虚假责任。

第四套是 十岗能力投资组合。品牌、内容、创意、媒介、投放、社媒、CRM、研究、营销运营和管理岗位共享任务建模、证据、数据、协同与治理底座，但专业纵深、人工智能新增能力、不可下放判断和证据任务必须不同。企业不能用同一张“人工智能能力表”覆盖所有岗位。

第五套是 30天实战认证包，由任务说明、允许资源、过程记录、产出样本、异常注入、评审量规、复盘答辩和再认证周期构成。其目标不是制造另一种证书，而是用受控、可比、接近真实工作的任务，判断员工能否稳定交付并承担责任。

整套能力模型最终必须进入六个经营机制：岗位卡、项目授权、绩效评审、晋升评议、薪酬与人才盘点。若能力模型只停留在培训目录、能力词典或人力资源系统页面，它不会改变工作，也不会改变组织结果。

BOARD BRIEF 02

十个核心结论

竹势 AI 营销智库出品

核心结论一：人工智能素养是进入条件，不是经营能力

了解人工智能的基本原理、局限、数据风险和适用边界，是多数营销岗位未来的共同基线。但“理解人工智能”与“能够负责营销任务”属于不同层级。

人工智能素养回答的是：员工是否知道模型可能出错、知道敏感数据不能随意输入、知道输出需要验证、知道不同工具具有不同边界。

任务能力回答的是：员工能否把经营目标转成任务，组织所需证据和数据，选择合适的人机分工，评测结果并处理异常。

专业判断回答的是：员工能否识别品牌战略是否偏移、顾客洞察是否成立、创意是否有区分度、媒介组合是否合理、实验是否能够支持因果判断、客户触达是否损害长期关系。

三者不能互相替代。企业可以把人工智能素养作为全员基线，但不能把基线培训结业视为岗位胜任。

核心结论二：工具熟练度会快速折旧，任务建模与专业判断更稳定

具体产品的界面、功能、价格、模型能力和自动化方式会持续变化。企业若围绕产品名称建立能力词典，就会不断追赶版本，且难以判断技能是否能够迁移。

更稳定的能力包括：

从经营目标识别真正问题；

把任务拆解为输入、动作、判断、输出与异常；

确定证据标准和质量门槛；

识别数据缺口与权限边界；

比较人、模型、规则和自动化各自适合承担的工作；

发现事实错误、逻辑跳跃、品牌偏差与客户风险；

根据结果进行复盘和系统改进。

工具熟练度仍然重要，但应被视为实现任务的局部条件，而不是人才价值的核心定义。

核心结论三：会生成不等于会判断，生成能力越强，审阅门槛越高

当系统能够快速产生更多文本、图像、视频、研究摘要、受众分群和优化建议时，组织面临的主要瓶颈会从“能否产出”转向“能否筛选、验证、组合、授权和负责”。

低成本生成会同时放大两类结果：高质量方案的探索空间，以及低质量内容、错误事实、同质化创意、错误归因和品牌风险。

因此，判断与质量不应作为任务末端的审批动作，而应成为每个岗位的核心能力。员工必须能够解释：

为什么这个结论可信；

使用了哪些证据；

哪些事实尚未确认；

哪些选项被排除以及原因；

哪个指标能够验证效果；

哪些情况需要停止自动执行；

失败后如何恢复和修正。

核心结论四：企业必须区分个人能力、团队能力和系统条件

一名员工无法独立弥补以下问题：

客户数据长期缺失；

品牌规范没有形成；

市场与销售目标冲突；

审批人长期不响应；

系统之间不能交换数据；

员工没有必要权限；

关键指标没有定义；

同一任务由多个部门重复拥有；

供应商交付缺少可验证标准。

如果这些组织问题被记入个人能力评价，企业会得到错误的人才结论，并诱发防御行为。员工会更倾向于选择容易完成、风险较低、证据要求较弱的任务，而不是解决真正困难的经营问题。

核心结论五：能力等级必须由工作证据决定，而不是由年限、职级或自评决定

“工作五年”“参加高级课程”“自评熟练”均不能稳定证明员工处于较高能力等级。

五级能力阶梯必须通过工作样本识别：

一级能解释概念和风险；

二级能按既定方法完成受控任务；

三级能独立、连续、稳定地交付；

四级能应对数据缺失、事实冲突、系统故障、目标变化和品牌风险；

五级能设计任务系统、建立质量机制、培养他人并提高团队总体表现。

晋升评审应重点审查员工处理过的复杂任务、异常案例、判断依据、系统改进和对他人的带动，而不是只审查常规产出数量。

核心结论六：真实工作样本比泛化证书更接近岗位表现

证书可以证明学习经历，但无法充分证明员工能够在真实约束下工作。更可靠的评测应包含：

不完整或相互冲突的信息；

明确的时间和预算限制；

真实品牌与渠道规则；

需要说明取舍的多个方案；

事实、合规或客户风险；

中途发生的异常；

过程记录与复盘答辩。

工作样本也并非天然公平。企业必须控制任务泄露、资源差异、评委偏差、残障可及性和不同岗位的机会差异，并允许候选人解释其选择和条件限制。

核心结论七：十类岗位需要共同底座，但不能共享同一套专业能力

品牌岗位的核心是品牌意义、长期一致性与声誉责任；内容岗位的核心是叙事结构、事实表达与渠道适配；创意岗位的核心是概念区分度、表现系统与制作可行性；投放岗位的核心是实验、预算、边际回报和风险止损；CRM岗位的核心是客户生命周期、授权同意、分群逻辑和触达节奏。

若企业用同一组“会写、会画、会分析、会自动化”描述十个岗位，就无法支持招聘、授权与晋升。共同底座应保持有限，专业纵深必须依据岗位真实任务定义。

核心结论八：能力投资要形成组合，而不是要求每个人成为全能者

企业有五种能力供给方式：

在岗培养；

跨岗迁移；

外部招聘；

共享专家或外部服务；

人工智能员工与自动化系统。

最优组合取决于能力邻接、训练距离、任务风险、需求频率、业务窗口和知识保密性。低频、高专业、短窗口任务可能更适合共享专家；高频、规则清晰、可验证任务可能适合系统化；高风险、强品牌责任任务必须保留明确的人类判断所有者。

核心结论九：能力模型必须直接控制项目授权

能力评测的经营价值，不是把员工分成若干等级，而是决定谁可以做什么、在什么条件下做、需要谁复核。

例如：

二级员工可以在模板和指导下完成低风险内容初稿；

三级员工可以独立完成既定品牌范围内的常规任务；

四级员工可以在异常发生时调整方案，并承担跨系统任务；

五级员工可以设计 workflow、设定质量门和审批机制。

授权应依据任务风险和能力证据动态调整，而不是只依据职位名称。

核心结论十：能力治理必须设置透明、最小必要和可申诉边界

当能力数据进入绩效、晋升、薪酬和人才盘点后，它会直接影响员工利益，因此必须设置明确边界：

- 仅收集完成明确人才目的所必需的数据；
- 不把个人私密对话、非工作内容和无关行为纳入评价；
- 不以自动化评分直接作出重大人事决定；
- 员工应知道评测目的、数据来源、保存周期和使用范围；
- 对评分、证据归属和异常情况提供人工复核与申诉；
- 对组织条件不足导致的失败进行分账；
- 认证结果设置有效期和复证机制；
- 禁止以隐蔽监控制造“持续表演”。

第一章 从工具清单到经营能力：AI营销人才模型的重新定义

竹势 AI 营销智库出品

AI 素养

共同基线

工具熟练

局部条件

任务能力

稳定交付

专业判断

不可下放

团队能力

互补复核

系统条件

数据流程权限

1.1 董事会真正需要解决的不是培训问题，而是人才配置问题

许多企业把人工智能人才建设交给学习发展部门，形成“课程采购—员工报名—结业统计—优秀案例展示”的闭环。这个闭环可以组织学习活动，却不足以回答经营层最关心的问题：

下一场品牌战役由谁负责；

哪些岗位可以减少手工执行；

哪些任务可以交给人工智能员工；

哪些判断必须由专业人员保留；

谁具备跨系统编排能力；

哪些团队存在关键能力单点；

招聘预算应投向专业纵深还是系统能力；

培训后是否能够迁移到真实工作；

员工表现不佳究竟是能力不足还是组织条件不足。

这类问题本质上属于人才配置、项目授权和经营治理，而不是课程管理。

AI营销能力模型必须成为经营系统的一部分。它的最小使用闭环应为：

经营任务定义 → 能力要求 → 证据任务 → 人员与系统配置 → 项目授权 → 过程观察 → 结果评审 → 能力更新。

若缺少其中任何一环，能力模型都可能退化为静态文件。

1.2 六个必须区分的概念

表1：AI营销人才能力的六层边界

表1：AI营销人才能力的六层边界			
层次	定义	可以证明什么	不能证明什么
AI素养	理解基本原理、适用范围、风险、数据与责任边界	知道如何安全、审慎地使用人工智能	不能证明能完成具体营销任务
工具熟练度	能够操作特定产品、功能或界面	能较顺利地使用某项工具	不能证明能力可迁移，也不能证明经营结果
任务能力	能把输入转化为符合标准的工作输出	能在特定任务中交付	不一定包含深厚专业判断或复杂异常处置
专业判断	能依据顾客、品牌、渠道、数据和商业规律作出取舍	能识别结果是否正确、合适和有价值	不代表能够独立完成跨系统执行
团队能力	多角色形成互补、交接、复核和备份	能完成个人无法独立完成的复杂任务	不能归为某个成员的个人分数
组织系统条件	数据、流程、知识、权限、技术、治理和资源环境	决定能力能否被稳定调用和放大	不能被当成员工个人素质

这六层中，前四层主要描述个人，后两层描述个人之外的生产条件。但四类个人能力也并非彼此包含。例如，一名资深品牌负责人可能具有很强的品牌判断，但对新工具并不熟悉；一名年轻运营人员可能工具使用熟练，却缺少对品牌资产和顾客关系的判断。合理的人才配置应识别这种差异，而不是用单一“人工智能能力分”覆盖。

1.3 AI营销能力与传统营销能力的关系

人工智能并未取消传统营销能力，而是改变了它们被调用和评测的方式。

顾客理解、市场细分、品牌定位、产品价值、传播策略、创意判断、媒介效率、客户关系和商业衡量，仍然构成营销专业的核心。变化在于：

- 第一，更多信息可以被快速收集和初步整理，因此专业人员必须更擅长判断来源质量、偏差和适用边界。
- 第二，更多内容和创意选项可以被低成本产生，因此稀缺能力从“产生一个选项”转向“提出正确问题、形成有区分度的方向、筛选高潜方案并建立学习机制”。
- 第三，更多执行动作可以跨系统自动运行，因此员工需要理解权限、审批、日志、异常恢复与责任分配。

第四，工作反馈速度提高，营销人员必须更快地把结果转为下一轮假设，而不是把项目复盘停留在总结。

因此，AI营销能力不是在传统营销能力旁边增加一列工具名称，而是重新设计专业能力如何进入任务、证据与系统。

1.4 AI营销能力与数字化、数据能力的边界

数字化能力主要关注流程是否被结构化、数据化和系统化；数据能力关注数据定义、质量、分析、实验和解释；人工智能能力关注如何在不确定输出和人机协作环境下完成任务。

三者存在交集，但不能混同。

表2：四类能力的边界与交集

表2：四类能力的边界与交集			
能力领域	核心对象	典型问题	对营销人才的要求
传统营销能力	顾客、品牌、市场与增长	做什么、为谁做、为何有效	洞察、策略、创意、渠道与商业判断
数字化能力	流程、系统与连接	工作如何被标准化并跨系统流动	流程设计、系统意识、接口与权限理解
数据能力	数据、指标与因果	事实是什么、变化为何发生	数据质量、分析、实验、解释与归因
AI任务能力	人、模型、知识、工具与规则的协同	哪部分由谁完成，如何评测和控制	任务建模、编排、质量、异常与治理

成熟人才不必在四个领域都达到专家水平，但必须知道何时需要其他专业角色介入。识别自身能力边界，本身就是高等级能力的证据。

1.5 企业常见的五种错误定义

错误一：把使用率定义为能力

使用频次只能说明员工使用了某项产品。员工可能反复执行低难度任务，也可能因为高风险任务尚未获得授权而使用较少。使用率适合作为采用观察，不适宜作为能力结论。

错误二：把生成数量定义为能力

内容、图片、视频或报告数量增加，可能来自更低的生成成本，也可能带来更高的审阅负担、品牌稀释和渠道噪声。必须结合采纳率、返工率、事实错误、品牌一致性、客户反应和业务影响判断。

错误三：把节省工时定义为价值

节省工时首先是过程变化。被节省的时间是否转移到更高价值工作、是否降低总成本、是否提高响应速度、是否增加收入或降低风险，需要独立验证。

错误四：把结业证书定义为胜任

结业证明学习活动完成，不证明员工能在真实约束下稳定交付。培训后必须通过任务迁移和工作证据进行验证。

错误五：把所有失败归因于个人

若企业没有提供品牌资料、客户数据、任务标准、系统权限和审批响应，员工不可能仅凭个人努力稳定交付。能力评审必须先完成三层分账。

1.6 管理含义

董事会与经营层应作出三项制度性改变：

不再批准仅以课程、工具和采用率为主体的人才计划；

要求每项能力绑定真实经营任务、可观察证据和授权用途；

要求人力、营销、数据、技术、法务与安全共同维护能力模型。

BOARD BRIEF 04

第二章 MARKET-8：AI营销人才的八维能力组合

竹势 AI 营销智库出品

M
任务建模

A
洞察定义

R
研究证据

K
知识数据

E
系统编排

T
判断质量

E
执行协同

R
责任治理

2.1 框架总览

MARKET-8不是八项彼此独立的技能，而是一条从经营目标到责任闭环的能力链。

表3： MARKET-8 能力组合总表

表3： MARKET-8 能力组合总表			
维度	核心问题	关键行为	典型工作样本
M 任务建模	这项经营目标应被转成什么任务	定义目标、对象、输入、输出、约束、指标和异常	任务卡、验收标准、风险清单
A 洞察与问题定义	真正需要解决的问题是什么	区分表象与根因，识别顾客、品牌与商业矛盾	问题陈述、假设树、机会定义
R 研究与证据	哪些证据足以支持判断	设计来源计划，验证事实，区分证据与推断	证据表、来源记录、冲突解释
K 知识与数据	任务依赖哪些上下文与数据	识别知识缺口、数据质量、访问权限和适用边界	数据需求表、知识包、质量说明
E 系统编排	人、模型、工具和规则如何配合	设计步骤、角色、交接、审批、日志和恢复路径	工作流、责任矩阵、运行说明
T 判断与质量	如何知道结果足够好	建立事实、品牌、创意、效果和风险评测	评审量规、测试集、缺陷记录
E 执行与协同	如何把方案稳定变成结果	推进节奏、交付、沟通、变更和复盘	行动计划、交接记录、复盘
R 责任与治理	谁能决定、谁需批准、谁最终负责	控制权限、升级异常、留痕、纠错和申诉	授权表、审批记录、事故复盘

