



017 / NETWORK

AI MARKETING / CHANGE ADOPTION
NETWORK

AI营销变革管理白皮书

真正的采用发生在工作现场：领导联盟、关键行为、流程接口、同伴信任、责任边界和反馈必须成为同一套可观测系统。

识别阻力的真实来源

区分理性、身份、能力、制度与风险五类断点

让采用证据逐层升级

接触、任务、质量、客户与经营结果分层审议

重画关键行为与工作

从经营结果倒推岗位动作、接口和责任

形成可持续强化回路

90 / 180 / 365 天推进、调整、回退或停止

出版说明

竹势 AI 营销智库出品

本报告讨论企业在采购人工智能工具、完成通识培训和启动试点之后，如何改变关键岗位的日常行为、端到端流程与经营结果。报告把变革对象界定为工作系统、权力关系、风险分配和强化机制，而非把采用困难简单归因于员工态度。

报告中的企业案例依据公开披露、监管文件、同行评审或严谨现场实验整理。企业自报、供应商案例和独立研究在文中明确区分；未公开的数据不作推断。所有数字均按来源可核验口径呈现。

版权归竹势 AI 营销智库所有。引用本报告时请注明报告名称、出品方与版本时间。

阅读导航

部分	核心问题	建议读者
执行摘要	董事会应形成什么共同判断与决议	全体经营层
第一至三章	失败为何被误诊；领导联盟与利益契约如何建立	董事长、CEO、CHRO
第四至八章	流程、能力、网络、权责与激励如何协同	CMO、CIO/CDO、一线经理
第九至十一章	采用证据、扩张闸门与治理如何运行	变革办公室、营销运营、风控合规
管理者工具	直接用于会议、诊断和阶段决议的表单	项目负责人及业务负责人

BOARD BRIEF 02

执行摘要

竹势 AI 营销智库出品

CHANGE ADOPTION RULE / 变革采用原则

AI 变革只有同时改变关键行为、工作接口、权责和强化机制，才会从热情试用变成稳定经营能力。

先锁经营结果

再定义关键行为

重画工作系统

按证据强化

AI营销变革的首要风险，不是员工明确反对，而是组织在表面采用之下维持旧目标、旧流程、旧审批、旧绩效和旧责任。工具进入了岗位，工作系统却没有改变。员工于是形成理性选择：在低风险、低价值的边缘任务上使用AI，在影响客户、收入、预算和职业声誉的关键任务上继续沿用旧方法。登录数和内容产量因此上升，经营结果没有变化。

董事会需要把AI采用视为一套可运营的“采用系统”。该系统从经营结果倒推关键行为，再配置工作流程、角色权责、同伴网络、能力练习、质量阈值、风险接管、激励与反馈。任何一环缺失，采用都会停留在个人技巧或短期活动层面。

十项可证伪核心结论

竹势 AI 营销智库出品

编号	结论	可证伪条件
1	AI变革的主要对象是工作系统。	若岗位目标、输入数据、审批权、质量标准、例外处理和绩效奖励均未改变，则培训与工具普及不会稳定改变关键任务行为。
2	单一高层赞助不足以跨越组织边界。	若赞助人不能调动业务、技术、人力与合规资源，不能在约定时间内解决跨部门冲突，项目会在“都支持、无人决策”中停滞。
3	员工按“工作损益表”而非宣传口号判断采用。	当节省时间归组织、错误责任归个人、优秀产出又削弱岗位稀缺性时，员工的最优策略是低风险试用而非关键任务迁移。
4	把AI叠加进旧流程通常先增加工作量。	双轨录入、重复审核、人工复制粘贴和并行留档会吞噬效率收益；流程没有删除步骤，工具只会制造新的等待和返工。
5	通识培训不能自动迁移为岗位能力。	能力形成依赖真实任务、可比较样例、即时反馈、经理教练、沙盘和质量阈值。培训完成率与关键任务胜任度之间不存在必然关系。
6	采用沿可信关系扩散，而非沿组织图扩散。	一线员工更愿意复制可信同伴在相似任务中的可见做法。只有“冠军”而没有岗位制度与接班机制，会形成脆弱个人孤岛。
7	责任与决策权不匹配会压制真实试错。	员工被要求承担AI错误后果，却无权修改流程、拒绝低质输出或请求人工接管时，心理安全 and 高质量反馈都会消失。
8	探索期与规模期必须采用不同绩效逻辑。	探索期需要奖励可验证学习和问题暴露；规模期才适合奖励稳定采用、质量、客户结果与单位经济性。过早绑定奖金会诱发刷量和隐瞒例外。
9	采用证据必须穿透到关键任务和客户行为。	接触、活跃、内容产量只证明工具活动；只有关键任务渗透、质量与例外、客户行为和经营结果的连续证据才能支持扩张。
10	规模化是一组阶段决议，不是一次上线。	90、180、365天必须设置推进、调整、回退和停止闸门。没有停止条件的项目会把沉没成本伪装成战略耐心。

董事会应在首次决议中明确的六件事

确定一个可归因的经营结果与2—4个关键岗位行为，避免以“全员使用AI”作为目标。

任命联合赞助人：业务结果负责人、技术与数据负责人、组织与人才负责人、风险与合规负责人；明确各自可调动资源与升级时限。

批准流程重构权限，允许项目组删除旧步骤、调整审核点、改变角色分工，而不仅是配置工具。

批准探索期的安全边界：任务分级、数据边界、人工接管、留痕、回退、客户告知和事故升级。

批准采用证据阶梯与基线，明确哪些指标支持扩张，哪些信号触发回退或停止。

要求业务负责人而非技术团队对关键行为迁移和经营结果负责。

第一章 误诊：为什么“员工不接受工具”常常是错误结论

竹势 AI 营销智库出品

1.1 采用困难通常是系统性理性反应

管理层最容易看到的是员工没有登录、没有提交AI产出或仍保留手工做法，于是把原因解释为保守、恐惧或学习意愿不足。这个解释把组织设计问题个人化。员工面对的是一组具体约束：绩效仍按旧产量计算；经理仍要求旧格式；AI输出需要额外复核；错误责任没有重新分配；客户投诉会直接落到个人；数据访问受限；优秀员工担心经验被系统吸收后失去稀缺性。此时不迁移关键任务，往往是保护绩效与职业风险的理性选择。

Lewin在《Frontiers in Group Dynamics》中提出，行为是个体与环境共同作用的结果。把这一场论连接到AI变革，管理含义是先改变约束力与驱动力的结构，再讨论态度。边界在于：Lewin的“解冻—改变—再冻结”来自相对稳定的组织环境，而AI能力和监管持续变化，今天更适合把“再冻结”理解为可审计、可更新的运行标准，而非永久固定。[S1]

1.2 变革采用网络：五层因果链

层级	要回答的问题	可观察证据	常见泄漏
经营结果	收入、利润、客户价值、现金流或风险结果为何应改变？	可归因指标、对照或基线、时间滞后	目标过宽或不可归因
关键行为	哪些岗位在什么触发下必须改变什么动作？	任务渗透、行为日志、审核记录	只统计登录和产量
工作系统	流程、数据、权限、质量与例外如何支持新行为？	流程时长、返工、等待、接管	旧流程与新工具双轨
关系网络	员工向谁学习、向谁求助、谁有可信度？	咨询网络、桥接者、同伴示范	冠军孤岛、经理缺席
强化机制	绩效、认可、制度和反馈如何让行为持续？	绩效口径、例会、预算、制度文本	奖励旧行为、问题被隐藏

该模型的关键不是层数，而是反向设计顺序：先定义经营结果，再确定关键行为；再为行为配置工作系统、关系网络和强化机制。若项目从工具功能开始，往往只能产生使用活动，无法证明经营因果。

第二章 领导联盟：赞助不是表态，而是跨边界决策能力

竹势 AI 营销智库出品

CEO 方向与取舍 经营结果、资源、冲突裁决
CMO / 业务 价值流 owner 关键任务、客户与结果
CHRO 岗位与强化 能力、经理、激励、公平
CIO / 治理 可靠与边界 数据、权限、日志、降级
赞助不是出席启动会；每个闸门都需要明确决策权、资源承诺和升级时限。

Kotter在《Leading Change》中强调指导联盟。AI营销项目比传统系统上线更需要联合赞助，因为它同时改变业务目标、岗位边界、数据权限、客户触点、内容责任和风险暴露。单一CMO通常无法调整身份权限，CIO不能决定营销流程，CHRO不能单独改变客户风险规则，法务也不能替业务定义可接受质量。[S2]

2.1 联合赞助人的四项真实权力

赞助角色	必须拥有的权力	不可委托的决策	建议升级SLA
业务赞助人	目标、预算、流程与岗位优先级	关键行为、客户价值、资源取舍	重大业务冲突48小时
技术/数据赞助人	架构、数据、身份权限与集成资源	数据边界、可用性、可观测性	关键依赖5个工作日
组织/人才赞助人	岗位、绩效、学习与劳资沟通	角色变化、能力认证、激励原则	岗位与绩效争议10个工作日
风险/合规赞助人	风险分级、控制、审计与事故处置	禁止场景、人工控制点、回退条件	高风险事项24小时

联盟有效性的检验不是会议出席率，而是跨边界问题被解决的速度、不可逆决策是否有人承担、资源是否随优先级移动。所谓“高层支持”若没有决策权与升级时限，只会把冲突推迟到一线。

2.2 领导联盟的反例：人人支持、无人承担

许多试点在演示阶段顺利，因为不触碰真实客户、真实预算和真实数据。一旦进入生产环境，业务要求速度，技术要求稳定，合规要求证明，员工要求责任边界。若联合赞助机制未提前定义，项目组会用更多会议替代决策，最终将失败解释为“组织协同不足”。管理层应把协同具体化为决策清单、授权额度和升级SLA。

第三章 意义与利益：员工用“工作损益表”判断AI

竹势 AI 营销智库出品

岗位收益 更少等待 · 更好判断 · 更强客户响应
岗位成本 学习 · 复核 · 身份变化 · 暴露风险
保护条件 可逆 · 公平 · 清晰责任 · 真实支持
员工接受度来自可感知的净工作收益，而非统一宣讲后的态度表态。

员工的采用意愿由可见收益、身份影响、控制感、公平感和风险共同决定。前景理论说明，人们对潜在损失通常比同等收益更敏感；心理契约研究则提示，正式制度之外，员工还会根据组织是否兑现公平、发展与安全承诺调整投入。AI项目只宣传节省时间，却不说明节省出的时间如何使用、岗位如何发展、错误由谁承担，会放大损失预期。[S3][S4]

工作损益项	员工可能的判断	管理设计
时间收益	节省时间是否会转化为更高工作量？	明确释放时间的用途与负荷上限
绩效收益	AI产出是否被认可，还是仍按旧格式考核？	把关键行为与质量纳入绩效
身份与地位	经验是否被贬值，岗位是否被削弱？	定义更高价值的新职责与晋升路径
控制感	能否拒绝低质输出、请求接管、修改规则？	赋予停止、升级和反馈权
公平感	不同岗位是否承担不同风险却得到同样奖励？	按风险和贡献差异化认可
责任风险	客户、合规或品牌错误由谁承担？	明确执行、审核、批准与事故责任

3.1 高意愿不等于正确使用

员工可能积极使用AI，却集中在写摘要、改文案和查资料等低风险任务；这能产生显著主观节省时间，却未必触及影响商机、客户体验和预算效率的关键任务。微软与LinkedIn在2024年对31个国家约3.1万名知识工作者的研究显示，使用扩张快于组织战略和治理建设；该类调查能说明普及与感知，但不能直接证明收入或利润改善。[S5]

第四章 流程重构：把AI塞进旧流程，通常先制造返工

竹势 AI 营销智库出品

01

删除

无价值等待与重复

02

简化

输入、交接与标准

03

重画

人与 AI 的责任接口

04

嵌入

真实任务与状态回写

05

接管

例外、暂停与复盘

Hammer在1990年提出业务流程再造，核心是围绕结果重新设计流程，而非自动化既有步骤。AI变革的当代应用不是追求“大爆炸式重构”，而是对高价值流程逐步执行四个动作：删除不再需要的步骤，合并重复判断，前移质量控制，重新配置人工控制点。边界在于，经典流程再造常低估组织学习、员工参与和持续风险校准；AI流程必须保留可回退与可更新能力。[S6]

流程动作	判断问题	AI营销示例	风险控制
删除	该步骤是否只为弥补旧系统缺陷？	删除人工复制平台数据到周报	保留数据校验与抽样
简化	是否存在多层重复审核？	将三轮格式审核合并为规则校验+业务审核	重大内容保留法务门
重画	任务顺序是否仍合理？	先生成多版本并小样测试，再投入制作	测试样本与停止阈值
前移	错误能否在更早阶段发现？	在生成前加载品牌禁区与证据要求	输出后仍做事实核验
分级	所有任务是否需要同等人工控制？	低风险内部摘要自动，高风险对外承诺双审	任务分级与权限隔离

案例1：财富500强软件公司的客户支持AI——生产率提升集中在新人，高手收益有限

2020—2021年前后，一家财富500强企业向5,179名客户支持人员分阶段部署生成式AI助手。系统根据历史高质量对话向坐席提供实时建议，嵌入既有客户支持 workflow。

证据与结果：NBER研究报告显示，问题每小时解决量平均提高约14%；低技能和低经验员工增益更大，最高约35%，高绩效员工的增益很小，部分质量指标甚至轻微下降。研究还观察到客户情绪改善与人员流失下降。[S7]

归因与外推边界：这是单一企业、特定支持任务和特定部署方式的准实验结果。生产率口径是解决问题数量与质量的组合，不能外推到所有知识工作；工具可能通过复制优秀员工隐性知识产生收益，也可能使高手受到不合适建议干扰。

可复制条件：把AI嵌入任务时应按经验水平差异化配置。新人可获得更强建议和训练反馈，专家应拥有更高的忽略权、规则修订权和知识回流责任。

案例2：P&G“网络化队友”现场实验——AI改变团队边界，但自动替代协作设计

2024年前后，哈佛商学院与宝洁合作开展大规模现场实验，791名专业人员围绕真实产品创新任务被随机分配到个人、团队以及有无生成式AI的不同条件。

证据与结果：研究显示，使用AI的个人能够达到接近无AI团队的表现，AI也缩小了商业与技术职能之间的知识边界，并影响参与者的积极情绪。论文为2025年HBS工作论文。[S8]

归因与外推边界：实验任务具有明确时间和产出标准，不能直接证明长期组织绩效；参与者知道自己处于实验环境，且工具、任务和评价方式受到控制。它不能支持“取消团队”，更不能证明复杂跨职能治理可以被模型替代。

可复制条件：把AI作为团队成员需要重新定义信息共享、责任和评审方式。组织可让个人先借助AI形成跨职能初稿，再由真实团队处理利益冲突、资源承诺和不可逆决策。

第五章 能力与练习：通识培训为何难以迁移为岗位行为

竹势 AI 营销智库出品



Bandura的自我效能理论指出，人们对自身能否完成任务的判断会影响选择、投入和坚持。最强来源通常是亲身成功经验，其次是相似他人的示范、社会说服与情绪状态。Kolb的经验学习循环进一步强调，能力来自具体经验、反思、概念化和再实验。对AI变革而言，培训必须围绕真实任务和反馈闭环，而非只讲工具功能。[S9][S10]