八维共同构成一项完整的AI营销任务能力。某些岗位可能在其中两到三项具有专业纵深，但所有承担真实任务的人员都需要达到最低基线。

2.2 M：任务建模

任务建模是整套模型的起点。许多失败并非执行不力，而是任务从一开始就不可评测。

“提高品牌影响力”“多做一些好内容”“用人工智能优化投放”都不是合格任务。合格任务需要回答：

服务哪个经营目标；

面向哪个顾客或业务对象；

当前问题和基线是什么；

可用输入有哪些；

输出必须采用什么形式；

哪些约束不能违反；

用什么标准验收；

哪些异常需要停止或升级；

谁拥有最终判断权。

表4：从模糊要求到可评测任务

表4：从模糊要求到可评测任务		
模糊要求	主要缺陷	可评测任务表达
提升品牌声量	无受众、周期、渠道和质量边界	在八周内围绕指定人群建立三个可验证传播命题，形成渠道试验并评估认知变化
多做短视频	把数量当目标	围绕两个顾客问题设计十二条内容实验，比较完播、有效互动、线索质量和品牌一致性
优化广告	无预算与决策边界	在既定预算和品牌限制内，提出三项可回滚调整，经审批后测试增量效果
做竞品研究	无决策用途	为下一季度产品定位决策识别五项竞争变化，区分事实、推断与待验证假设
激活私域客户	无同意、频率和客户价值边界	对具备有效授权的沉默客户分层，设计差异化触达并控制投诉、退订和客户价值风险

任务建模能力的高级表现，不是把任务拆得越细越好。过度拆分会破坏整体判断，使员工只对局部动作负责。高级人员必须在可执行性与整体性之间取得平衡。

2.3 A：洞察与问题定义

人工智能可以快速总结大量信息，却不能自动保证问题值得解决。问题定义要求营销人员区分：

顾客说出的需求与真实行为矛盾；

渠道指标变化与商业问题；

内容表现下降与品牌、产品、受众或分发原因；

销售抱怨线索质量与线索定义、跟进速度或价值主张问题；

投放成本上升与竞争、创意疲劳、落地页、产品价格或测量误差。

高水平问题定义应保留不确定性，而不是过早给出答案。员工需要提出多个竞争性解释，并说明下一步需要什么证据来区分。

表5：问题定义的可观察行为

表5：问题定义的可观察行为		
低水平行为	中等水平行为	高水平行为
接受需求方原始表述	追问目标、对象和背景	重构问题并识别系统性矛盾
快速产生解决方案	先形成假设	比较竞争性假设与决策价值
以单一指标解释结果	结合多个指标	区分相关、因果、机制与边界
把异常视为噪声	记录异常	用异常挑战原有假设

2.4 R：研究与证据

研究与证据能力要求员工能够区分：

原始来源与转述；

事实、估计、观点和推断；

样本结果与总体结论；

相关性与因果性；

供应商披露与独立证据；

最新信息与已经失效的信息；

可以公开使用的材料与受限制信息。

工作样本不应只看报告是否完整，还要检查证据链是否能够复核。高等级人员应主动说明数字的时间、对象、样本、口径、来源性质和限制。

表6：证据质量检查表

表6：证据质量检查表		
检查项	关键问题	失败信号
来源权威性	是否优先使用法规、标准、论文和官方披露	主要依赖无出处转述
时间适用性	信息是否仍适用于当前市场与产品	使用旧数据解释新变化
样本与对象	研究对象是否与当前任务相近	把局部样本外推到所有客户
口径一致性	不同数字能否比较	混用收入、效率、采用和满意度
利益相关性	来源是否存在销售或宣传动机	把供应商自报当成独立结论
可追溯性	其他评审者能否复核	只有结论，没有证据记录
不确定性	是否说明未知和限制	把推断写成确定事实

2.5 K：知识与数据

知识与数据能力不是“会查资料”或“会做表”，而是知道任务依赖什么信息、信息能否被合法和正确使用，以及缺失信息如何影响判断。

营销任务通常需要四类上下文：

品牌与产品知识；

顾客、渠道和竞争知识；

历史表现与业务数据；

合规、权限和使用规则。

员工应能够识别数据质量问题，包括缺失、重复、定义漂移、时间错位、归因偏差和权限不明。不能在数据不足时制造精确结论。

2.6 E：系统编排

系统编排是AI营销人才区别于单点工具使用者的关键能力。它要求员工决定：

哪一步由人完成；

哪一步由模型辅助；

哪一步适合规则自动化；

哪一步需要外部数据或工具；

哪一步必须人工审批；

结果写回哪里；

发生错误后如何停止、回滚和恢复；

谁对跨步骤结果负责。

编排不等于把任务做得复杂。小团队和低频任务不应为了展示技术而建立庞大工作流。编排质量取决于是否降低交接损耗、提高稳定性和控制风险。

2.7 T：判断与质量

质量评测至少需要覆盖六个维度：

事实质量；

逻辑质量；

品牌质量；

顾客与体验质量；

商业与效果质量；

合规与风险质量。

不同任务的权重不同。创意任务不能只追求一致性；研究任务不能只追求表达流畅；投放任务不能只看平台回报；CRM任务不能只看转化而忽视退订、投诉和长期客户价值。

2.8 E：执行与协同

执行能力包括节奏管理、跨角色交接、版本控制、决策记录、利益相关者沟通和复盘。员工需要把任务从“有一个好方案”推进到“结果被正确交付并进入下一轮学习”。

人工智能增加了执行速度，也可能增加版本混乱和责任模糊。因此，高水平执行者必须保留清晰的状态、责任和变更记录。

2.9 R：责任与治理

责任与治理决定组织是否能够放心扩大授权。员工不仅要知道哪些行为被允许，还要知道：

哪些任务可以自动运行；

哪些任务只能形成草稿；

哪些信息不得进入外部系统；

哪些客户触达需要有效授权；

哪些品牌、预算和公开表达必须由指定人员批准；

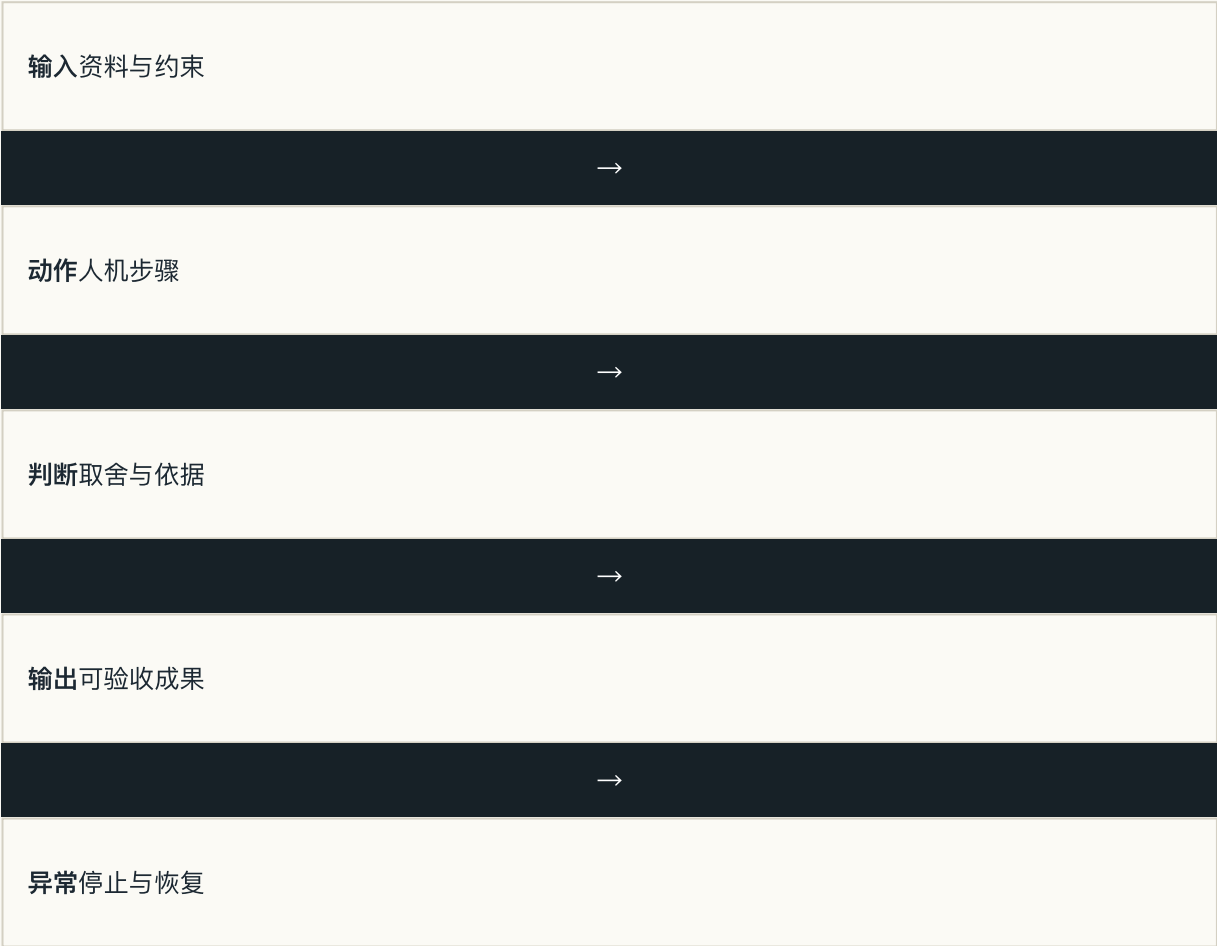
出现风险时向谁升级；

如何记录和修复错误。

责任不是对工具输出“兜底”的模糊要求，而是具体的判断权、审批权、执行权和停止权设计。

第三章 从真实经营任务反推能力项

竹势 AI 营销智库出品



3.1 为什么不能从热门工具清单出发

从工具清单出发会产生三个结构性问题。

第一，能力随产品变化而频繁重写。企业难以形成长期人才资产。

第二，不同岗位被迫学习大量与其核心责任无关的功能，造成培训投入分散。

第三，组织容易把“员工不会用工具”当作问题，而忽视任务没有定义、数据不可用、流程未调整和授权不清。

正确顺序应是：

经营结果 → 关键任务 → 任务证据 → 能力要求 → 人机分工 → 工具与系统。

工具位于链条末端，而不是起点。

3.2 五要素任务单元

每项能力都必须绑定一个可观察任务单元。建议使用“输入—动作—判断—输出—异常”五要素。

表7：五要素任务单元

表7：五要素任务单元

要素	定义	评测关注点
输入	完成任务所需资料、数据、规则和上下文	是否识别缺失、冲突和质量问题
动作	人员与系统执行的步骤	是否合理分工、留痕并控制成本
判断	需要专业取舍的关键节点	是否说明依据、替代方案和风险
输出	可验收的工作成果	是否满足事实、品牌、效果和格式要求
异常	偏离正常路径的情况	是否能够停止、升级、恢复和复盘

3.3 以经营结果为起点的任务反推

表8：经营问题到能力证据的反推示例

表8：经营问题到能力证据的反推示例

经营问题	关键任务	主要能力	工作证据
新产品认知不足	定义目标人群、核心价值和传播命题	洞察、品牌、证据、任务建模	定位论证、证据表、传播测试模
内容产出多但转化弱	识别内容与顾客旅程断点	数据、内容、渠道、实验	漏斗诊断、内容组合、实验复盘
投放成本持续上升	区分竞价、创意、受众、落地页和测量问题	分析、实验、预算判断	假设树、调整方案、止损条件
私域触达引发投诉	重构分群、频率、授权和人工接管机制	CRM、客户体验、治理	同意记录、分群规则、投诉复盘
品牌表达不一致	建立品牌知识、质量门和例外处理	品牌、知识、系统编排	品牌规则、评审量规、异常样本
市场团队返工严重	重构任务说明、交接与审批	任务建模、协同、运营	责任矩阵、任务卡、周期数据

3.4 能力项的合格标准

一项合格能力必须满足六个条件：

能够对应真实任务；

能够被行为观察；

能够通过工作样本评测；

能够区分不同熟练等级；

能够说明适用边界；

能够连接岗位、授权或人才决策。

“创新思维强”“善于使用人工智能”“数据感觉好”均不符合这些条件，因为它们无法直接观察和评测。

合格表达示例：

能将模糊营销目标转化为包含对象、输入、约束、指标和异常规则的任务说明；

能对关键结论建立可追溯证据链，并区分事实、推断和待验证假设；

能在品牌规范、客户授权和预算边界内设计人机协同工作流；

能识别输出中的事实错误、品牌偏差和效果测量缺陷；

能在数据中断、目标变化或系统失败时调整任务路径并完成复盘。

3.5 防止过度原子化

任务分析容易走向另一个极端：把工作拆成大量微小动作，并为每个动作建立评分。这会导致三类损失：

员工只优化局部动作，不对整体结果负责；

复杂专业判断被分解为机械检查；

评测成本大幅增加，模型难以维护。

企业应采用“最小完整责任单元”：任务足够具体，可以观察和评测；同时足够完整，保留结果责任与专业判断。

3.6 关键事件法的应用

企业可以从过去一年高成功、高失败和高风险项目中提取关键事件：

哪个行为显著提高了结果；

哪个错误造成返工、损失或风险；

哪个异常最能区分普通与高水平人员；

哪个决策依赖真正的专业判断；

哪项工作只能由团队共同完成；

哪个组织条件决定了能力能否发挥。

关键事件不应只收集“优秀故事”。失败、冲突、延误、错误归因和风险事件往往更能揭示真实能力。

表9：关键事件记录模板

表9：关键事件记录模板	
字段	记录要求
业务背景	目标、对象、时间、资源与限制
触发事件	发生了什么变化或异常
行为	当事人实际采取了什么动作
判断依据	使用了哪些信息和专业原则
结果	对质量、客户、预算、时间和风险的影响
替代行为	更优或更差的可能做法
能力映射	对应MARKET-8的哪些维度
系统条件	哪些流程、数据、权限影响了结果