能力形成环节	最低设计要求	证据	常见失败
真实任务	使用本岗位真实输入与质量标准	任务样本与产出对比	课堂案例过于理想化
即时反馈	在任务后快速指出事实、结构、风险与客户影响	审核意见、错误标签	只给“好/不好”
经理教练	经理能观察、示范、纠偏和分配任务	教练记录、任务复盘	经理本人不使用
安全沙盒	允许低成本试错且不影响真实客户	隔离环境、模拟数据	要么完全禁止，要么直接上线
质量阈值	定义何时可用、何时复核、何时回退	通过率、例外率	以主观满意代替标准
能力认证	用任务表现而非课时判断胜任	盲测、情景测试	以培训完成率代替能力

案例3：BCG顾问实验——能力边界决定收益与风险

2023年，研究者与波士顿咨询合作，让758名顾问完成18项接近真实咨询工作的任务，并随机分配不同AI使用条件。

证据与结果：在位于模型能力边界内的任务上，使用GPT-4的顾问完成更多任务、速度更快、质量更高；在模型能力边界外的一项任务上，AI组更容易给出错误答案。该研究形成“锯齿状技术前沿”概念。[S11]

归因与外推边界：任务持续时间有限，参与者是咨询专业人士，实验使用特定时期模型；结果不能直接外推到长期客户项目、团队政治或高监管决策。

可复制条件：培训不能只教“怎么用”，必须教员工识别能力边界、验证事实、在低置信度时升级，并对不同任务设定不同人工控制点。

第六章 同伴网络：采用沿可信关系扩散

竹势 AI 营销智库出品

可信桥接者 连接业务、技术与一线
先行小组 验证机制
经理层 重排日常工作
一线同伴 降低学习与身份风险
制度 owner 把经验写进系统
冠军的价值是缩短信任距离；经验若不能进入流程、标准和经理节奏，冠军会变成单点依赖。

Rogers在《Diffusion of Innovations》中说明，创新扩散受到相对优势、兼容性、复杂度、可试验性和可观察性影响，并通过社会系统中的意见领袖传播。AI采用的新增难点是输出质量具有任务依赖性，员工无法仅凭“别人用了”判断是否适合自己的高风险任务。因此，需要把可观察性具体化为相似岗位、相似客户、相似风险条件下的可复核案例。[S12]

网络角色	价值	机制	脆弱点
可信同伴	降低试用不确定性	展示相似任务的完整做法与错误	只展示成功案例
桥接者	连接业务、技术、合规与数据群体	翻译需求与约束	角色过载、离职断点
一线经理	把采用嵌入任务分配与复盘	观察、教练、资源协调	只催登录不改任务
领域专家	定义质量标准与例外	审核难例、更新知识	被当成无限审核资源
社区运营者	沉淀问题、样例与规则	例会、知识库、案例复盘	沦为宣传群

冠军机制只有在三项条件成立时才可扩张：冠军的做法被转化为岗位标准；冠军拥有明确时间与认可；至少有一名接班者和一个跨团队桥接关系。否则，组织会把制度缺口寄托在少数热心员工身上。

第七章 权责与心理安全：允许暴露问题，才能获得真实反馈

竹势 AI 营销智库出品

执行 允许 AI 完成的动作
审核 证据与质量责任
批准 风险与商业后果
接管 异常、暂停、回退
复盘 无责上报与系统修正

Edmondson将心理安全界定为团队成员相信提出问题、承认错误和寻求帮助不会遭受人际惩罚。AI项目尤其需要这一条件，因为模型错误、数据缺陷和流程例外只有一线人员最先看到。若管理层把问题上报视为抵触、把回退视为失败，员工会隐藏错误或绕开系统。心理安全并不等于降低绩效标准，而是把“说出风险”与“对质量负责”同时制度化。[S13]

任务风险级别	执行	审核	批准	人工接管	复盘
L1 内部低风险	AI可自动执行	抽样	无需逐项	异常时	周度
L2 可逆对外	人机协同	业务审核	岗位负责人	随时	日/周
L3 重要客户/预算	AI建议、人执行	双人或专业审核	业务负责人	默认可用	逐项
L4 法律/监管/重大品牌	仅辅助分析	专业审核	授权高管/法务	强制	逐项+事故机制

案例4：加拿大航空聊天机器人纠纷——责任不能外包给系统

2022年，一名客户依据加拿大航空网站聊天机器人提供的丧亲票价说明购票，随后申请退款被拒。2024年，不列颠哥伦比亚省民事解决法庭认定公司应对其网站聊天机器人提供的信息负责。

证据与结果：裁决要求公司赔偿差价及相关费用。案件表明，对外自动化触点的错误不会因为“由聊天机器人生成”而转移企业责任。[S14]

归因与外推边界：这是特定司法辖区的小额纠纷裁决，不等于所有国家对AI输出责任的统一规则；但它提供了清晰的治理信号：企业不能把自动化触点视为独立主体。

可复制条件：对客户承诺、价格、权益和政策解释，应建立版本化知识源、事实校验、人工接管和定期回归测试；责任矩阵必须指向企业岗位，而非系统。

第八章 激励与绩效：旧奖励会把新工具拉回旧行为

竹势 AI 营销智库出品

目标

关键行为进入经营节奏

可见反馈

质量、例外与客户结果

经理认可

奖励学习与协作

制度对齐

绩效、资源与晋升

下一轮决议

扩大、调整或停止

Goodhart定律常被概括为：当指标成为目标，它就不再是好指标。AI项目若把登录数、生成篇数或节省小时数直接绑定奖金，员工会优化可见活动，而不是客户价值和质量。探索期还会出现更严重的问题：员工为了保住奖励而减少问题上报，组织失去学习信号。[S15]

阶段	绩效重点	可奖励行为	不宜直接奖励
探索期 0—90天	学习速度与风险暴露	完成真实任务、记录例外、提出流程改进、帮助同伴	登录数、产量、未经验证的节省时长
验证期 90—180天	稳定关键任务采用与质量	达到阈值、降低返工、提高响应、减少高风险例外	单一采用率
规模期 180—365天	客户与经营结果、单位经济性	客户行为改善、转化、成本、风险下降、知识沉淀	只看短期收入
稳态期	持续改进与治理	更新规则、审计通过、能力覆盖、跨团队复制	零错误口号

绩效设计应把团队结果与个人责任结合：团队对流程结果负责，个人对执行、审核、升级和学习职责负责。完全以团队指标衡量会产生搭便车，完全以个人产量衡量会破坏协作和知识共享。

BOARD BRIEF 12

第九章 采用证据：五层阶梯穿透到经营结果

竹势 AI 营销智库出品

L1

接触

权限、培训、登录

L2

任务渗透

关键工作真实使用

L3

质量与例外

一次通过、返工、接管

L4

客户行为

响应、选择、信任

L5

经营结果

增量、成本、风险

低层信号证明接触，不证明采用；拨款与授权必须和证据层级匹配。

采用不是一个数字，而是一条证据链。不同频率的任务需要不同分母：高频任务可按周计算任务渗透；月度战役应按事件计算；低频重大任务应按案例和控制有效性评估。没有分母的“活跃用户”无法告诉管理层关键工作是否迁移。

层级	核心问题	领先指标	滞后指标	不充分替代指标
L1 接触与可用	人能否进入并完成最小操作？	权限开通、可用率、首次成功	稳定活跃	培训完成率
L2 关键任务渗透	目标任务中多少按新方式完成？	任务覆盖、岗位覆盖、经理分配	持续渗透率	登录数、生成量
L3 质量与例外	结果是否达标，风险是否可控？	一次通过、事实错误、接管、返工	投诉、事故、审计	主观满意
L4 客户与业务行为	客户是否响应不同？	打开、互动、响应、线索质量	转化、留存、客诉	内容产量
L5 经营结果	收入、利润、现金流或风险是否改善？	管道、周期、单位成本	收入、毛利、现金回收、损失	节省小时估算