第四章 通用底座：所有AI营销岗位必须具备的最小能力

竹势 AI 营销智库出品

问题定义

任务建模

证据审查

数据理解

质量判断

人机协同

责任治理

4.1 通用底座的设计原则

通用底座不是让所有岗位变得相同，而是建立最低协作语言。它应满足三个原则：

数量有限，避免挤压专业能力；

与风险和责任直接相关；

能通过真实任务评测。

建议企业把通用底座分为七类：问题定义、任务建模、证据、数据、质量、协同和治理。

4.2 七类共同能力

表10：AI营销岗位共同能力底座

表10：AI营销岗位共同能力底座

共同能力	最低要求	不合格表现	证据任务
问题定义	能区分需求表述与真实业务问题	直接接受模糊要求并开始生成	将业务要求重写为问题与假设
任务建模	能明确输入、输出、约束和验收标准	任务无法判断是否完成	编制任务卡和异常规则
证据审查	能核验事实并标记不确定性	把流畅表达视为正确	建立来源表与冲突说明
数据理解	能识别口径、质量和权限问题	对数字不问来源与定义	完成数据质量检查
质量判断	能按岗位标准评审输出	只看速度、数量或形式	使用量规评审多个样本
人机协同	能合理分配人、系统与模型任务	全部人工或全部自动	设计含审核与恢复的流程
责任治理	能遵守授权、审批、隐私和升级规则	越权执行或责任不清	完成风险情景与处置答辩

4.3 人工智能素养的最低基线

所有营销岗位至少需要理解：

输出可能存在事实错误、遗漏、偏见和不稳定；

同一任务不同运行结果可能不同；

敏感信息、个人信息和商业秘密具有使用限制；

外部内容可能存在版权、商标、肖像和授权问题；

自动化执行会放大速度，也会放大错误；

关键结果必须由明确责任人评审；

模型输出不是证据本身；

工具的可用性、价格和能力可能变化；

员工必须记录重要判断和人工修改；

高风险任务需要更严格的授权和复核。

理解这些原则，只代表达到通用基线，不能直接证明具备三级以上任务能力。

4.4 数据素养的岗位化要求

营销人员不必全部成为数据科学家，但必须能够回答：

指标如何定义；

数据来自哪里；

是否完整、及时和可比；

是否存在选择偏差；

结果是相关还是因果；

平台归因与企业实际收入是否一致；

平均值是否掩盖重要分群；

是否存在短期指标改善、长期价值受损；

数据是否具有合法使用基础。

品牌与内容岗位的数据要求更偏向受众、内容表现和品牌测量；投放岗位更强调实验、增量、边际成本和止损；CRM岗位更强调客户身份、生命周期、授权和长期价值；管理岗位更强调指标体系、资源配置与跨部门解释。

4.5 研究与证据的共同要求

所有岗位都应建立“证据不是装饰”的意识。证据必须影响决策。

例如，内容人员发现某个观点缺少可靠支持，应调整表达或删除；品牌人员发现热门趋势与目标顾客不符，应拒绝追随；投放人员发现平台归因口径变化，应暂停直接比较；CRM人员发现客户授权状态不明确，应停止自动触达；管理者发现项目结果主要受系统条件影响，不应把失败直接记入个人绩效。

4.6 人机协同的四种基本模式

表11： 人机协同模式与适用边界

表11： 人机协同模式与适用边界			
模式	描述	适用任务	主要风险
人主导、系统辅助	人定义问题和方案，系统处理局部工作	高风险、强专业判断任务	效率提升有限
系统产出、人审核	系统生成候选，人负责筛选和批准	内容初稿、研究整理、常规分析	审核疲劳、自动化偏见
系统执行、人监督	系统按规则持续执行，人监控异常	规则清晰、可回滚的运营任务	异常发现不及时
人机共同探索	人与系统反复提出假设、比较方案	策略、洞察、创意探索	责任边界模糊

企业不应预设“自动化程度越高越先进”。最佳模式取决于任务风险、可验证性、可回滚性和专业责任。

4.7 判断与质量的共同基线

员工至少应掌握“五问评审法”：

事实问： 哪些陈述可以核验，来源是什么？

逻辑问： 结论是否由证据支持， 是否存在跳跃？

品牌问： 是否符合品牌定位、语气、价值与长期资产？

客户问： 是否真正帮助客户， 是否造成误导、打扰或不公平？

经营问： 如何影响收入、成本、现金、客户价值、速度或风险？

不同岗位可在此基础上增加专业评审维度。

4.8 治理与责任的共同基线

员工必须知道四类边界：

数据边界： 哪些信息可以使用、共享和保存；

执行边界： 哪些动作可以自动完成， 哪些只能形成草稿；

判断边界：哪些决定必须由指定专业人员作出；

升级边界：哪些异常必须停止任务并上报。

表12：通用风险情景与合格行为

表12：通用风险情景与合格行为		
风险情景	不合格行为	合格行为
来源相互冲突	选择更符合预期的一项	记录冲突、比较权威性并降低结论强度
顾客数据缺失	用相似人群数据替代且不说明	标记缺口、说明影响并提出补充方案
系统建议大幅增加预算	直接执行	验证边界、模拟影响、设置审批与止损
自动生成内容涉及敏感声明	修改措辞后直接发布	进入专业审核，保留证据和批准记录
客户要求删除个人信息	继续保留以便分析	按制度停止使用并触发处理流程
品牌热点机会与价值观冲突	为流量快速跟进	依据品牌责任拒绝或调整参与方式

4.9 不同行业的基线不能完全相同

金融、医疗、汽车、教育、公共服务等领域涉及更高的信息、合规和客户影响风险，其共同能力门槛应更高。高风险岗位可能要求：

更严格的来源等级；

强制双人复核；

更小的自动执行范围；

更完整的过程日志；

更短的复证周期；

对法律、伦理与行业规则进行专项认证。

低风险任务也不能取消质量责任，只是可以采用更轻量的控制方式。

4.10 管理含义：共同底座只能解决“能否安全协作”，不能替代专业纵深

建立共同底座后，企业不应宣布所有营销人员已经具备“AI营销能力”。共同底座解决的是最低协作语言和安全使用问题。真正决定品牌、内容、创意、投放、客户运营和经营管理质量的，仍是岗位专业纵深。

下一部分将进一步展开专业判断、系统编排、能力证据五级阶梯、十岗差异化能力投资组合，以及完整企业与组织案例。

来源

第五章 专业判断与系统编排：从“能产出”到“能负责”

竹势 AI 营销智库出品

事实门

来源与不确定性

顾客门

真实需求与体验

品牌门

长期一致性

效果门

增量与边际价值

责任门

权限与可恢复

5.1 人工智能提高的是方案供给，专业判断决定的是方案价值

在传统营销组织中，人员和时间限制往往使团队只能形成少量研究结论、传播命题、创意方向或投放方案。人工智能显著降低了候选方案的生产成本，但没有自动降低判断成本。

当候选方案从三个增加到三十个，管理难题不再是“有没有想法”，而是：

哪些方案真正对应顾客问题；

哪些结论具有可靠证据；

- 哪些创意具有独特性，而非表面新奇；
 - 哪些优化建议只对平台指标有利，却损害利润或客户价值；
 - 哪些内容符合品牌短期传播需求，却侵蚀长期品牌资产；
 - 哪些自动化动作可以回滚，哪些错误一旦发生就不可逆；
 - 哪些决策可以下放给系统，哪些必须保留明确的人类责任人。
- 因此，专业判断不是人工智能应用之前的旧能力，而是人工智能扩大产出规模之后更加稀缺的控制能力。
- 专业判断至少包含五种取舍：
- 问题取舍：决定什么问题值得解决；
 - 证据取舍：决定哪些信息足以支持行动；
 - 方案取舍：决定哪种方案更符合顾客、品牌和经营目标；
 - 风险取舍：决定效率、创新与控制之间的边界；
 - 时间取舍：决定何时继续分析、何时试验、何时停止。
- 如果员工只负责调用系统、生成候选和整理输出，却不承担上述取舍，就不能被定义为高等级AI营销人才。

5.2 专业判断的五个质量门

表13：专业判断五道质量门

表13：专业判断五道质量门			
质量门	核心判断	典型问题	主要责任人
顾客真实性门	任务是否建立在真实顾客问题上	这是顾客真实障碍，还是内部想象	洞察、品牌、CRM、业务负责人
品牌一致性门	行动是否积累而非消耗品牌资产	是否与定位、承诺和长期记忆结构一致	品牌负责人、创意负责人
证据可信门	结论是否有足够证据支持	来源、口径、样本和边界是否可靠	研究、数据、专业评审者
经营可行门	方案是否能改善真实经营结果	是否考虑收入、成本、现金、能力和机会成本	管理者、投放、运营负责人
责任可控门	执行是否可授权、追踪和纠错	谁批准、谁执行、谁停止、谁负责	业务负责人、法务、安全、运营

五道门不应全部集中在流程末端。成熟组织会把它们嵌入任务设计、过程评审和自动化规则中。例如，品牌禁区应在内容生成之前进入知识与规则层；客户授权状态应在CRM触达之前被系统检查；预算止损条件应在投放调整之前确定，而不是损失发生后再追责。

5.3 系统编排的最小结构

系统编排不是搭建复杂技术架构，而是将任务转换为可运行的责任链。最小结构包括七个组件。

表14：AI营销任务编排七组件

表14：AI营销任务编排七组件		
组件	需要定义的内容	缺失后的典型风险
目标与验收	经营目标、任务输出、指标和质量门	产出很多但无法判断价值
输入与上下文	数据、知识、规则、历史案例和限制	输出泛化、事实错误、品牌偏移
角色与分工	人员、系统、外部专家和AI员工职责	重复劳动、责任真空
步骤与交接	顺序、并行关系、状态与交付接口	任务卡住、版本混乱
判断与审批	哪些节点需要专业判断或管理批准	自动执行越权
异常与恢复	中断、数据缺失、冲突、错误和回滚路径	局部错误被高速放大
日志与学习	保存输入、修改、决策、结果和复盘	无法审计，经验不能沉淀

5.4 自动化深度必须服从任务风险

表15：任务风险与授权深度

表15：任务风险与授权深度			
风险等级	任务特征	适宜授权方式	人类责任要求
低风险	内部使用、可快速纠错、不涉及敏感数据	系统可独立完成并抽检	明确抽检人和错误阈值
中低风险	对外但影响有限、可撤回	系统生成或执行，人定期审核	设品牌与事实检查
中风险	涉及客户体验、预算或公开表达	系统形成建议，人逐项批准	指定业务责任人
高风险	涉及敏感数据、重大预算、监管声明	人主导，系统辅助	双重审核、完整留痕
极高风险	可能造成重大法律、声誉或人身影响	禁止自主执行	高层或专业委员会决策

人工智能人才的高阶表现，不是追求最大自动化，而是能够依据风险选择适当自动化深度。

5.5 异常处理是区分普通执行者与高级人才的关键

常规任务可以通过模板、示例和固定流程完成，异常才真正暴露能力。常见异常包括：

- 数据突然中断；
- 两个权威来源结论冲突；
- 顾客行为与调研表达不一致；
- 品牌热点迅速转为争议；
- 平台算法或归因规则变化；
- 自动化内容触发投诉；
- 预算消耗速度异常；
- 模型输出在特定人群中表现出偏差；
- 外部供应商未按约定提供证据；
- 高层在项目中途改变目标。

高级员工必须能够判断异常属于数据问题、任务问题、模型问题、流程问题、专业问题还是治理问题，并采取停止、降级、切换、补证、人工接管或重新定义任务等行动。

案例一：可口可乐“Create Real Magic”——扩大创意参与，不等于下放品牌责任

主体与时间： 可口可乐公司，2023年推出面向公众的“Create Real Magic”生成式创意平台，并在后续营销活动中持续探索人工智能辅助创意。

业务情境： 全球品牌希望利用生成式技术扩大创意参与度，让外部创作者使用可口可乐部分品牌资产形成作品，同时保持品牌识别和资产安全。

具体动作：

提供受控的品牌素材和创作环境；

允许创作者在限定资产范围内生成视觉作品；

通过筛选与人工评审选择可进一步传播的内容；

将公众创作与品牌活动连接，而不是允许系统直接代表品牌发布全部结果。

可核验结果： 该项目形成了全球范围的创作者参与和大量候选作品，并成为大型消费品牌公开探索生成式创意的早期案例。其直接价值主要体现为参与、创意探索和品牌话题，而非可以直接归因的长期收入增长。

条件边界：

品牌开放的是限定素材和受控参与，不是完整品牌决策权；

生成数量不能等同于品牌资产增长；

公众参与作品仍需人工筛选、版权和品牌审查；

成果不能外推为所有品牌都适合开放生成。

管理启示： 创意岗位的AI能力不应定义为“会生成更多视觉”，而应包括品牌资产边界设计、候选筛选、权利审查和对外发布责任。

案例二：Harvard Business School与宝洁“The Cybernetic Teammate”现场实验——个人能力与团队能力并非简单相加

主体与时间： Harvard Business School研究团队与宝洁公司合作开展的现场实验，研究成果于2025年前后公开传播。

业务情境： 研究关注生成式人工智能如何影响专业人员在产品创新与商业问题任务中的表现，以及个人使用人工智能后是否，能达到传统跨职能团队的部分效果。

具体动作：

将数百名宝洁专业人员分配到个人或团队、使用或不使用生成式人工智能等不同条件；

要求参与者完成接近真实工作的产品创新任务；

比较解决方案质量、时间、知识边界和情绪体验；

分析人工智能是否帮助不同专业背景人员跨越职能知识边界。

可核验结果： 公开研究显示，使用人工智能的个人能够在部分任务上达到或接近传统双人团队表现；人工智能还可能帮助商业与技术人员形成更综合的方案。研究样本为791名专业人员，是人机协作研究中较具代表性的企业现场实验之一。