9.1 指标树与基线规则

每个经营指标必须连接到可观察的关键行为，并说明预期滞后。

基线至少覆盖一个完整业务周期；季节性强的业务需要同比或匹配对照。

任务渗透率的分母是符合条件的目标任务，不是员工人数。

质量指标同时覆盖平均水平与尾部风险，避免均值掩盖严重错误。

例外率下降不一定代表系统更好，也可能代表员工停止上报；需结合审计抽样与心理安全信号。

经营结果没有改善时，应分解目标、行为、系统、网络和强化机制，而不是直接扩大培训。

第十章 中国企业案例：变革机制比“上线AI”更重要

竹势 AI 营销智库出品

案例5：海尔智家——数字营销与组织机制同步推进

海尔智家在2023—2025年年度报告中持续披露国内市场的数字营销、数字库存和用户直连转型，并以“人单合一”把员工价值与用户价值连接。公司将数字化延伸到经销商、零售触点、物流、服务与营销资源配置，而非仅提供单点工具。

证据与结果：2023年公司收入同比增长7.3%，归母净利润同比增长12.8%；2025年上半年披露新增100多家经销商，畅销型号占比提高6个百分点。公司同时强调数字库存和数字营销对运营效率与用户体验的作用。结果为公司综合经营结果，不能完全归因于AI或单一变革措施。[S16][S17]

归因与外推边界：公开披露没有给出员工关键任务采用率、对照组或独立因果识别。海尔的组织模式具有长期文化基础，其他企业不能只复制口号或平台。

可复制条件：可复制的是连接机制：把用户结果、经销商经营、岗位责任和数字流程放入同一目标体系；在扩张前验证经销商和员工是否减少等待、库存与返工。

案例6：平安银行——敏捷流程与数据驱动运营先于规模化智能应用

平安银行在2019年年报中披露推进全行开发运维一体化和安全开发生命周期项目，建立从需求到生产的端到端研发流程，并扩大科技人员与投入。该案例属于数字化与智能化基础建设，对后续AI采用具有直接启示。

证据与结果：公司披露相关项目使其可响应的业务开发需求比2018年增加约30%；零售人均营收同比提高17.7%。这些数字来自企业年报，且受到业务增长、组织投入和市场环境共同影响。[S18]

归因与外推边界：这是金融机构较早阶段的数字化转型案例，不是生成式AI营销项目；它不能证明增加科技投入必然带来经营回报。

可复制条件：AI营销项目需要端到端交付流程、安全开发、数据驱动运营和业务技术融合。若这些基础缺失，生成式AI会放大需求积压与权限冲突。

案例7：招商银行——以客户经营与组织能力承接金融科技

招商银行多年在年报中把金融科技投入、客户经营、流程线上化和员工能力建设作为同一经营体系推进。其零售金融强调客户分层、数字渠道、风险管理和服务流程协同。

证据与结果：公开年报披露了金融科技投入、数字客户与零售经营指标，但并未把具体经营变化单独归因于生成式AI。[S19]

归因与外推边界：银行业务受监管、风险偏好和客户结构影响显著；公开披露缺少岗位级采用证据。

可复制条件：高监管企业应先建立任务分级、数据权限、可解释审核和客户沟通规则，再扩张生成式AI；营销与服务采用不能脱离风险治理。

案例8：百度内部智能代码助手——高频任务采用依赖工作流嵌入与反馈

百度在官方披露和开发者大会中介绍智能代码助手在内部研发中的规模化使用，强调代码补全、解释、测试与研发流程结合。

证据与结果：公司披露的内部代码生成比例和覆盖数据属于企业自报口径，公开材料未提供独立对照或长期质量成本全量数据。
[S20]

归因与外推边界：软件研发与营销任务不同；代码可以通过测试和编译获得更强反馈，营销质量和客户反应更难即时验证。

可复制条件：可复制条件是高频任务、即时质量反馈、版本控制和可回退。营销团队需要建立等价机制，例如事实检查、品牌规则、A/B测试与发布回滚。

第十一章 治理与扩张：制度、流程、角色、数据和反馈必须同步强化

竹势 AI 营销智库出品

NIST AI风险管理框架把治理、映射、测量和管理视为持续循环。欧盟《人工智能法》要求雇主在工作场所使用特定高风险AI系统前告知受影响员工及其代表；中国《生成式人工智能服务管理暂行办法》强调发展与安全并重、分类分级治理。对企业而言，治理不能停留在政策文本，必须转化为任务级控制、角色责任和可审计证据。[S21][S22][S23]

强化对象	90天	180天	365天
制度	试点章程、任务分级、数据边界	岗位制度、审核与事故机制	纳入内控、采购、审计与人才制度
流程	重画1—2条关键流程	复制到相邻流程并淘汰双轨	跨单元标准与本地例外
角色	联合赞助、产品负责人、领域专家	岗位职责与经理教练	职业路径与能力认证
数据	最小必要数据、日志、基线	质量标签、例外库、客户信号	统一指标、持续监控、模型/供应商比较
反馈	周度问题清单与快速修复	月度学习评审、季度闸门	年度组合复盘与预算重配

11.1 90/180/365天路线图与决议门

阶段	核心任务	推进条件	调整/回退条件	停止条件
0—90天 证明行为迁移	选定流程、建立基线、重构任务、沙盒练习、上线有限真实任务	关键任务渗透持续上升；质量达到最低阈值；例外可解释	返工增加、经理缺席、数据不可用、责任不清	高风险事故；业务目标失效；无可行数据与流程基础
91—180天 证明流程与客户效果	扩大岗位覆盖、优化控制点、连接客户行为、调整绩效	质量稳定；客户领先指标改善；单位成本可接受	尾部错误上升；客户无反应；双轨系统未退出	连续两个周期无行为或客户改善且无可信修复路径
181—365天 证明经营与可复制性	跨团队复制、制度化岗位与治理、预算重配	经营指标出现符合滞后的改善；复制团队达到门槛	不同团队结果高度不稳定；治理成本过高	经营价值不足以覆盖全生命周期成本与风险

第十二章 当代专家视角：主张、依据与边界

竹势 AI 营销智库出品

专家/主体	具体主张与依据	适用边界	本报告判断
Erik Brynjolfsson	数字技术价值取决于与流程、技能和组织资本的互补投资；生成式AI现场研究显示收益分布不均。	适用于解释为何工具采购不等于生产率；不能把单个现场实验外推为全行业ROI。	把互补资产列入投资案，而非把全部预算留给软件许可。
Raffaella Sadun	管理实践和组织能力影响技术价值兑现；P&G实验关注AI如何改变专业边界与协作。	实验环境能识别机制，但长期组织政治与激励需要另行验证。	把团队边界、知识共享和决策责任纳入设计。
Amy Edmondson	心理安全使团队能够报告错误、提问和学习，同时不排斥高标准。	不适合被误用为降低问责或允许低质量。	奖励早期暴露问题，区分善意试错与违规绕过。
Ethan Mollick	AI是一种通用能力，个人应通过真实工作实验理解边界；组织需要制定使用规范。	个人实验经验不能替代企业级治理与因果测量。	把实验转化为任务标准、知识库与审核机制。
Tsedal Neeley	数字化与远程协作变革需要领导沟通、信任和团队规范。	沟通不能替代流程和权限调整。	让经理成为变革翻译者和教练，而非信息转发器。
Daron Acemoglu	技术选择会影响劳动分工与权力分配；自动化并非唯一方向。	宏观分析不能直接决定单一企业场景。	优先选择扩展判断与客户价值的任务，警惕只为监控或削减岗位的设计。
NIST AI RMF团队	AI风险管理应持续执行治理、映射、测量与管理。	框架是自愿、跨行业的，需要企业按任务具体化。	将风险控制嵌入流程与指标，不以政策文件代替运行证据。
欧盟立法者与监管机构	高风险工作场所AI需要透明、告知、人类监督与责任安排。	具体适用范围取决于系统用途、地区和实施日期。	跨国企业应建立最严格适用要求下的可配置治理基线。