条件边界：

任务处于受控实验环境，不能直接等同于长期业务绩效；

研究主要观察特定创新任务，不能外推到所有营销、管理或高风险决策；

人工智能提高方案表现，不代表取消专业团队；

对事实责任、执行落地、组织政治和长期顾客结果的影响仍需独立验证。

管理启示： 企业应把人工智能视为团队认知结构的一部分，但不能因为个人在受控任务中表现提升，就取消交叉评审、专业责任和系统条件。

第六章 能力证据五级阶梯：知道、完成、稳定、应变、带动

竹势 AI 营销智库出品

L1

知道

L2

指导下完成

L3

独立稳定交付

L4

处理异常并优化

L5

设计系统并带动他人

6.1 等级设计原则

能力等级应由四个变量共同决定：

独立性：需要多少指导和复核；

复杂度：能够处理多复杂的任务；

异常能力：是否能在偏离正常路径时恢复；

影响范围：只影响个人产出，还是影响团队和系统。

年限、职级、证书和自评可以作为背景信息，但不能代替工作证据。

表16：能力证据五级阶梯总表

表16：能力证据五级阶梯总表				
等级	能力定义	任务条件	可观察工作样本	典型授权
L1 知道	理解概念、规则与风险	解释性、低复杂度	概念说明、风险识别题	仅观察或辅助
L2 指导下完成	能按模板和指导完成任务	范围固定、异常较少	受控任务、修改记录	低风险任务，需复核
L3 独立稳定交付	能持续满足质量标准	常规真实任务	连续项目样本、质量数据	常规任务独立授权
L4 处理异常并优化	能诊断异常、调整流程并提高结果	跨角色、目标变化、数据冲突	异常复盘、优化前后证据	中高复杂任务授权
L5 设计系统并带动他人	能建立标准、机制与人才能力	组织级、多任务组合	系统设计、培养记录、团队结果	体系与重大项目设计权

6.2 一级：知道

一级人员能够：

解释基本概念；

识别明显风险；

描述任务流程；

知道何时寻求帮助。

一级不应获得独立对外发布、预算调整、客户自动触达或敏感数据处理权限。

表17：一级工作样本

表17：一级工作样本	
证据类型	合格表现
概念说明	能区分AI素养、工具熟练度与任务能力
风险识别	能指出事实、隐私、版权、品牌和自动执行风险
流程复述	能说明任务输入、输出、审批和责任
情景判断	遇到超出权限的情况能够停止并升级

6.3 二级：可在指导下完成

二级人员能在明确模板、资料、规则和评审者支持下完成任务。其核心证据不是一次成功，而是能够接受反馈并正确修改。

二级适合承担：

内容初稿；

研究资料整理；

标准化报表；

受控创意变体；

低风险社媒候选回复；

固定规则下的客户分群建议。

6.4 三级：可独立稳定交付

三级是岗位独立胜任的关键门槛。员工不仅能完成任务，还能在多个周期保持稳定质量。

三级评测应至少检查：

连续三次以上真实任务；

返工率和缺陷类型；

事实与品牌质量；

时间与资源控制；

是否能主动识别缺失输入；

是否留下足够的判断记录；

是否在权限范围内独立完成。

6.5 四级：可处理异常并优化

四级人员能在常规流程失效时工作。他们能够诊断异常，选择降级、切换或重新设计方案，并把经验转为流程改进。

表18：四级异常评测情景

表18：四级异常评测情景	
异常	评测重点
数据突然缺失	是否降低结论强度、补充替代证据或停止行动
品牌热点转为争议	是否快速识别风险并调整内容计划
投放成本异常	是否区分市场变化、素材疲劳与测量问题
客户投诉增加	是否暂停自动触达并检查分群、频率和授权
两个系统记录冲突	是否识别主数据和责任系统
高层改变目标	是否重设任务、指标和资源，而非机械继续

6.6 五级：可设计系统并带动他人

五级人员的价值主要体现为组织能力。他们能够：

从经营战略定义能力组合；

设计任务标准和质量门；

建立人机分工与授权；

培养和校准评审者；

设计异常管理和复证周期；

提高团队整体表现；

降低关键人才单点风险；

识别哪些问题应通过系统改造而非培训解决。

五级不等于管理职位。专业专家同样可以达到五级，例如设计品牌评审系统、实验框架或研究证据标准的资深人员。

6.7 等级不能按所有能力平均

一名人员可能在内容专业判断达到四级，在系统编排只有二级，在治理方面为三级。企业应形成能力剖面，而不是简单平均。

表19：能力剖面示例

表19：能力剖面示例		
能力维度	等级	授权含义
任务建模	L4	可设计复杂内容战役任务
研究与证据	L3	可独立完成常规证据核验
系统编排	L2	复杂自动化需运营专家支持
判断与质量	L4	可担任内容质量责任人
责任与治理	L3	可批准常规内容，不可批准高风险声明

6.8 再认证与能力折旧

能力不会因为技术更新而整体失效，但部分证据会折旧。

建议将能力分为：

稳定核心：顾客理解、问题定义、专业判断、证据意识、责任；

中期变化： workflow 设计、数据接口、评测方法；

快速变化：具体工具、模型、界面和功能。

复证周期可按风险设置：

表20：再认证周期建议

表20：再认证周期建议		
能力类别	建议周期	复证方式
通用AI素养	12个月	风险情景与政策更新
常规任务能力	12—18个月	工作样本抽检
高风险授权能力	6—12个月	异常演练、案例答辩
工具操作能力	3—6个月或版本变化后	快速实操
系统设计能力	以重大变更为触发	架构评审与事故复盘

BOARD BRIEF 09

第七章 十岗差异化能力投资组合

竹势 AI 营销智库出品

品牌

内容

创意

媒介

投放

社媒

CRM

研究

营销运营

管理

十个岗位共享MARKET-8共同语言，但专业纵深、AI新增能力、不可下放判断和证据任务必须分别定义。

7.1 品牌岗位

表21：品牌岗位能力组合

表21：品牌岗位能力组合	
类别	具体要求
共同底座	问题定义、证据核验、品牌数据理解、人机协同、风险升级
专业纵深	品牌定位、价值主张、品牌架构、长期记忆结构、声誉与一致性
AI新增能力	将品牌规范转为可执行知识与规则；评审大规模生成内容；识别品牌漂移；设计例外审批
不可下放判断	品牌承诺、定位变化、重大公共表达、价值观冲突、声誉风险
证据任务	品牌知识包、表达边界表、候选内容评审、品牌异常复盘

品牌人员不能只负责“统一语气”。高等级品牌人才需要判断短期流量是否损害长期资产，并有权否决不符合品牌责任的高表现内容。

案例三：中国平安——大规模智能化必须建立在治理与专业责任之上

主体与时间： 中国平安，长期在金融、保险、健康和客户服务场景推进人工智能与数字化能力，并在年度报告和公开披露中持续介绍相关应用。

业务情境： 金融服务企业需要在大量客户交互、营销、服务与风控任务中提高效率，同时受到个人信息、金融合规、客户公平和错误决策风险约束。

- 具体动作：
- 建设企业级数据与智能能力；
 - 将智能技术应用于客户服务、经营分析、销售支持和风险管理；
 - 在关键金融决策和客户服务环节保留专业规则、人工管理与审计要求；
 - 通过组织治理而非个人自由使用管理高风险应用。
- 可核验结果： 中国平安公开披露了广泛的智能化应用和数字化经营成果，但不同场景的业务结果、口径和归因各不相同，不能把技术覆盖直接解释为人才能力提升。

条件边界：

金融场景风险高，不能照搬到普通内容生产；

企业披露属于主体自报，需要与监管规则和独立证据共同理解；

系统能力属于组织资产，不能全部计入员工个人能力；

效率改善不自动等于客户价值、公平或长期品牌提升。

管理启示： 品牌、CRM与管理岗位必须把客户信任和责任治理列为不可下放判断，不能只考核智能化使用率。

7.2 内容岗位

表22：内容岗位能力组合

表22：内容岗位能力组合	
类别	具体要求
共同底座	任务建模、来源核验、品牌规范、渠道规则、版权意识
专业纵深	叙事结构、信息层级、事实表达、受众语言、编辑与改写
AI新增能力	设计内容任务包；建立事实与品牌检查；管理多版本生成；从表现数据提炼学习
不可下放判断	事实声明、核心观点、敏感表达、标题与正文一致性、是否值得发布
证据任务	一题多平台内容包、来源记录、编辑修改轨迹、发布后复盘

内容岗位不能以初稿数量衡量。更关键的指标包括采纳率、实质修改量、事实缺陷、品牌一致性、有效阅读、客户行动和内容资产复用。

7.3 创意岗位

表23：创意岗位能力组合

表23：创意岗位能力组合	
类别	具体要求
共同底座	问题定义、品牌理解、证据、协作、权利与风险意识
专业纵深	创意概念、符号与叙事、视觉系统、制作可行性、文化敏感度
AI新增能力	扩展创意空间；管理参考与生成资产；评测同质化；建立可生产的变体系统
不可下放判断	核心创意命题、文化与伦理边界、最终视觉方向、重大肖像与权利使用
证据任务	创意领地、概念板、候选筛选依据、制作包、投放后学习

人工智能可以生成大量视觉变体，但不能用“看起来新颖”替代创意是否服务战略、是否可识别、是否具有文化合理性。

7.4 媒介岗位

表24：媒介岗位能力组合

表24：媒介岗位能力组合	
类别	具体要求
共同底座	数据口径、任务目标、渠道规则、证据和治理
专业纵深	受众覆盖、频次、触点角色、渠道组合、价格与资源谈判
AI新增能力	跨平台数据整理；情景模拟；动态资源建议；识别算法偏差和测量断点
不可下放判断	重大媒介组合、品牌安全、长期渠道关系、跨平台资源取舍
证据任务	媒介组合方案、覆盖频次模型、渠道边界说明、事后增量分析

媒介岗位必须防止把平台自动建议当作企业最优方案。平台通常优化其可观察指标，企业需要对利润、品牌和全渠道结果负责。

7.5 投放岗位

表25：投放岗位能力组合

表25：投放岗位能力组合	
类别	具体要求
共同底座	数据质量、预算责任、实验意识、日志与审批
专业纵深	出价、预算、定向、素材、落地页、归因、增量与边际回报
AI新增能力	自动诊断；异常预警；建议生成；实验编排；预算变更草稿与回滚
不可下放判断	重大预算迁移、止损阈值、归因解释、高风险人群定向
证据任务	账户体检、假设树、实验方案、预算审批包、异常复盘

案例四：Meta广告投放系统——平台自动化提高执行效率，也可能弱化企业解释能力

主体与时间： Meta，持续推进Advantage系列自动化广告产品，并在2022年后强化人工智能驱动的受众、版位、预算和创意优化。

业务情境： 广告主希望减少手工设置，提高跨版位和受众匹配效率。

具体动作：

将受众扩展、版位、预算和创意组合交给平台算法优化；

提供自动化建议和一体化投放产品；

用平台内行为和转化信号持续调整投放。

可核验结果： Meta官方案例和产品材料报告了部分广告主的效率改善，但这些数字通常属于供应商自报，且结果受行业、数据质量、创意、归因窗口和市场环境影响。

条件边界：

平台优化目标不一定等于企业利润或长期客户价值；

自动化降低了部分操作要求，却提高了实验设计、数据质量和结果解释要求；

平台归因不能自动证明增量；

广告主仍需控制品牌安全、预算和高风险定向。

管理启示： 投放人才的价值应从手工调参转向实验、预算责任、数据解释和跨平台经营判断。

7.6 社媒岗位

表26：社媒岗位能力组合

表26：社媒岗位能力组合

类别	具体要求
共同底座	平台规则、品牌表达、事实核验、客户尊重、升级机制
专业纵深	社群语境、互动节奏、平台文化、热点判断、危机识别
AI新增能力	舆情聚类；候选回复；内容排期；异常监测；人工接管编排
不可下放判断	危机回应、敏感评论、社会议题参与、删除或封禁重大决定
证据任务	社媒响应包、热点判断、危机模拟、评论分类与接管记录

案例五：微软Tay聊天机器人——高速互动会高速放大治理缺陷

主体与时间： 微软，2016年在Twitter推出Tay聊天机器人，上线后不足一天即被下线。

业务情境： 微软尝试通过开放社交互动测试能够学习年轻人语言风格的聊天机器人。

具体动作：

允许机器人在公开社交平台与用户高频互动；

系统从互动内容中学习表达；

未能充分防范恶意用户集中诱导和有害内容放大；

事件发生后迅速停止服务并公开回应。

可核验结果： Tay在短时间内产生大量不当和冒犯性表达，成为自动化社交互动失控的经典案例。

条件边界：

2016年的技术环境与今天不同，但开放互动、恶意操纵和高速放大的治理问题仍然存在；

事件不能用于证明所有社交自动化都不可行；

关键问题不是单纯模型能力，而是测试、边界、监控、停止权和责任设计不足。

管理启示： 社媒岗位的高级能力必须包括异常识别、人工接管和危机升级，而不只是自动回复效率。

7.7 CRM岗位

表27：CRM岗位能力组合

表27：CRM岗位能力组合	
类别	具体要求
共同底座	数据质量、个人信息保护、客户授权、任务建模和治理
专业纵深	客户生命周期、分群、触达、忠诚、流失、客户价值和服务体验
AI新增能力	动态分群；下一行动建议；会话摘要；触达草稿；客户信号识别
不可下放判断	授权合法性、高影响客户处置、敏感客户分群、重大服务补救
证据任务	生命周期方案、分群规则、授权检查、触达实验、投诉复盘