管理者工具包

竹势 AI 营销智库出品

工具1：变革准备度诊断

维度	关键证据	评分说明（0—3）	封顶/否决规则
经营目标	目标、基线、时间滞后、负责人	0无目标；1活动目标；2可测业务目标；3可归因经营目标	无业务负责人则总分封顶40%
关键行为	岗位、触发、动作、质量标准	0未定义；1笼统；2部分任务；3完整行为地图	未定义目标任务分母则否决扩张
流程系统	新流程、删除步骤、控制点、例外	0叠加工具；1局部；2已重画；3运行验证	双轨流程未退出则不得规模化
数据技术	数据可用、权限、日志、稳定性	0不可用；1手工；2有限集成；3可观测	关键数据违法或不可审计则否决
领导联盟	联合赞助、决策权、SLA	0无；1名义；2部分权力；3可调资源	无业务与风险共同赞助则封顶50%
能力练习	真实任务、反馈、经理教练、认证	0讲座；1演示；2练习；3胜任认证	仅培训完成率不得判定就绪
关系网络	可信同伴、桥接者、接班机制	0无；1单冠军；2网络；3跨团队复制	关键冠军无接班则高风险
权责安全	RACI、接管、事故、心理安全	0不清；1原则；2流程；3演练	高风险任务责任不清则否决上线
激励强化	绩效、认可、例会、预算	0冲突；1中性；2部分一致；3制度化	绩效明确奖励旧行为则封顶60%
证据与闸门	五层指标、基线、停止条件	0无；1活动；2业务链；3完整闸门	无停止条件不得追加年度预算

使用规则：各维度0—3分，总分仅用于内部比较，不代表行业分位。任何否决项优先于总分；封顶规则防止高平均分掩盖关键缺口。

工具2：利益相关者与影响网络图谱

对象	权力（高/低）	受影响程度	可信度	桥接关系	主要损益	参与策略
董事长/CEO						
CMO/业务负责人						
CHRO/HRBP						
CIO/CDO/数据团队						
法务/合规/风控						
一线经理						
关键岗位员工						
客户/渠道伙伴						

会议用法：先标出正式权力，再标出员工实际信任的人；至少识别两个跨群体桥接者。对高受影响、低权力群体安排持续反馈渠道，而非只做宣讲。

工具3：关键行为迁移地图

关键任务	旧行为	新触发与新行为	所需能力	环境/数据	即时反馈	责任/接管
示例：活动复盘	人工汇总平台数据、凭经验写结论	活动结束后自动汇总，负责人验证异常并形成下一轮动作	异常判断、归因基础	统一事件与成本数据	对账差异、业务审核	运营执行；负责人批准；数据异常人工接管
任务1						
任务2						
任务3						

工具4：五类阻力矩阵

阻力类型	诊断信号	可用干预	误用风险
理性阻力	看不到岗位收益；流程更慢；质量不稳	公布任务损益、删除旧步骤、提供对照证据	用宣传代替真实收益
情绪/身份阻力	担心贬值、失控、失去专业地位	定义新职责、同伴示范、参与设计、职业路径	把所有担忧心理化
能力阻力	只会低风险任务；遇异常无法处理	真实任务练习、反馈、经理教练、认证	重复通识培训
制度/激励阻力	绩效仍奖励旧产量；经理不分配新任务	调整目标、例会、认可、资源与预算	过早绑定奖金导致刷量
风险/责任阻力	错误责任归个人；无接管权；高风险任务绕开	任务分级、RACI、人工接管、事故机制	用“心理安全”掩盖控制缺失

工具5：领导联盟RACI与升级机制

决策事项	CEO/董事会	CMO/业务	CHRO	CIO/CDO	法务/风险	一线经理	升级时限
经营目标与预算	A	R/C	C	C	C	I	季度/重大变更
关键流程重构	I	A/R	C	C	C	R	5个工作日
岗位与绩效	I	C	A/R	I	C	R	10个工作日
数据与权限	I	C	C	A/R	C	I	5个工作日
高风险控制与事故	I	C	C	R	A/R	I	24小时
采用证据与扩张	A	R	C	C	C	C	阶段闸门

工具6：五层采用证据看板模板

层级	指标	定义/分母	基线	本期	阈值	异常解释	决议
L1 接触与可用							
L2 关键任务渗透							
L3 质量与例外							
L4 客户行为							
L5 经营结果							

工具7：角色差异化行动清单

角色	行动清单
CEO	指定联合赞助；要求一个经营结果；批准流程删除权；主持90/180/365天闸门；对无证据扩张说不。
CMO	选择高价值任务；定义客户与品牌质量；让经理分配真实任务；退出旧流程；对业务结果负责。
CHRO	重写岗位与绩效；建立能力认证；训练经理教练；监测公平与心理安全；处理职业路径和劳动沟通。
CIO/CDO	提供最小必要数据与身份权限；保证日志、稳定性、回退；建立模型和供应商可替换性；公布服务SLA。
一线经理	把新行为写进任务分配、例会和复盘；观察真实操作；保护问题上报；处理负荷与优先级冲突。
员工	在授权范围内使用；验证事实与质量；记录例外；请求接管；贡献可复用样例与改进建议。

BOARD BRIEF 17

董事会结语：把变革管理变成经营系统

竹势 AI 营销智库出品

董事会最终审议一张变革契约

经营结果 + 关键行为 + 工作接口 + 责任边界 + 采用证据 + 下一次决议日

AI营销变革不会因为工具更强而自动完成。模型能力提高，反而会放大组织中原有的目标冲突、责任模糊、数据缺陷和审批迟滞。真正的管理任务，是把经营结果、关键行为、工作系统、关系网络和强化机制连成一条可测量、可治理、可回退的因果链。

董事会不应再问“有多少员工使用AI”，而应问：哪些关键任务已经迁移；质量和例外是否受控；客户行为是否变化；经营结果是否符合预期滞后；若没有，泄漏发生在哪一层；下一阶段是推进、调整、回退还是停止。

当这些问题进入预算、绩效、流程、审计和人才制度，AI才从工具活动变成组织能力。

BOARD BRIEF 18

专题深化：把变革采用网络转化为可运营机制

竹势 AI 营销智库出品



一、关键行为必须具有触发、动作、对象和质量边界

“使用AI完成营销工作”不是可管理的行为定义。它缺少何时发生、由谁发起、针对什么对象、完成到什么质量、异常时如何处理。可执行的行为应写成：当销售机会进入方案阶段时，客户经理在四小时内调用经批准的方案流程生成第一版，核验客户事实与报价边界，交由行业负责人审核；涉及未公开客户数据或非标准条款时，立即转人工。只有这样的描述，系统才能记录任务分母、经理才能观察、员工才能练习、风险团队才能审计。

关键行为数量应受控。单个试点建议只迁移2—4个行为，因为每个行为都需要流程、数据、权限、培训、质量、例外和绩效支持。一次定义二十个行为，通常意味着每个行为都只获得浅层支持。董事会应要求项目组说明行为选择逻辑：经营杠杆、发生频率、可测量性、可逆性、数据可用性以及跨部门依赖。

行为设计字段	管理问题	合格示例	不合格示例
触发	何时必须发生？	新线索评分达到阈值后15分钟内	需要时使用
执行者	谁实际完成？	区域营销运营	市场部
动作	具体改变什么？	生成并核验三类跟进建议，选择一类发送	用AI跟进
对象	作用于谁或什么？	进入培育阶段且已授权触达的线索	所有客户
质量	什么算合格？	事实准确、品牌合规、CTA明确、一次审核通过	效果更好
例外	何时停止或接管？	高价值客户、敏感行业、低置信度信息	出问题再说
证据	如何记录？	任务日志、审核、发送、客户响应	员工自报

二、工作系统的七个最小组成

工作系统至少包含目标、输入、流程、工具、角色、控制和反馈七个组成。任何一个组成仍按旧逻辑运行，员工都会在新旧方式之间折返。例如，内容团队使用AI生成多版本，但预算审批仍要求先确定唯一创意；数据团队提供实时表现，但复盘会议仍按月召开；销售团队收到线索建议，但CRM字段不允许记录建议来源和接管原因。局部自动化因此没有形成端到端收益。