案例六：星巴克Deep Brew——个性化能力依赖数据、门店和运营系统共同工作

主体与时间： 星巴克，2019年前后公开介绍Deep Brew人工智能平台，并持续用于个性化、门店运营和预测相关场景。

业务情境： 星巴克拥有会员、移动应用、门店交易和产品数据，希望提高个性化推荐和运营效率。

具体动作：

结合会员历史、门店、时间、天气和产品等信息形成个性化建议；

将智能能力嵌入移动应用与运营流程；

让门店与数字渠道共同承接推荐结果；

通过持续数据反馈更新客户理解。

可核验结果： 星巴克公开披露Deep Brew支持个性化和门店决策，但具体财务贡献难以从其他会员、门店和产品因素中完全分离。

条件边界：

个性化能力依赖高质量第一方数据和成熟会员体系；

推荐效果不能简单归因于算法；

不同国家的数据和隐私要求不同；

个性化不能损害客户控制感和品牌体验。

管理启示： CRM人才能力必须与数据、门店执行和客户授权共同评测，不能把组织基础设施算作个人能力。

7.8 研究岗位

表28：研究岗位能力组合

表28：研究岗位能力组合	
类别	具体要求
共同底座	问题定义、来源评价、数据伦理、表达边界、复核
专业纵深	研究设计、抽样、访谈、观察、实验、综合分析与洞察
AI新增能力	来源发现；资料编码；证据冲突识别；假设扩展；可追溯研究档案
不可下放判断	研究问题、样本适用性、因果解释、洞察成立与否
证据任务	研究计划、证据库、编码记录、洞察链、限制说明

案例七：浙江大学与阿里巴巴等中国研究与产业主体的人工智能人才实践——教育、产业和任务必须连接

主体与时间： 浙江大学、阿里巴巴等中国高校与企业近年来持续推进人工智能课程、产业实践和复合型人才培养。

业务情境： 人工智能发展速度快，单纯课堂知识难以覆盖真实产业任务；企业也需要既理解专业领域又能使用智能技术的人才。

具体动作：

高校设置人工智能通识、专业课程和交叉培养；

企业提供平台、真实问题、竞赛、实习或产业项目；

通过项目成果而非单纯考试观察部分应用能力；

强调技术与行业问题结合。

可核验结果： 中国高校和企业已形成大量人工智能课程与产业合作项目，但结业规模、竞赛成绩和项目数量不能直接证明长期岗位绩效。

条件边界：

学生项目与企业真实责任仍存在差异；

产业平台可能影响课程方向；

不同专业迁移距离不同；

教育成果需要进入真实工作后复证。

管理启示： 企业招聘研究和营销人才时，应重视真实任务样本、问题定义与证据能力，而不是仅看专业名称和课程证书。

7.9 营销运营岗位

表29：营销运营岗位能力组合

表29：营销运营岗位能力组合	
类别	具体要求
共同底座	流程、数据、协同、权限、质量和复盘
专业纵深	需求管理、资源配置、流程设计、技术栈、指标体系、供应商管理
AI新增能力	工作流编排；任务路由；知识管理；监控与异常恢复；成本治理
不可下放判断	组织流程、权限架构、核心系统选择、重大供应商依赖
证据任务	端到端流程图、能力账本、运行指标、事故复盘、系统退出方案

营销运营岗位是连接个人能力、团队能力和组织系统条件的核心岗位。其绩效不能只看上线多少工具，而应看需求周期、返工、质量、可用性、成本、风险和能力沉淀。

7.10 管理岗位

表30：管理岗位能力组合

表30：管理岗位能力组合	
类别	具体要求
共同底座	经营目标、证据、资源、协同、风险与公平
专业纵深	战略取舍、组织设计、预算、人事决策、跨部门治理
AI新增能力	建立任务组合；设计授权梯度；配置人、AI员工与专家；审查能力证据
不可下放判断	战略目标、重大预算、组织责任、人员评价和高风险授权
证据任务	人才组合方案、授权矩阵、投资决策、能力盘点、申诉处理

管理者不能把人工智能人才建设完全委托给人力资源部门。管理者必须亲自决定哪些判断权不可下放、哪些能力需要冗余、哪些岗位应重构以及哪些问题属于系统责任。

案例八：IBM Watson for Oncology——能力宣传超过真实边界会产生反噬

主体与时间： IBM Watson for Oncology在2010年代被多家医疗机构测试和引入，随后受到关于建议质量、训练数据、适用性和商业表现的广泛质疑。

业务情境： 医疗机构希望使用人工智能辅助复杂诊疗决策，供应商强调系统能够整合知识并提供建议。

具体动作：

将医学文献、专家知识和病例信息用于生成治疗建议；

与医院合作开展测试和部署；

在部分市场以智能决策支持能力进行推广；

后续因适用性、可靠性和商业问题缩减相关业务。

可核验结果： 系统在部分研究和机构中显示出一定一致性或辅助价值，但也出现不安全或不适用建议的报道，并最终未达到早期市场预期。

条件边界：

医疗决策风险远高于一般营销任务；

不同医院、国家和病例条件差异巨大；

专家共识一致性不等于真实临床结果；

供应商宣传不能替代独立验证。

管理启示： 企业不能因系统展示效果良好就扩大授权。高风险任务必须以真实环境证据、异常案例和专业责任为基础。

第八章 招聘工作样本与评测：从“说自己会”到“证明能胜任”

竹势 AI 营销智库出品

任务 Brief

允许资源

过程证据

异常注入

评审量规

复盘答辩

8.1 工作样本为何优于泛化问答

传统面试容易奖励表达流畅、术语熟悉和经验包装。AI营销岗位尤其容易出现候选人能够描述工具、展示漂亮产出，却无法说明任务目标、事实依据、修改过程和责任边界。

工作样本的优势在于接近真实工作，但必须设计得当。

表31：AI营销招聘评测组合

表31：AI营销招聘评测组合		
评测方式	主要观察	局限
结构化经历面试	历史行为、角色和结果	候选人可能夸大个人贡献
工作样本	实际任务表现	任务设计和资源条件影响结果
过程答辩	判断依据、取舍与边界	评委需要专业训练
历史作品核验	长期质量与复杂度	保密限制、团队贡献难拆分
异常注入	应变、风险和升级	需防止过度戏剧化
参考核验	稳定性、协作和责任	参考人存在偏差

8.2 工作样本设计七原则

- 对应真实岗位任务；
- 时间控制合理；
- 允许使用岗位中真实可用资源；
- 明确禁止和限制；
- 不要求候选人免费完成可直接商用成果；
- 设置可解释的评审量规；
- 允许答辩和申诉。

8.3 十岗招聘工作样本

表32：十岗差异化工作样本

表32：十岗差异化工作样本

岗位	工作样本任务	核心观察
品牌	评审一组品牌候选内容并提出边界规则	长期一致性、品牌责任、否决依据
内容	将一份资料转成多平台内容并完成事实表	叙事、编辑、证据、渠道适配
创意	从同一商业问题形成三个创意领地	区分度、战略连接、筛选能力
媒介	在预算约束下设计渠道组合	覆盖、频次、触点角色和权衡
投放	诊断一组异常账户数据	实验、归因、预算与止损
社媒	处理热点、投诉和普通互动混合队列	语境、升级、速度与品牌
CRM	设计客户分层与触达规则	授权、生命周期、客户价值
研究	基于相互冲突来源形成判断	证据质量、限制和洞察
营销运营	重构一条高返工流程	系统思维、责任、成本和恢复
管理	设计团队能力组合与授权矩阵	经营取舍、公平、风险和组织设计

8.4 评审量规

表33：工作样本通用评审量规

表33：工作样本通用评审量规

维度	权重建议	高分表现
问题定义	15%	识别真实问题并说明假设
专业判断	20%	取舍有明确专业依据
证据与数据	15%	来源可靠，限制透明
结果质量	20%	输出可用且符合岗位标准
人机分工	10%	工具使用适度、可解释
异常与风险	10%	能识别、停止、升级和恢复
复盘与表达	10%	能解释修改与下一步

权重必须按岗位调整。创意岗位不能让一致性和格式压倒原创判断；投放岗位不能让演示表达压倒预算与实验能力；管理岗位不能只看方案完整，而忽视组织公平与执行条件。

8.5 防止评测失真

常见失真包括：

题目提前泄露；

候选人使用外部人员代做；

不同候选人获得资源不同；

评委偏好特定表达风格；

任务过度接近企业真实商业需求；

候选人因残障、语言或设备条件处于不利地位；

只评最终产出，不看过程；

只评一次任务，不看历史稳定性。

企业应保存评分理由，而不只保存分数。

第九章 培养、能力邻接、迁移与团队组合

竹势 AI 营销智库出品

PERSON

个人

任务与判断证据

TEAM

团队

互补、复核与备份

SYSTEM

系统

数据、流程与权限

9.1 培训必须嵌入真实任务

知识课程可以建立共同语言，但能力主要通过真实任务、反馈、复盘和重复实践形成。

建议采用“20—30—50”结构：

20%用于概念、规则和示范；

30%用于受控练习和反馈；

50%用于真实项目与复盘。

比例不是固定标准，核心是防止培训全部停留在课堂。

9.2 能力邻接与迁移图

能力迁移不应只看岗位名称，而要比对五项因素：

现有能力基础；

目标任务要求；

训练距离；

错误风险；

业务窗口。

表34：能力迁移决策矩阵

表34：能力迁移决策矩阵

邻接程度	风险	业务窗口	建议方式
高邻接	低	宽松	在岗自学与任务实践
高邻接	高	适中	项目训练加专家复核
中邻接	低	适中	结构化训练与轮岗
中邻接	高	较短	外部专家加内部影子学习
低邻接	低	宽松	长周期培养或岗位转换
低邻接	高	很短	外部招聘或共享专家

9.3 十种典型迁移路径

表35：营销岗位能力迁移示例

表35：营销岗位能力迁移示例

原岗位	目标能力方向	邻接优势	主要缺口
内容	社媒运营	平台表达和内容生产	实时互动、危机和社群判断
内容	品牌	叙事与语言	定位、长期资产和组织影响
研究	品牌策略	洞察和证据	品牌决策与创意转化
媒介	投放	渠道和数据	账户实验、预算细节
投放	营销运营	数据和流程意识	跨职能治理与系统设计
CRM	研究	顾客数据和分群	研究设计、质性理解
社媒	CRM	客户互动经验	生命周期、授权和数据治理
创意	内容领导	概念和表达	编辑系统、稳定产能管理
营销运营	管理	系统和协同	战略、人才和经营责任
品牌	管理	长期视角和跨部门影响	财务、资源和组织设计

9.4 30天实战认证包

表36：30天实战认证结构

表36：30天实战认证结构

模块	内容	证据
任务说明	真实经营目标、对象、边界和验收	正式任务卡
允许资源	数据、知识、工具、专家和时间	资源清单
过程记录	关键输入、修改、决策和系统调用	过程日志
产出样本	中间成果与最终交付	版本档案
异常注入	数据缺失、目标变化、冲突或风险	应变记录
评审量规	专业、证据、质量、风险和结果	多评委评分
复盘答辩	解释取舍、错误和下一步	答辩记录
再认证	明确有效期和复证条件	认证档案

30天认证不能证明长期绩效，但可以比一次考试更有效地支持初步授权。

9.5 个人—团队—系统三层能力账本

表37：三层能力账本

表37：三层能力账本

层级	记录内容	典型问题	责任主体
个人能力	任务、判断、证据、协作和异常能力	员工能否胜任	员工与直属管理者
团队能力	互补、复核、交接、备份和共享专家	团队是否存在关键缺口	团队负责人
系统条件	数据、流程、知识、工具、权限和治理	环境是否支持稳定交付	业务、技术、数据和管理层

表38：失败归因分账示例

表38：失败归因分账示例			
失败现象	个人可能责任	团队可能责任	系统可能责任
内容事实错误	未核验来源	无专业复核者	无可信知识库
投放预算超支	未遵守止损	无轮班监控	告警与权限失效
CRM投诉增加	错误分群判断	客服未及时接管	授权状态未接入
战役延期	任务管理不足	交接责任不清	审批流程过长
品牌表达漂移	忽视规范	品牌评审缺位	品牌知识未结构化

9.6 团队组合的五种角色

复杂AI营销任务通常需要五种能力角色，不一定对应五个全职岗位：

任务所有者：对经营目标负责；

专业判断者：对品牌、内容、投放或客户专业质量负责；

证据与数据者：对来源、口径和测量负责；

系统编排者：对流程、工具、权限和恢复负责；

治理复核者：对法律、安全、公平和重大风险负责。

小团队可以一人承担多个角色，但每个角色必须被明确拥有。

9.7 人、外部专家与AI员工的投资组合

表39：五种能力供给方式

表39：五种能力供给方式			
供给方式	最适合	不适合	管理重点
在岗培养	高频、邻接度高、长期需要	极短窗口、高风险陌生能力	项目实践和反馈
跨岗迁移	有相近专业基础的人才	专业距离过大	明确缺口与过渡授权
外部招聘	长期关键、内部缺失	需求短暂或高度不确定	工作样本和文化适配
共享专家	低频、高专业、高风险	日常高频执行	接口、责任和知识转移
AI员工或自动化	高频、规则清晰、可验证、可回滚	模糊、高影响、不可逆任务	权限、监控和人工接管

9.8 团队冗余与关键能力备份

企业不应以极致人效为由消除所有能力冗余。以下能力至少需要双人或替代机制：

重大品牌表达审核；

高预算投放管理；

客户数据权限管理；

核心 workflow 维护；

危机社媒响应；

关键研究结论复核；

人员能力评审与申诉。

适度冗余不是浪费，而是业务连续性与判断质量的保险。

9.9 培养投资的停止条件

不是所有培养都值得继续。出现以下情况时，应重新评估路径：

目标能力与员工现有能力距离过大；

- 长期缺少真实实践机会；
 - 业务窗口已经关闭；
 - 岗位需求频率不足；
 - 员工不愿承担目标岗位责任；
 - 系统条件长期不具备；
 - 外部专家或自动化明显更经济；
 - 培养成本超过长期能力价值。
- 停止某条培养路径不等于否定员工价值，而是把人员配置到更适合的能力组合。

9.10 管理含义

- 培养、迁移和团队组合必须围绕经营任务，而不是围绕培训目录。管理者需要同时回答：
- 哪些能力应成为全员基线；
 - 哪些能力只需要少数专家；
 - 哪些能力可以由AI员工提供；
 - 哪些判断必须保留在人类岗位；
 - 哪些组织条件必须先补齐；
 - 哪些员工适合跨岗迁移；
 - 哪些任务应通过外部专家解决；
 - 哪些认证结果可以转化为项目授权。
- 只有当这些答案进入资源、岗位和项目决策，能力模型才真正成为经营工具。