组成	必须明确	审计问题
目标	业务结果与任务目标	是否存在活动目标替代经营目标？
输入	数据、知识、客户授权与时效	输入是否完整、合法、可追溯？
流程	步骤、顺序、等待与删除项	是否真正退出旧步骤？
工具	系统能力、稳定性和可替换性	故障时是否可降级？
角色	执行、审核、批准、接管	责任是否与权力匹配？
控制	阈值、抽样、日志、事故机制	控制是否按风险分级？
反馈	客户、质量、成本与学习信号	反馈是否能改变规则与流程？

社会技术系统理论提醒管理者，技术子系统与社会子系统需要联合优化。该思想源于煤矿工作设计等研究，不能机械套用到数字营销；但其核心边界判断仍有效：单独优化算法准确率，可能增加审核负担、降低自主性或破坏团队知识流动。项目评审必须同时观察技术性能和工作质量。[S25]

三、经理层是采用系统的“传动轴”

一线经理决定任务如何分配、什么问题值得升级、谁获得练习机会、旧流程何时退出以及绩效如何解释。很多企业把经理当作沟通渠道，只要求其转发培训和催促登录，结果经理继续按旧标准验收工作，员工自然把AI放在边缘。经理需要三类能力：任务设计能力，能把目标转换成具体行为；质量教练能力，能基于样例给出反馈；冲突处理能力，能在速度、质量、客户与风险之间做升级。

经理的工作量也必须重估。变革初期，观察、复核和教练会增加时间投入。如果组织同时维持原有管理跨度和全部旧报表，经理会通过减少辅导、扩大形式审核来保护自身绩效。合理做法是暂时降低非关键汇报要求，给经理明确的教练时段，并用问题解决速度而非登录率评价其变革贡献。

经理每周机制	目的	最低产出
任务分配会	把新行为嵌入真实任务	本周目标任务清单与责任人
质量校准	统一“合格”标准	2—3个通过/不通过样例
例外复盘	发现系统与流程缺口	例外分类、责任、修复时限
同伴示范	提高可观察性和自我效能	相似岗位完整演示
负荷检查	防止AI增负与双轨	删除项、等待项、过载岗位
数据回顾	连接行为与客户结果	五层证据趋势与决议

四、沟通不是宣传，而是持续更新的组织契约

有效沟通需要回答员工反复提出的具体问题：为什么选择这个流程；哪些岗位和任务会变化；节省时间后工作量如何安排；哪些数据不可输入；哪些输出必须审核；错误如何处理；谁有权停止；绩效和晋升如何变化；项目何时评估和可能停止。若这些问题没有答案，频繁宣传只会提高不信任。

沟通还要承认不确定性。管理层可以明确区分已决定事项、试验假设和待协商事项。把尚未验证的效率承诺包装成确定事实，会形成新的心理契约风险。一旦员工发现实际返工更高或岗位安排与承诺不符，后续再提供真实信息也难以恢复信任。

沟通层级	内容	频率	责任人
董事会/经营层	目标、资源、风险偏好、阶段决议	季度与重大事件	CEO/联合赞助人
经理层	任务、质量、负荷、绩效解释	每周	业务负责人/HRBP
员工层	操作边界、案例、例外、职业影响	上线前+持续	一线经理/领域专家
风险与技术	数据、权限、事故、版本变化	持续与事件触发	CIO/CDO/风险负责人
客户与伙伴	自动化使用、授权、接管与申诉	按场景	业务/法务

五、阻力诊断必须先区分原因，再选择干预

同一个“不使用”行为可能来自完全不同原因。员工可能认为工具质量不足，这是理性阻力；可能担心专业身份下降，这是身份阻力；可能不会处理异常，这是能力阻力；可能因为绩效仍要求旧格式，这是制度阻力；也可能因为错误责任无法承受，这是风险阻力。用更多培训处理制度问题、用激励处理责任问题、用宣讲处理质量问题，都会加剧抵触。

诊断应优先使用行为和系统证据，而非泛化态度问卷。项目组可对比不同任务、不同经理和不同风险等级的采用差异；观察员工在低风险任务上是否积极、在高风险任务上是否回避；核查返工、等待和接管记录；分析问题是否集中在某一数据源、审核人或客户类型。态度数据可作为补充，但不能替代流程证据。

六、探索与规模化需要不同的组织节奏

探索阶段的目标是学习哪种任务、流程和控制组合可行。团队需要短周期实验、透明记录失败、快速调整权限与规则。规模化阶段的目标是稳定、复制和经济性，需要版本管理、服务SLA、能力认证、审计与预算纪律。把规模期控制过早施加到探索期，会让实验变慢；把探索期的灵活做法带入规模期，则会造成不可审计和个人依赖。

管理要素	探索期	规模期
范围	单流程、少岗位、可逆任务	跨团队、跨系统、真实客户规模
目标	验证机制和最低可行控制	稳定质量、客户结果与经济性
决策	快速、授权项目团队	标准化、例外治理、变更委员会
数据	足够验证假设的最小数据	统一口径、持续监控、审计保留
人员	高参与的先行者+代表性员工	岗位覆盖、经理体系、接班机制
绩效	学习、问题暴露、流程改进	稳定采用、质量、客户与经营结果
退出	随时回退	受控迁移、业务连续性方案

七、全生命周期经济性必须包含“隐性变革成本”

AI项目商业论证常只计算许可费和推理成本，忽略集成、数据清理、审核、经理教练、合规、事故处置、双轨运行、供应商切换和员工时间。尤其在试点早期，人工复核可能高于节省的执行时间。该现象不代表项目必然失败，但要求管理层明确成本下降的机制与时间：质量是否会通过知识回流提高；审核是否会从全量转抽样；旧系统何时退出；异常是否会减少；经理教练是否形成标准。

成本类别	容易遗漏的项目	管理要求
获取与部署	许可、模型、连接器、环境	分开一次性与持续成本
数据与知识	清理、授权、标签、更新	明确数据产品负责人
流程变更	设计、测试、双轨、退出旧系统	设定双轨结束日期
人员能力	培训、练习、经理教练、专家审核	计入岗位时间而非视为免费
质量与风险	抽样、审计、事故、法律与客户补救	按尾部风险预留
运营维护	监控、版本、供应商、故障与回退	服务SLA和替换方案
机会成本	延迟项目、员工注意力、客户试错	与其他增长投资比较

八、双环学习：不仅修正输出，还要修正规则

Argyris和Schön区分单环学习与双环学习。单环学习是在既定目标和规则下纠正偏差，例如优化一个生成模板；双环学习会质疑目标、假设和治理规则，例如发现“每周发布十篇内容”本身与客户价值无关，转而重设内容组合和销售接口。AI系统容易形成大量低成本输出，如果组织只做单环优化，就会更高效地执行错误目标。[S26]

每月学习评审应至少讨论四类问题：哪些错误来自模型或数据；哪些来自流程和角色；哪些来自目标和指标；哪些来自客户和市场假设。只有前两类问题被讨论时，变革会退化为技术优化。双环学习的边界是不能把每次异常都升级为战略重审，否则组织失去稳定性。应以重复出现、跨团队出现或影响经营结果的模式作为触发。

九、失败、受限和停止是组合管理的一部分

并非所有AI营销场景都值得规模化。低频、数据稀缺、质量难测、错误成本高且缺乏人工接管的任务，可能长期不具备经济性。成熟的变革组合应允许项目被停止，并把停止理由转化为知识资产：目标不成立、技术能力不足、数据不可获得、流程无法改变、风险超过收益、客户不接受或全生命周期成本过高。

停止不等于否定AI战略。它释放预算和管理注意力，让资源进入更适合的任务。相反，没有停止机制的项目会利用“技术还在进步”“员工需要更多时间”等无法证伪的解释不断延期。董事会应要求每个项目在启动时写下停止条件，并在阶段闸门上由不承担沉没成本责任的人参与评审。

十、跨团队复制必须证明“机制可移植”

一个团队成功并不意味着另一个团队可以复制。复制前需要区分核心机制与本地条件。核心机制可能是快速反馈、明确接管和经理教练；本地条件可能是客户类型、数据成熟度、岗位经验和监管要求。项目组应形成“最小不可变核心”和“可配置参数”：哪些质量阈值不可降低，哪些工作流可本地调整，哪些数据必须统一，哪些审批可按风险变化。

跨团队复制应采用阶梯式样本：先在相似团队验证，再进入差异更大的地区、业务或客户群。若结果下降，应判断是执行偏差还是机制依赖本地条件。只有在至少两个具有差异的团队中达到行为、质量和客户门槛，才能声称具备组织级可复制性。

十一、营销场景的差异化变革设计

营销并非单一职能。品牌内容、广告投放、销售线索、客户运营、市场研究和高管沟通的任务频率、错误成本、反馈速度与责任主体不同，因此不能使用同一采用方案。内容生成高频、可逆，但品牌和事实错误会累积信任损失；广告投放反馈快，却涉及预算和平台权限；线索跟进直接影响客户关系，要求授权和接管；市场研究容易出现来源与推理混淆；高管沟通低频、高影响，需要更严格的事实与责任控制。

场景	主要价值机制	关键行为	核心风险	采用证据
品牌内容	缩短构思与适配周期、提高测试密度	基于品牌上下文生成、核验、选择并发布	事实错误、同质化、品牌漂移	一次通过率、发布周期、客户互动、品牌风险事件
广告投放	更快发现异常并形成优化建议	按阈值巡检、生成变更草案、人工批准	预算损失、错误归因、权限越界	建议采纳率、异常响应、增量效果、回退次数
线索运营	提高响应速度与跟进一致性	线索进入后生成建议、人工确认、记录结果	未经授权触达、客户冒犯、敏感信息	响应时长、覆盖率、有效对话、转化与投诉
市场研究	扩大证据覆盖并缩短分析周期	检索、筛选、标注、综合、提出可证伪判断	伪来源、过时信息、推理越界	来源可核验率、事实错误、决策采用、后验准确性
高管沟通	快速形成结构化初稿与情景分析	生成草案、事实核验、责任人签发	重大误导、保密、责任模糊	修改幅度、核验通过、决策时效、事故为零

场景差异决定经理教练方式。内容团队需要样例校准和品牌判断；投放团队需要实验设计与因果意识；线索团队需要客户同理心、授权和升级；研究团队需要证据分级与不确定性表达。用统一课程覆盖所有岗位，只能形成共同语言，不能形成岗位胜任。

十二、对员工影响的公平性与可持续性审查

AI变革可能重新分配高价值任务、学习机会和可见度。若只有少数“数字原住民”获得关键项目，其他员工被留在重复劳动中，组织会扩大能力差距；若系统主要学习资深员工的经验，却不承认其贡献，知识共享意愿会下降；若节省时间全部转换为更高工作量，员工会把效率收益理解为隐性裁员或劳动强化。

CHRO与业务负责人应在阶段闸门上审查五项公平性：不同岗位是否拥有合理的学习机会；高风险审核工作是否被隐形压给少数专家；绩效评价是否考虑任务难度和风险；员工贡献的知识是否被确认；工作负荷和监控强度是否发生不透明变化。公平性不是附加的员工体验议题，它直接影响知识回流、问题上报和长期采用。

审查主题	证据	风险信号	纠偏动作
学习机会	练习任务与教练覆盖	机会集中在少数人	轮换、最低练习量、经理责任
知识贡献	规则、样例和审核贡献记录	专家持续输出但无认可	贡献署名、时间预算、绩效认可
工作负荷	净工时、返工、审核与会议	节省时间被新增任务吞噬	删除旧工作、负荷上限、岗位调整
监控与自主	日志用途、可见范围、申诉	员工不知道数据如何被使用	透明告知、最小化、访问与申诉机制
绩效公平	任务难度、风险与结果	只看数量或系统活跃	组合指标、情境校准、人工复核

十三、供应商与技术变化下的组织连续性

AI能力、价格、数据政策和产品接口变化很快。企业若把岗位流程、知识和审计记录锁在单一供应商中，会把变革成果变成新的依赖。技术治理应把业务流程、知识资产、评测集、权限规则和日志保留在企业可控制的层面；模型与工具作为可替换组件。

供应商切换能力不是要求多云多模型的复杂架构，而是保持最低可移植性：关键输入输出采用开放格式；业务规则 and 知识来源可导出；评测集不依赖单一平台；身份与审批由企业系统控制；终止服务后数据删除和留存责任明确。否则，员工形成的新行为可能随着产品变化而失效，组织再次回到手工。

连续性要素	最低要求	董事会风险
业务规则	规则、阈值、审批可导出	关键控制被供应商黑箱化
知识产权	来源、版本、权限可管理	经验无法迁移或权属不清
评测与基线	企业拥有固定测试集	模型升级后无法判断质量变化
身份与日志	企业统一身份、完整留痕	离职权限与审计失控
数据退出	删除、返还、保留期限明确	商业敏感信息长期残留
降级方案	故障时可转人工或替代服务	关键营销流程中断

十四、董事会问题库：避免被“采用热闹”误导

董事会和经营层应使用能够暴露因果链的问题，而不是只要求更漂亮的仪表盘。以下问题可用于季度审议：目标经营结果的基线是什么，预期滞后多长；目标任务分母是多少，多少任务真正按新方式完成；质量改善是否来自筛选简单任务；返工和人工审核是否被计入成本；客户是否知道或感受到自动化；高价值员工是否因系统建议降低判断质量；哪些旧步骤已经正式退出；最近三次人工接管说明了什么；哪个风险或成本信号会触发停止；若供应商明天不可用，流程如何继续。

这些问题的价值在于迫使项目团队说明机制与边界。无法回答并不必然意味着项目失败，但意味着当前证据不足以扩大预算、权限或客户范围。董事会最重要的纪律，是把“有前景”与“已证明”分开，把战略信念转化为分阶段、可回退的资本配置。

十五、变革办公室的最小运行制度

变革办公室不应成为独立于业务的汇报层。其职责是维护因果链、推动跨部门决策、管理证据和阶段闸门。业务团队拥有结果，技术团队拥有可靠性与数据，HR拥有岗位与能力，风险团队拥有控制框架；变革办公室负责把这些责任连接起来，并在问题超出单一部门权限时触发升级。

最小运行制度包括一份统一问题清单、一个决策日志、一套五层证据看板、一组版本化流程和每两周一次的跨职能解决会议。问题清单必须记录影响、责任人、截止时间和升级状态；决策日志记录为何选择某种控制或阈值，防止人员变化后重复争论；流程版本记录旧步骤何时退出，避免双轨长期存在。

运行机制	输入	输出	失效信号
每周业务复盘	任务、质量、客户与负荷数据	本周改进与责任	只讨论工具使用
双周跨职能解决会	未解决依赖与风险	决策、资源、SLA	问题反复延期
月度学习评审	错误模式、例外、客户反馈	规则与流程版本更新	只修模型参数
阶段闸门	五层证据、经济性、风险	推进、调整、回退、停止	默认自动续期
季度董事会审议	经营结果、组合比较、重大风险	预算与优先级重配	用案例故事替代数据

变革办公室还应定期检查“影子AI”与绕行行为。员工私下使用未经批准工具，可能说明正式工具不可用、流程过慢或需求未被满足，也可能造成数据和合规风险。治理不应只靠封禁，而应分析绕行原因：对合理需求提供受控替代，对高风险行为明确后果，并持续缩短正式流程的响应时间。