第十章 让能力模型进入经营机制：岗位卡、授权、绩效、晋升、薪酬与人才盘点

竹势 AI 营销智库出品

岗位卡

项目授权

绩效

晋升

薪酬

人才盘点

10.1 能力模型只有改变决策，才具有经营价值

能力模型最常见的失败，不是框架不完整，而是框架没有进入真实管理动作。企业可能完成能力词典、等级定义、测评问卷和培训地图，却仍然按照旧方式招聘、分工、授权、考核和晋升。在这种情况下，能力模型只是人力资源文件，不是经营系统。

要让模型进入组织日常，至少需要建立六条连接：

能力与岗位卡连接；

能力与项目授权连接；

能力与绩效证据连接；

能力与晋升评议连接；

能力与薪酬及稀缺性连接；

能力与人才盘点、继任和团队重组连接。

这六条连接不能机械地把每项能力变成分数。能力数据的用途，是帮助管理者更准确地作出岗位与资源决策，而不是制造一个看似精确、实则掩盖条件差异的总排名。

10.2 岗位卡：从职责清单升级为“任务—判断权—证据”契约

传统岗位说明通常包括岗位目的、主要职责、任职要求和汇报关系，但很少说明人工智能参与后工作如何变化。新的岗位卡应增加五项内容：

该岗位对哪些经营任务负责；

哪些步骤可交给系统或AI员工；

哪些判断不可下放；

哪些能力证据才能获得授权；

当系统或数据失效时，该岗位承担什么恢复责任。

表40：AI营销岗位卡标准结构

表40：AI营销岗位卡标准结构		
模块	必须回答的问题	示例
岗位经营使命	该岗位最终改善什么经营结果	提高品牌资产积累与传播一致性
核心任务组合	该岗位持续承担哪些完整责任单元	品牌定位、战役评审、表达治理
MARKET-8要求	八维能力分别达到什么等级	判断与质量L4，系统编排L3
专业纵深	哪些营销专业能力不可缺少	品牌架构、定位、声誉管理
AI协作范围	哪些任务可由系统辅助或执行	资料整理、候选内容扫描
不可下放判断	哪些决定必须由岗位责任人作出	核心品牌承诺、重大公共表达
工作证据	用什么样本证明胜任	品牌评审记录、异常复盘
授权边界	可以独立完成什么，需要谁复核	常规内容可批准，重大声明需升级
系统依赖	需要哪些数据、知识、权限和流程	品牌知识库、法务审核接口
复证条件	多久或在什么变化后重新评测	12个月或品牌战略重大调整后

岗位卡不是对员工无限加责。若岗位要求员工对品牌一致性负责，企业就必须提供品牌资料、评审权和停止发布的权限；若要求投放人员对预算负责，就必须提供及时数据、止损机制和变更审批通道。责任必须与资源、信息和权力相匹配。

10.3 项目授权：按任务风险和能力证据动态分配判断权

职位名称不能自动证明某人具备所有项目授权。一名资深内容负责人可能从未管理自动发布系统，一名年轻投放人员可能在实验设计方面已达到较高水平。企业需要把授权从职位附属物转化为“能力证据与任务风险的匹配结果”。

表41：项目授权四类权力

表41：项目授权四类权力		
权力	定义	典型例子
建议权	可以形成方案，但不能批准执行	形成预算调整建议
制作权	可以完成工作成果，但需他人评审	完成对外内容草稿
批准权	可以在明确边界内批准执行	批准常规品牌内容
停止权	发现异常时可立即中止任务	暂停异常客户触达

很多组织只关注批准权，却忽视停止权。人工智能系统可以高速执行，若一线专业人员没有停止权，异常会在层层请示中被放大。高风险任务必须明确谁可以立即暂停，而不必等待最终责任人在线。

表42：能力等级与授权深度

表42：能力等级与授权深度				
能力等级	建议权	制作权	批准权	停止权
L1 知道	可在观察下参与	不独立拥有	无	发现明显风险可上报
L2 指导下完成	可提出受控建议	低风险任务，需复核	无或极有限	可暂停本人执行
L3 独立稳定交付	可独立提出常规方案	常规任务独立	可在既定边界内批准	可暂停所属任务
L4 处理异常并优化	可提出复杂方案	跨角色任务独立	可批准中等风险任务	可暂停端到端流程
L5 设计系统并带动他人	可设计决策机制	可配置任务系统	可设定授权规则	可建立组织级停机机制

授权不是永久荣誉。出现重大缺陷、岗位变化、系统升级、法规变化或长期缺少实践时，应重新评估。企业也不应因一次错误永久取消授权，而应区分疏忽、系统缺陷、合理试验失败和故意违规。

10.4 绩效：从产出活动转向“任务质量、经营贡献与系统责任”

人工智能时代最危险的绩效做法，是继续用数量奖励岗位，却让系统大幅降低数量成本。若内容人员仍按稿件数量考核，组织会得到更多低价值内容；若研究人员按报告数量考核，组织会得到更多摘要而非洞察；若投放人员按平台回报考核，可能诱导短期归因优化而非真实增量；若营销运营按上线工具数量考核，会形成工具堆叠。

绩效应至少分为四层：

- 任务结果：工作成果是否达到约定目标；
- 质量与责任：事实、品牌、客户、预算和风险是否受控；
- 能力与行为：是否展示相应等级的独立性和异常处理；
- 系统贡献：是否减少返工、沉淀知识、改善流程并带动他人。

表43：AI营销绩效四层结构

表43：AI营销绩效四层结构

层级	评审内容	不宜替代为
经营结果	收入、利润、客户价值、品牌、速度或风险结果	单纯使用率
任务质量	采纳、准确、稳定、及时和可恢复	生成数量
判断责任	关键取舍、异常处置、风险升级	工具操作记录
系统贡献	知识沉淀、流程改善、人才带动	会议发言次数

绩效归因必须保留边界。营销结果受到产品、价格、销售、市场环境和渠道变化影响，不能把全部结果归给单一岗位；同样，也不能因结果多因素决定，就完全不要求员工对任务质量和可控贡献负责。

表44：十岗绩效证据建议

表44：十岗绩效证据建议			
岗位	结果证据	质量与责任证据	系统贡献证据
品牌	品牌认知、偏好、资产一致性	重大表达评审、品牌风险	品牌规则与评审体系
内容	有效阅读、行动、资产复用	事实缺陷、返工、品牌一致性	内容模板与知识沉淀
创意	测试表现、记忆度、可生产性	权利风险、同质化、战略一致	创意学习库与评审机制
媒介	增量覆盖、有效频次、资源效率	品牌安全、口径一致性	跨渠道测量机制
投放	增量结果、边际回报、预算控制	止损、异常、归因解释	实验与预警体系
社媒	有效互动、客户响应、风险控制	危机处置、错误回复	分类、接管和升级机制
CRM	留存、复购、客户价值、服务结果	授权、退订、投诉	分群与旅程规则
研究	决策采纳、预测或洞察价值	来源、限制、可复核性	证据库和研究规范
营销运营	周期、返工、稳定性、成本	权限、事故、供应商风险	工作流和能力资产
管理	经营组合结果、资源效率	公平决策、重大风险	团队能力与继任体系

10.5 晋升：评审影响范围，而不是奖励熟练执行

晋升应反映员工承担更复杂责任和扩大组织影响的能力。人工智能提高个人产出后，企业尤其要防止把“一个人做得更多”直接等同于“可以管理更大范围”。

专业晋升和管理晋升应保持双通道。

表45：专业与管理双通道

表45：专业与管理双通道		
层级	专业通道证据	管理通道证据
初级	能在指导下完成任务	能管理本人任务
中级	能独立稳定交付	能协调小型项目
高级	能处理复杂异常	能配置跨角色资源
专家	能建立专业标准与系统	能设计团队与授权机制
首席/负责人	能影响企业级专业能力	能承担战略、人才和治理责任

晋升答辩应重点审查以下证据：

处理过哪些复杂任务；

哪些结果由本人真正负责；

如何识别和处理异常；

哪些经验被转成团队标准；

如何培养、评审或带动他人；

是否曾主动停止错误方向；

如何看待一次失败；

哪些系统条件限制了结果；

下一层级将扩大何种责任。

仅展示成功项目容易制造幸存者偏差。高级晋升材料必须包含至少一个失败、受限或重大纠偏案例。

10.6 薪酬：奖励稀缺责任，不奖励工具热度

某项工具短期热门，可能导致市场薪酬快速上升，但产品功能会扩散，操作门槛会下降。企业应把薪酬溢价放在更持久的稀缺性上：

能承担高影响专业判断；

- 能连接经营目标与系统执行；
- 能处理异常和重大风险；
- 能建立组织标准；
- 能培养他人；
- 能在数据、品牌、客户和预算之间作出可信取舍；
- 能降低企业对单一人员或供应商的依赖。

表46：薪酬价值的四类来源

表46：薪酬价值的四类来源

价值来源	是否稳定	薪酬含义
特定工具操作	低至中	可给予短期技能津贴，不宜形成永久高溢价
稳定任务能力	中	应进入岗位等级和绩效
专业判断责任	高	应体现岗位价值和风险责任
系统设计与人才带动	高	应体现专家或管理层级价值

对于短期紧缺技能，可以设置有期限的技能津贴，并明确复审日期，避免工具折旧后薪酬结构固化。

10.7 人才盘点：从“高潜—高绩效”转向“能力—责任—系统适配”

传统九宫格常用绩效和潜力两个维度，但在人工智能转型中还需要增加：

- 当前任务能力；
- 未来目标任务；
- 能力邻接距离；
- 授权风险；
- 团队关键性；
- 系统依赖；

迁移意愿；

能力折旧速度。

表47：AI营销人才盘点六问

表47：AI营销人才盘点六问	
盘点问题	管理用途
他现在能独立承担什么任务	当前授权
他在哪些异常场景中仍能稳定工作	复杂项目配置
哪些表现依赖特定系统或团队	避免高估个人
哪些能力与未来岗位高度邻接	培养和迁移
哪些判断一旦离开他就无人承担	继任与冗余
哪些能力正在因技术变化折旧	再认证与投资调整

人才盘点的结果不应只有“保留、培养、淘汰”。更有效的行动包括：调整任务、补充系统条件、改变团队组合、降低或扩大授权、安排共享专家、建立备份和重设培养路径。

BOARD BRIEF 13

第十一章 公平、隐私、员工沟通与申诉

竹势 AI 营销智库出品

能力数据影响员工利益，就必须透明、最小必要、人工复核并可申诉。

不以自动评分直接作重大人事决定；不把组织条件缺口记成个人失败。

11.1 能力治理是一项劳动治理，不只是测评技术

当企业开始记录员工的人工智能使用、工作过程、修改轨迹、任务表现和异常处理时，容易产生一种错觉：数据更多，评价就更客观。事实上，数据的采集方式、任务机会、系统条件和算法规则都可能产生新的不公平。

能力治理必须遵循四个基本原则：

目的明确；

最小必要；

人工复核；

可解释、可申诉。

11.2 公平风险来自哪里

表48：能力评测的主要公平风险

表48：能力评测的主要公平风险

风险	形成机制	控制措施
任务机会不均	有人获得高价值项目，有人长期做低风险执行	记录机会差异，提供轮岗与替代证据
资源条件不均	不同团队的数据、工具和支持不同	使用三层能力账本分账
评委偏差	偏好特定表达、背景或风格	结构化量规、多评委、校准会议
自动评分偏差	系统以历史模式定义“优秀”	禁止自动作出重大人事决定
监控行为扭曲	员工为指标制造可见活动	少看使用频次，多看任务证据
残障与语言障碍	统一任务形式造成不必要困难	提供合理便利和等价任务
申诉成本过高	员工不敢挑战评分	独立复核、禁止报复
供应商黑箱	外部工具无法说明评分逻辑	要求可审计、退出和人工替代

11.3 隐私与个人信息边界

中国《个人信息保护法》确立了明确目的、最小范围、公开透明、质量保障和安全保护等原则。涉及员工个人信息的能力评测，应有明确的管理目的和制度依据，避免把工作监控扩展为无限制行为追踪。

企业不得因为技术上能够采集，就默认可以采集。以下数据尤其需要谨慎：

私人对话与非工作内容；

与岗位无关的情绪或人格推断；

精细到不必要程度的键盘、鼠标和屏幕监控；

健康、家庭、身份等敏感信息；

未经说明的跨系统行为画像；

使用个人账号完成工作留下的数据；

与能力任务无关的社交关系分析。

表49：能力数据用途边界

表49：能力数据用途边界		
数据类型	可接受用途	禁止或高风险用途
工作样本	评估任务能力	未经同意对外公开
评审记录	授权、培养和复证	用于无关的人格推断
过程日志	复盘关键判断和异常	持续监视全部工作行为
使用记录	排查采用和支持问题	直接等同能力或绩效
客户数据操作记录	安全审计	扩展用于员工非必要画像
申诉材料	人工复核	对申诉者实施不利待遇

能力数据应设置保存期限。已过期、无继续用途或员工离职后不再需要的数据，应依据制度删除或匿名化。

11.4 员工沟通：先解释工作为何变化，再解释如何评测

企业若以“人工智能能力认证”突然进入员工绩效和晋升，容易被理解为裁员筛选、隐性监控或替代计划。有效沟通应回答七个问题：

为什么岗位工作正在变化；

哪些能力仍然长期重要；

哪些任务将由系统承担；

哪些判断仍由人负责；

能力如何被评测；

组织将提供什么资源；

结果如何影响授权、发展和人事决定。

表50：员工沟通四阶段

表50：员工沟通四阶段		
阶段	沟通重点	管理者应避免
启动前	目标、边界、非惩罚性盘点原则	宣称所有人必须快速转型
盘点期	证据要求、隐私和申诉	偷偷收集行为数据
认证期	任务规则、资源和公平	把一次失败定义为不胜任
应用期	授权、发展、复证和调整	只公布分数，不解释决策

11.5 申诉机制

能力评测进入重大人事决策后，员工必须获得有效申诉渠道。

表51：能力评测申诉流程

表51：能力评测申诉流程	
步骤	处理要求
提交异议	员工可对事实、证据、评分和系统条件提出异议
暂停重大决定	对争议较大的结果，暂缓不可逆人事动作
独立复核	由未参与初评的专业人员和人力代表复核
补充证据	允许提交历史样本、条件说明或等价任务
书面结论	说明维持或调整评分的理由
纠正与恢复	修正授权、档案或人事决定
反报复保护	不得因合理申诉降低机会或评价

申诉不是削弱管理权，而是提高能力数据质量。员工往往掌握系统故障、资源差异和任务背景，申诉可以帮助企业发现错误归因。

11.6 自动化人事决策的边界

人工智能可以辅助整理证据、发现评分差异、提醒复证，但不应独立作出招聘拒绝、降薪、晋升否决、调岗或解除劳动关系等重大决定。

人的复核不能只是形式。复核者必须有权改变结果，并能访问足够证据。若管理者只看到系统总分，不理解数据来源和限制，就不构成实质复核。

第十二章 14天盘点、30天实战认证、90天团队重组路线

竹势 AI 营销智库出品

DAY 01-14

任务盘点

任务图与三层分账

DAY 15-44

实战认证

样本、异常与答辩

DAY 45-90

团队重组

岗位、授权与复证

12.1 路线总览

实施不应从建立庞大能力词典开始，而应选择真实经营任务，形成最小闭环。

表52：14/30/90天路线

表52：14/30/90天路线

阶段	核心目标	主要交付物	阶段门
14天能力盘点	看清任务、人员与系统缺口	任务地图、三层账本、风险清单	是否找到高价值试点
30天实战认证	用真实任务验证能力与授权	工作样本、等级证据、授权建议	是否具备稳定交付证据
90天团队重组	重构岗位、团队、流程与能力供给	新岗位卡、团队组合、复证机制	是否形成可运行体系

12.2 前14天：能力盘点

第1—3天：确定经营任务

选择一至三个符合以下条件的任务：

与经营目标直接相关；

周期足够短；

有真实数据和使用场景；

可以观察结果；

风险可以控制；

涉及多种能力，而非单纯工具操作。

不宜选择纯展示性任务，例如“生成一批海报”或“让所有人体验某个产品”。更适宜的任务包括：

新产品内容到线索闭环；

投放账户异常诊断；

CRM沉默客户激活；

品牌内容质量治理；

竞品变化到策略响应；

社媒舆情到人工接管。

第4—7天：建立任务地图

使用“输入—动作—判断—输出—异常”拆解，并标记MARKET-8要求。

表53：14天任务盘点卡

表53：14天任务盘点卡	
字段	内容
经营目标	任务服务的结果
当前基线	时间、成本、质量、结果与风险
关键输入	数据、知识、规则和权限
关键判断	不可机械执行的取舍
当前人员	谁实际承担任务
AI与工具	当前由系统承担的部分
异常	历史上最常见失败
证据	可获得的工作样本
系统缺口	数据、流程、权限和治理问题
试点建议	适合认证或改造的任务

第8—11天：个人—团队—系统分账

禁止直接对员工打总分。先判断：

个人是否缺少知识、判断或执行能力；

团队是否缺少复核、交接或备份；

系统是否缺少数据、规则、权限或稳定性。

第12—14天：确定投资组合

每项缺口从五种方式中选择：

在岗培养；

跨岗迁移；

外部招聘；

共享专家；

AI员工或系统能力。

14天结束时，董事会不需要看到复杂模型，而需要看到三张图：关键任务图、三层能力账本和优先投资组合。

12.3 第15—44天：30天实战认证

30天认证的目标是形成初步授权证据，而不是一次性给员工贴标签。

表54：30天认证周节奏

表54：30天认证周节奏

周次	重点	交付
第1周	任务说明、基线与受控练习	任务卡、初始样本
第2周	独立执行与过程记录	中间产出、修改日志
第3周	异常注入与跨角色协作	异常处置、交接证据
第4周	真实交付、复盘与答辩	最终样本、评分与授权建议

异常注入必须与真实岗位相关，不应设计成与工作无关的智力游戏。例如，投放岗位可注入归因口径变化、预算突增或数据中断；内容岗位可注入事实冲突、品牌规则变化或敏感声明；CRM岗位可注入授权状态不明和投诉增加。

表55：认证结果的五类处置

表55：认证结果的五类处置

结果	处置
证据充分、风险匹配	扩大或确认授权
能力基本达到但稳定性不足	限定授权，继续项目训练
个人能力足够但系统条件不足	优先修复系统，不降低个人评价
能力邻接但存在关键缺口	安排专项训练或共享专家
与目标岗位距离过大	调整岗位或采用其他供给方式

12.4 第45—90天：团队重组

团队重组不是按认证结果简单淘汰人员，而是重新组合人、AI员工、专家和系统。

表56：90天团队重组的六项工作

表56：90天团队重组的六项工作	
工作	关键产出
重写岗位卡	任务、能力、判断权和证据
调整项目授权	建议权、制作权、批准权、停止权
配置能力组合	人员、共享专家、AI员工和外部服务
重构关键流程	审批、交接、日志、异常与恢复
更新绩效机制	结果、质量、责任和系统贡献
建立复证机制	周期、触发条件、申诉和档案

90天结束时，应完成至少一个真实业务闭环的组织化运行，而不是只完成评测。

12.5 90天后的扩展

扩展顺序应是：

从一个任务扩展到相邻任务；

从一个团队扩展到同类团队；

从低中风险扩展到更高风险；

从局部规则扩展到组织级岗位体系；

从人工复核扩展到可审计的自动化治理。

不能因为试点成功，就无差别复制到所有岗位。每次扩展都要重新评估数据、顾客、品牌、风险和专业条件。

BOARD BRIEF 15

第十三章 董事会、CMO与CHRO管理工具

竹势 AI 营销智库出品

董事会十问

CMO 季度检查

CHRO 治理清单

联合责任表

投资组合画布

13.1 董事会十问

表57： 董事会AI营销人才十问

表57： 董事会AI营销人才十问

序号	董事会问题
1	我们正在改善哪些真实经营任务，而不是学习哪些工具？
2	哪些能力是全员基线，哪些只需要专家拥有？
3	哪些判断绝不能下放给系统？
4	哪些失败来自个人，哪些来自团队与系统？
5	高风险任务由谁拥有停止权？
6	工作样本如何进入招聘、授权和晋升？
7	我们是否把使用率和节省工时误写成经营价值？
8	哪些能力存在单点风险和继任缺口？
9	员工能力数据是否透明、最小必要、可申诉？
10	技术变化后，哪些能力仍会长期保值？

董事会不需要审核所有能力条目，但必须决定能力建设的经营目标、风险边界和资源原则。

13.2 CMO能力组合清单

表58：CMO季度检查清单

表58：CMO季度检查清单		
检查项	是/否	证据
核心营销任务是否已转成可评测任务		
	任务卡	
十岗是否有差异化能力要求		
	岗位卡	
重大品牌、预算和客户判断是否有责任人		
	授权矩阵	
内容和创意是否有事实与品牌质量门		
	评审量规	
投放是否区分平台归因与真实增量		
	实验报告	
CRM是否核验授权和客户影响		
	同意与投诉记录	
社媒是否建立人工接管与停止机制		
	危机流程	
是否保留失败、异常和纠偏证据		
	复盘库	
团队是否存在关键能力单点		
	能力账本	
AI员工是否有权限、日志与退出方案		
	治理档案	

13.3 CHRO人才机制清单

表59：CHRO能力治理清单

表59：CHRO能力治理清单

领域	检查问题
招聘	是否使用结构化工作样本，而非只看证书
培养	是否嵌入真实项目和反馈
绩效	是否避免使用率、数量和工时单一指标
晋升	是否要求复杂任务、异常和带动证据
薪酬	是否把溢价放在稳定稀缺责任上
人才盘点	是否考虑能力邻接、系统依赖和关键性
公平	是否控制机会、资源和评委偏差
隐私	是否遵守最小必要与保存期限
申诉	是否存在独立人工复核
退出	是否能删除过期标签并调整错误结论

13.4 CMO—CHRO联合责任矩阵

表60：能力体系联合责任

表60：能力体系联合责任

事项	CMO	CHRO	CIO/CDO	法务与安全	业务负责人
关键任务定义	A	C	C	C	R
岗位能力设计	A/R	A/R	C	C	C
工作样本设计	R	R	C	C	C
数据与工具条件	C	C	A/R	C	C
授权与审批	A/R	C	C	C	A/R
公平与隐私	C	A/R	C	A/R	C
绩效与晋升	A/R	A/R	C	C	C
系统事故复盘	A	C	A/R	A/R	R

A表示最终负责，R表示执行负责，C表示协同参与。实际企业可按治理结构调整，但不得让任何重大事项处于无人负责状态。

13.5 能力投资组合画布

表61：能力投资组合决策画布

表61：能力投资组合决策画布

字段	决策内容
目标任务	需要改善的真实经营任务
业务价值	收入、利润、客户、品牌、速度或风险
当前能力	个人、团队和系统现状
缺口	哪些能力或条件不足
邻接距离	现有人才到目标能力的距离
风险	错误影响、可逆性与监管要求
频率	高频持续或低频专项
窗口	需要多快形成能力
供给方式	培养、迁移、招聘、专家或AI员工
退出条件	何时停止、切换或重新评估

BOARD BRIEF 16

第十四章 结论：人才竞争的核心是判断责任与组织能力

竹势 AI 营销智库出品

人才竞争的核心，是判断责任与组织能力。

让任务有人负责、证据可以审阅、系统受到约束、员工获得公平发展。

企业不会因为员工使用更多人工智能产品，就自动形成AI营销能力。真正的组织能力来自五个连接：

经营目标与任务连接；

任务与能力连接；

能力与工作证据连接；

证据与授权连接；

授权与经营责任连接。

AI营销人才不是一种孤立的新职业，也不是所有营销岗位外加同一组工具技能。它是一种新的工作组织方式：专业人员把模糊目标转化为任务，系统扩展研究、生成和执行能力，人类保留关键判断，团队形成互补，组织提供数据、流程和治理条件，最终由明确责任人对结果负责。

企业应警惕三种幻觉。

第一种是工具幻觉：认为采购先进产品、完成全员培训，就已经具备能力。

第二种是个人英雄幻觉：认为少数高水平个人可以弥补数据、流程、权限和组织设计问题。

第三种是全面自动化幻觉：认为更多自主执行天然意味着更高生产力。

成熟企业采取相反路径：从任务出发，以证据评测，以风险决定授权，以三层账本进行归因，以团队组合形成能力，以治理保护客户、员工与企业。

最终，AI营销人才模型应帮助经营层作出五个更好的决定：

招谁；

培养谁；

把什么任务交给谁；

哪些工作交给系统；

谁对最终结果承担判断责任。

当这五个决定能够被真实工作证据支持，能力模型才不再是一份人力资源文件，而成为企业的营销生产系统。

附录一 经典思想的当代转译与适用边界

竹势 AI 营销智库出品

本报告没有把经典思想作为装饰性引语，而是将其转化为能力设计原则。经典理论形成于特定时代，其价值在于提供稳定的思考结构，而不是直接给出人工智能时代的操作答案。

一、Peter Drucker：知识工作与成果责任

原始思想背景： 德鲁克在知识工作和管理研究中强调，知识工作者不能仅以活动量衡量，管理应关注贡献与成果；知识工作者也需要明确“我能作出什么贡献”。

当代转译：

AI营销人才不能按操作次数和生成数量定义；

专业人员必须把知识、判断转化为经营贡献；

员工需要知道其任务服务于何种客户和组织结果；

管理者要为知识工作提供目标、信息与责任边界。

适用边界： 营销结果具有多因素性，不能把所有商业结果直接归因于单个知识工作者。德鲁克的成果导向不等于简单结果主义，更不意味着忽视过程质量和系统条件。

二、David McClelland与Richard Boyatzis：以胜任特征替代单纯智力与资历判断

原始思想背景： McClelland提出应测试与工作表现相关的胜任特征，而非过度依赖智力测验和学历；Boyatzis进一步发展了管理胜任模型，关注导致有效表现的个人特征与行为。