十六、最终决议标准：何时可以称为组织能力

只有当新行为不依赖少数热心员工、在经理任务分配中稳定出现、由系统提供所需数据和权限、质量与例外可以监测、客户与经营结果出现符合逻辑的变化、人员变动或工具升级后仍能继续运行，企业才能把试点称为组织能力。任何一项依赖个人记忆、临时审批或手工补丁，都意味着能力尚未制度化。

组织能力也不是永久状态。市场变化、模型升级、监管更新和客户偏好会使原有控制失效。企业需要把评测集、流程、岗位和风险规则作为持续维护的经营资产。年度预算不只购买技术，还要覆盖知识更新、能力复训、审计、供应商切换和流程再设计。

补充判定：当项目只能证明员工更频繁地接触工具，却无法证明目标任务迁移、质量稳定、客户反应或经营结果时，管理层应把它归类为能力探索，而非业务转型。探索可以继续，但预算、权限和外部客户范围必须与证据等级匹配。任何扩大都应说明新增风险由谁承担、旧流程何时退出、下一次闸门将用什么数据作决定。

PUBLISHER

关于竹势 AI 营销智库

竹势 AI 营销智库出品

竹势 AI 营销智库关注企业增长、营销组织变革、智能体运营与 AI 治理，帮助企业老板、CEO、CMO 和数字化管理者把复杂技术议题转化为可讨论、可比较、可执行的经营决策。

本报告的判断框架、责任边界和实施工具，旨在帮助管理团队用受控方式验证 Agentic Marketing 的价值，并在规模化之前建立必要的证据、权限与问责机制。

shichangbu.ai

专题中心 · 返回顶部 ↑

纯文本来源索引

竹势 AI 营销智库出品

[S1] Kurt Lewin, *Frontiers in Group Dynamics, Human Relations*, 1947；场论与群体动力学。

[S2] John P. Kotter, *Leading Change*, Harvard Business School Press, 1996；指导联盟与变革步骤。

[S3] Daniel Kahneman and Amos Tversky, *Prospect Theory: An Analysis of Decision under Risk*, *Econometrica*, 1979。

[S4] Denise M. Rousseau, *Psychological Contracts in Organizations*, Sage, 1995。

[S5] Microsoft and LinkedIn, *2024 Work Trend Index: AI at Work Is Here. Now Comes the Hard Part*, 2024-05-08；31国约31,000名知识工作者及平台信号。

[S6] Michael Hammer, *Reengineering Work: Don't Automate, Obliterate*, *Harvard Business Review*, 1990。

[S7] Erik Brynjolfsson, Danielle Li, Lindsey R. Raymond, *Generative AI at Work*, NBER Working Paper 31161, 2023；5,179名客户支持人员分阶段部署研究。

[S8] Fabrizio Dell'Acqua et al., *The Cybernetic Teammate: A Field Experiment on Generative AI Reshaping Teamwork and Expertise*, Harvard Business School Working Paper 25-043, 2025；P&G 791名专业人员随机现场实验。

[S9] Albert Bandura, *Self-Efficacy: The Exercise of Control*, W.H. Freeman, 1997。

[S10] David A. Kolb, *Experiential Learning*, Prentice Hall, 1984。

[S11] Fabrizio Dell'Acqua et al., *Navigating the Jagged Technological Frontier*, Harvard Business School Working Paper 24-013, 2023；758名BCG顾问实验。

[S12] Everett M. Rogers, *Diffusion of Innovations*, 5th edition, Free Press, 2003。

[S13] Amy C. Edmondson, *Psychological Safety and Learning Behavior in Work Teams*, *Administrative Science Quarterly*, 1999。

[S14] *Moffatt v. Air Canada*, 2024 BCCRT 149, British Columbia Civil Resolution Tribunal, 2024-02-14。

[S15] Charles Goodhart, *Problems of Monetary Management: The U.K. Experience*, 1975；指标目标化问题的原始语境。

[S16] Haier Smart Home Co., Ltd., *Annual Report 2023*, 2024-03-27；数字转型、营销、用户与组织机制披露。

[S17] Haier Smart Home Co., Ltd., *Interim Report 2025 / Annual Report 2025*, 2025—2026；数字库存、数字营销、经销商与畅销型号口径。

[S18] Ping An Insurance (Group) Company of China, Ltd., Annual Report 2019, 2020-02-20；平安银行敏捷转型、端到端研发流程与经营指标。

[S19] China Merchants Bank Co., Ltd., Annual Reports 2023—2025；金融科技投入、零售客户经营、流程与风险管理披露。

[S20] Baidu, official developer conference and annual disclosures on internal intelligent coding assistant adoption, 2023—2025；企业自报口径。

[S21] NIST, Artificial Intelligence Risk Management Framework 1.0, NIST AI 100-1, 2023。

[S22] NIST, Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile, NIST AI 600-1, 2024-07-26。

[S23] Regulation (EU) 2024/1689, Artificial Intelligence Act, 2024-06-13；工作场所高风险AI的告知、人类监督与治理要求。

[S24] 国家互联网信息办公室等七部门，《生成式人工智能服务管理暂行办法》，2023-07-13发布，2023-08-15施行。

[S25] Fred Emery and Eric Trist, The Causal Texture of Organizational Environments, Human Relations, 1965；社会技术系统传统。

[S26] Chris Argyris and Donald Schön, Organizational Learning II, Addison-Wesley, 1996；单环与双环学习。

[S27] World Economic Forum, Future of Jobs Report 2025；雇主技能、岗位与转型预期调查，属调查口径。

[S28] McKinsey Global Survey, The State of AI, 2024—2025；企业自报采用与价值调查，非因果证据。

[S29] OECD, Employment Outlook 2023: Artificial Intelligence and the Labour Market；工作质量、技能与劳动影响。

[S30] International Labour Organization, Generative AI and Jobs: A Global Analysis of Potential Effects on Job Quantity and Quality, 2023；任务暴露分析。

[S31] 中华人民共和国个人信息保护法，2021-08-20通过，2021-11-01施行。

[S32] 中华人民共和国数据安全法，2021-06-10通过，2021-09-01施行。

[S33] 国务院，《关于深入实施“人工智能+”行动的意见》，2025；企业战略、组织和流程融合方向。

[S34] Microsoft, 2026 Work Trend Index: Agents, Human Agency, and the Opportunity for Every Organization, 2026-05-05；“吸收”与学习系统观点，属企业研究。

附录A：案例证据分级与使用边界

竹势 AI 营销智库出品

证据类型	可信用途	主要局限	写作标识
同行评审/严谨现场实验	识别因果与机制	任务、样本与时期外推有限	说明样本、任务、时间与外推边界
监管/法院文件	确认义务、责任与事件事实	司法辖区和适用范围有限	说明地区、制度与生效时间
上市公司财报/官方披露	确认公司动作与披露结果	自报、综合结果、缺少对照	明确“公司披露”“未单独归因”
供应商案例	理解实施方式与可能结果	选择性报告、利益冲突	明确供应商口径，不作行业基准
权威调查	描述感知、采用和趋势	自报偏差，不能证明因果	写明样本、地区、时间和问题口径

附录B：阶段会议议程模板

竹势 AI 营销智库出品

议题	必须回答的问题	输入材料	决议
经营结果	指标是否按预期滞后变化？	基线、趋势、对照	推进/调整/停止
关键行为	目标任务中多少已迁移？	任务分母、岗位覆盖	扩大/聚焦
质量与例外	平均质量和尾部风险如何？	抽样、错误、接管、投诉	调整控制点
流程与负荷	是否减少步骤、等待和返工？	流程时长、员工负荷	删除旧流程/回退
组织与能力	经理是否教练，岗位是否胜任？	认证、教练、网络	补能力/换负责人
风险与合规	控制是否有效，是否有新风险？	日志、审计、事故	限制/暂停
经济性	全生命周期成本是否可接受？	许可、集成、审核、治理成本	追加/重配预算