当代转译：

AI营销能力必须对应真实工作表现；

工具证书和自评不能替代工作样本；

能力项必须转成可观察行为；

评测应关注在特定情境下的有效表现。

适用边界：胜任模型容易被写成抽象行为词典，也可能忽略组织系统。现代转译必须加入个人—团队—系统分账。

三、John Flanagan：关键事件法

原始思想背景：Flanagan通过关键事件法识别导致成功或失败的关键行为，用于岗位分析和绩效研究。

当代转译：

从真实项目中的高成功、高失败和高风险事件提取能力；

异常处置比常规流程更能区分等级；

能力设计应保留事件背景、行为、结果和条件；

失败案例应进入晋升和认证证据。

适用边界：关键事件依赖参与者回忆，可能存在选择性叙述。企业需要结合日志、产出和多方证据，不应把单个故事当作普遍规律。

四、Herbert Simon：有限理性

原始思想背景：Simon指出决策者受到信息、时间和认知能力限制，现实决策往往追求足够满意，而非理论最优。

当代转译：

AI可以扩展信息处理，但不能消除信息质量和目标冲突；

任务必须明确何时继续研究、何时行动；

人机系统需要控制候选过载；

高水平人才应在不完全信息下作出透明、可修正的判断。

适用边界：有限理性不能成为低质量决策的借口。组织仍应改善数据、流程和评审机制。

五、W. Edwards Deming：系统观与质量责任

原始思想背景：Deming强调大多数绩效问题来自系统，管理者需要改善流程、反馈和质量机制，而非只责怪个人。

当代转译：

建立个人—团队—系统三层能力账本；

数据、流程、权限和工具缺陷不能全部归因于员工；

质量应嵌入流程，而不是依赖末端检查；

管理者对系统条件承担责任。

适用边界： 系统观不意味着个人无需负责。故意违规、疏忽和拒绝学习仍需被识别，关键是正确分账。

六、Chris Argyris：双环学习

原始思想背景： Argyris区分修正行动的单环学习，与挑战目标、规则和假设的双环学习。

当代转译：

低等级人员按流程修正错误；

高等级人员能够质疑任务定义、指标和工作流；

复盘不仅问“如何做得更快”，还问“是否做了正确的事”；

能力模型本身也需要根据实践修改。

适用边界： 持续挑战规则可能降低执行稳定性。双环学习应在明确责任和证据基础上进行。

七、Dreyfus兄弟：技能获得阶段

原始思想背景： Dreyfus模型描述从新手依赖规则，到专家在复杂情境中形成整体判断的发展过程。

当代转译：

五级证据阶梯按独立性、复杂度和异常能力区分；

新手需要规则与模板；

高水平人员能够处理规则失效和情境变化；

年限不能自动代表阶段。

适用边界： 专家直觉可能受到偏见影响。人工智能时代的专家仍需保留证据、复核和可解释性。

八、Karl Popper：可证伪性

原始思想背景： Popper强调科学命题应允许被证据反驳。

当代转译：

营销假设必须说明什么结果会证明其错误；

研究和投放不能只收集支持原结论的证据；

AI生成的解释必须与竞争性解释比较；

复盘应保留失败结果。

适用边界： 品牌、创意和复杂社会行为不总能形成严格可证伪命题，但仍可提高假设透明度和证据纪律。

附录二 具名专家视角及适用边界

竹势 AI 营销智库出品

1. Peter Drucker：成果而非活动

原始上下文：知识工作者的有效性与贡献。

本报告转译：使用量、内容量和工时不是最终能力证据。

边界：不能用单一短期结果替代专业过程和长期价值。

2. David McClelland：测试胜任而非只测试智力

原始上下文：人员甄选应关注与工作表现相关的特征。

本报告转译：使用工作样本、行为证据和真实任务。

边界：胜任特征必须岗位化，不能泛化为人格标签。

3. Richard Boyatzis：有效表现与情境匹配

原始上下文：胜任力与岗位、组织环境共同决定有效表现。

本报告转译：个人能力不能脱离团队和系统条件。

边界：不应把模型复杂化到无法应用。

4. John Flanagan：关键事件揭示有效行为

原始上下文：从成功和失败事件中提取岗位要求。

本报告转译：认证和晋升必须包含异常与失败证据。

边界：需防回忆偏差和只选择戏剧性事件。

5. Herbert Simon：有限理性与决策边界

原始上下文：决策者在有限信息和能力下行动。

本报告转译：AI扩大处理能力，但仍需任务边界和停止条件。

边界：不能因不确定性取消责任。

6. W. Edwards Deming：系统决定大部分表现条件

原始上下文：质量管理与系统改进。

本报告转译：用三层账本区分个人、团队和系统责任。

边界：不排除个人故意违规和持续不胜任。

7. Chris Argyris：双环学习

原始上下文：组织学习应挑战底层假设。

本报告转译：高等级人才应能重构错误任务和指标。

边界：挑战必须基于证据和责任，不是无边界质疑。

8. Hubert与Stuart Dreyfus：从规则依赖到情境判断

原始上下文：技能习得阶段。

本报告转译：能力等级按独立性、异常与影响范围划分。

边界：专家判断仍需校准，不能神化直觉。

9. Anders Ericsson：刻意练习

原始上下文：专业表现需要针对弱点、即时反馈和持续练习。

本报告转译：培养必须嵌入任务、反馈和复证。

边界：营销复杂任务不等同于结构稳定的竞技领域，练习效果受任务质量和业务机会影响。

10. Amy Edmondson：心理安全

原始上下文：团队成员能够报告错误、提出异议和承认不确定性。

本报告转译：员工应能停止自动化、报告缺陷和挑战错误结论。

边界：心理安全不等于降低绩效标准或回避责任。

11. Daniel Kahneman：认知偏差与噪声

原始上下文：人类判断存在系统偏差和不必要差异。

本报告转译：使用结构化量规、多评委和校准减少评审噪声。

边界： 标准化也可能压制创意和情境判断，需要保留专业答辩。

12. Erik Brynjolfsson：技术价值取决于互补性组织创新

原始上下文： 数字技术和人工智能的生产率价值通常需要流程、技能和组织资本配合。

本报告转译： AI能力不能脱离流程重构、数据和团队设计。

边界： 宏观生产率和企业局部项目之间不能直接等同。

附录三 前述案例的证据性质与限制总表

竹势 AI 营销智库出品

表62：案例证据性质与使用边界

表62：案例证据性质与使用边界

案例	证据性质	可支持的判断	不可支持的判断
可口可乐Create Real Magic	企业官方材料、公开活动信息	受控开放创意与品牌治理	不能证明长期销售或品牌资产增长
HBS/P&G Cybernetic Teammate	企业现场实验、学术研究	AI可改变个人与团队任务表现	不能外推至所有岗位和长期绩效
中国平安智能化实践	年报与企业官方披露	大规模应用需要组织治理	不能把覆盖率等同人才能力
Meta自动化广告	平台官方产品与案例	投放操作正趋向自动化	供应商自报不能证明普遍增量
微软Tay	企业回应与广泛公开记录	开放自动互动会放大治理缺陷	不能证明所有社媒自动化不可行
星巴克Deep Brew	企业官方披露与管理层材料	个性化依赖数据与运营体系	不能单独归因财务增长
浙江大学、阿里巴巴等人才实践	高校、企业公开材料	产业项目有助于复合能力培养	课程与竞赛不能证明岗位胜任
IBM Watson for Oncology	企业材料、研究与调查报道	高风险场景需独立验证和专业责任	不能直接类比所有营销任务

案例的用途是解释机制和管理边界，而不是提供企业照抄模板。供应商或企业自报数字必须与独立证据区别标注。

附录四 管理工具模板

竹势 AI 营销智库出品

工具一：能力证据卡

表63：能力证据卡

表63：能力证据卡

字段	填写内容
人员与岗位	

目标能力

任务背景

本人角色

输入条件

关键判断

最终输出

异常与处置

结果与限制

团队贡献

系统条件

建议等级

授权建议

复证日期

工具二：项目授权卡

表64：项目授权卡

表64：项目授权卡	
字段	填写内容
项目/任务	
风险等级	
建议权持有人	
制作权持有人	
批准权持有人	
停止权持有人	
系统执行范围	
强制人工节点	
异常升级路径	
日志要求	
授权有效期	

工具三：团队能力账本

表65：团队能力账本

表65：团队能力账本						
关键任务	主要人员	备份人员	AI/系统能力	共享专家	关键缺口	风险

工具四：实战认证评审表

表66：实战认证评审表

表66：实战认证评审表			
评审维度	评分	证据	评委说明
问题定义			
任务建模			
专业判断			
研究与证据			
知识与数据			
系统编排			
判断与质量			
执行与协同			
责任与治理			
异常处理			
综合授权建议			

公开来源索引

竹势 AI 营销智库出品

以下索引仅列机构、作者、标题、日期与主要使用口径，不含外部链接。

OECD, 《OECD Skills Outlook 2023: Skills for a Resilient Green and Digital Transition》, 2023年；数字转型、技能与任务适应框架。

OECD, 《OECD Employment Outlook 2023: Artificial Intelligence and the Labour Market》, 2023年；人工智能、工作任务与劳动市场影响。

World Economic Forum, 《Future of Jobs Report 2025》, 2025年；未来技能、岗位变化与企业人才策略。

International Labour Organization, 《Generative AI and Jobs: A Global Analysis of Potential Effects on Job Quantity and Quality》, 2023年；任务暴露、增强与替代边界。

European Commission, ESCO分类体系；职业、技能、能力和资格的结构化定义。

U.S. Department of Labor, O*NET Resource Center；岗位任务、工作活动、知识与技能分析体系。

UNESCO, 《Guidance for Generative AI in Education and Research》, 2023年；人工智能素养、责任与教育使用边界。

NIST, 《Artificial Intelligence Risk Management Framework 1.0》, 2023年；治理、测量、管理和可信人工智能框架。

NIST, 《Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile》, 2024年；生成式系统风险与控制。

ISO/IEC, 《ISO/IEC 42001:2023 Information technology—Artificial intelligence—Management system》, 2023年；人工智能管理体系。

全国人民代表大会常务委员会, 《中华人民共和国个人信息保护法》, 2021年；员工和客户个人信息处理原则。

全国人民代表大会常务委员会, 《中华人民共和国数据安全法》, 2021年；数据分类、处理与安全责任。

国家互联网信息办公室等, 《生成式人工智能服务管理暂行办法》, 2023年；生成式服务治理要求。

国家互联网信息办公室, 《个人信息保护合规审计管理办法》, 2025年；个人信息处理合规审计。

David C. McClelland, 《Testing for Competence Rather Than for Intelligence》, American Psychologist, 1973年；胜任特征与人员甄选。

Richard E. Boyatzis, 《The Competent Manager: A Model for Effective Performance》, 1982年; 胜任模型与有效表现。

John C. Flanagan, 《The Critical Incident Technique》, Psychological Bulletin, 1954年; 关键事件法。

Benjamin Bloom等, 《Taxonomy of Educational Objectives》, 1956年; 学习目标层级。

Lorin Anderson、David Krathwohl等, 《A Taxonomy for Learning, Teaching, and Assessing》, 2001年; 修订版学习分类。

Stuart Dreyfus、Hubert Dreyfus, 《A Five-Stage Model of the Mental Activities Involved in Directed Skill Acquisition》, 1980年; 技能获得阶段。

Anders Ericsson、Ralf Krampe、Clemens Tesch-Römer, 《The Role of Deliberate Practice in the Acquisition of Expert Performance》, Psychological Review, 1993年; 刻意练习与专业表现。

Frank Schmidt、John Hunter, 《The Validity and Utility of Selection Methods in Personnel Psychology》, Psychological Bulletin, 1998年; 人员甄选方法效度。

Levashina、Hartwell、Morgeson、Campion, 《The Structured Employment Interview: Narrative and Quantitative Review》, Personnel Psychology, 2014年; 结构化面试。

Peter Drucker, 《The Effective Executive》, 1967年; 知识工作者、贡献与有效性。

Herbert A. Simon, 《Administrative Behavior》, 1947年及后续版本; 有限理性与组织决策。

W. Edwards Deming, 《Out of the Crisis》, 1982年; 系统观、质量和管理责任。

Chris Argyris、Donald Schön, 《Organizational Learning》, 1978年; 单环与双环学习。

David Kolb, 《Experiential Learning》, 1984年; 经验学习循环。

Karl Popper, 《The Logic of Scientific Discovery》, 1959年英文版; 可证伪性与证据。

Joseph Juran, 《Juran on Quality by Design》, 1992年; 质量策划与管理。

Amy Edmondson, 《Psychological Safety and Learning Behavior in Work Teams》, Administrative Science Quarterly, 1999年; 心理安全与团队学习。

Daniel Kahneman、Olivier Sibony、Cass Sunstein, 《Noise: A Flaw in Human Judgment》, 2021年; 判断噪声与结构化评审。

Erik Brynjolfsson、Daniel Rock、Chad Syverson, 《Artificial Intelligence and the Modern Productivity Paradox》, 2017年; 技术、互补资产与生产率。

Harvard Business School研究团队,《The Cybernetic Teammate: A Field Experiment on Generative AI Reshaping Teamwork and Expertise》, 2025年前后公开版本; 791名宝洁专业人员现场实验。

The Coca-Cola Company,《Create Real Magic》相关官方发布, 2023年; 受控品牌资产与公众生成创意。

The Coca-Cola Company, 年度报告及营销创新相关官方材料, 2023—2025年; 品牌与人工智能应用背景。

Meta,《Advantage+》系列官方产品材料与广告主案例, 2022—2026年; 自动化投放、供应商自报效果及适用边界。

Microsoft, Tay事件官方回应, 2016年; 公开社交机器人失控与停止处置。

Starbucks, Deep Brew相关投资者材料、年度报告与官方技术说明, 2019—2025年; 个性化与运营智能。

IBM, Watson Health及Watson for Oncology相关企业材料, 2010年代; 智能决策支持背景。

STAT等专业媒体及公开研究, Watson for Oncology调查与评价, 2017—2022年; 风险、限制与反噬案例。

中国平安, 年度报告, 2023—2025年; 金融、保险、健康场景的智能化和组织能力。

中国平安, 投资者与企业官网人工智能相关材料, 2023—2026年; 企业自报应用范围。

浙江大学, 人工智能通识教育、交叉人才培养与产教融合公开材料, 2023—2026年; 高校人才培养。

阿里巴巴集团、阿里云, 人工智能人才、开发者和产业实践公开材料, 2023—2026年; 产业能力与项目实践。

Belbin,《Management Teams: Why They Succeed or Fail》, 1981年; 团队角色, 谨慎作为团队讨论工具, 不作为人格定论。

Melvin Conway,《How Do Committees Invent?》, 1968年; 组织沟通结构与系统设计关系。

Eric Trist、Fred Emery等, 社会技术系统相关研究; 人、技术与组织共同设计。

Parasuraman、Sheridan、Wickens,《A Model for Types and Levels of Human Interaction with Automation》, IEEE Transactions on Systems, Man, and Cybernetics, 2000年; 自动化层级与人类作用。

Theodore Levitt,《Marketing Myopia》, Harvard Business Review, 1960年; 顾客问题与产品视角边界。

Les Binet、Peter Field,《The Long and the Short of It》, 2013年; 品牌建设与短期激活平衡。

Byron Sharp,《How Brands Grow》, 2010年; 品牌增长、触达与记忆结构, 适用时需考虑品类和市场边界。

Philip Kotler、Kevin Lane Keller,《Marketing Management》相关版本; 营销管理、顾客价值与市场系统。

中国信息通信研究院, 人工智能治理、可信人工智能与产业人才相关报告, 2023—2026年; 中国企业应用与治理语境。

工业和信息化部等,《中小企业数字化赋能专项行动方案(2025—2027年)》, 2024年; 中小企业数字化和人工智能应用。

国务院,《关于深入实施“人工智能+”行动的意见》, 2025年; 人工智能融入战略、组织和业务流程。

PUBLISHER

关于竹势 AI 营销智库

竹势 AI 营销智库出品

竹势 AI 营销智库聚焦 AI 市场部与 AI 营销，面向企业创始人、CEO、CMO及市场增长团队，构建“6+1”知识架构。六大核心模块组成“道、法、术、器、技、例”六脉体系：“道”负责认知刷新，研究 AI 时代的营销本质与战略判断；“法”关注体系建设，沉淀可持续、可复制的方法论；“术”提供落地路径，把策略转化为具体行动；“器”评测工具与平台，帮助企业选择合适的能力组合；“技”分享可以立即应用的实战技巧；“例”通过标杆案例验证方法与成效。“+1”专题围绕 AI 营销的重要趋势与关键经营问题，推出系统、深入、可下载的专题白皮书。

竹势智库通过专业研究、实践经验与管理框架连接认知、决策和执行，帮助企业看清方向、减少试错，逐步建设能够创造真实经营价值的 AI 市场部。