



HUMAN × AI × PARTNER /
ORGANIZATION DESIGN

人机协同的AI市场部 组织设计白皮书

把岗位口号改写为可运行的任务契约：谁判断、谁执行、谁交接、谁批准、谁对结果与事故负责。

重写组织基本单元

用七字段任务契约连接目标、权限、验收与责任

分清三类角色边界

人的判断、AI 的执行与伙伴的专业能力各有条件

让交接与例外可运行

交接包、审批门、升级时限、暂停与回滚

选择适合的组织形态

中心化、嵌入式、联邦式、平台型与轻量网络

六项组织审议

每项审议都留下边界、证据、责任人和重新评估条件。

01 / TASK

工作如何被拆解

任务契约与验收

02 / ROLE

谁最适合承担

人、AI 与伙伴责任谱

03 / AUTHORITY

可以自主到哪一步

风险、可逆性与证据

04 / HANDOFF

上下文如何接住

交接包与接收确认

05 / EXCEPTION

异常如何升级

审批、暂停与回滚

06 / FORM

组织如何长期运行

形态、节奏与绩效

BOARD BRIEF 01

出版说明

竹势 AI 营销智库出品

本白皮书面向企业董事长、CEO、CMO、CHRO、CIO/CDO、法务与安全负责人、营销负责人、业务单元负责人，以及代理商与生态伙伴管理者。报告讨论的对象不是某一种模型或软件，而是企业如何把人、AI 员工与外部伙伴组织成可运行、可审计、可扩展的营销生产系统。

本报告的核心命题是：当 AI 能够持续参与判断、生成、执行与监测时，组织设计的基本单元需要从静态岗位说明扩展为任务契约。任务契约必须写清目标、输入、权限、输出、验收、升级与最终责任。岗位仍然存在，但岗位不再足以描述跨人机边界的真实工作。

报告中的数字均保留其公开来源的时间、样本和口径。企业自报与平台披露仅作为该主体的实施证据，不外推为行业平均。未公开的成本、人员、回报与系统细节不作推算。文末来源索引只保留机构、作者、标题、日期及口径，不含第三方可点击链接。

BOARD BRIEF 02

执行摘要

竹势 AI 营销智库出品

BOARD ORGANIZATION RULE / 董事会组织原则

组织设计的基本单元，是一份把目标、输入、权限、验收、升级与责任写清楚的任务契约。

任务可委派

权限可撤回

交接可确认

例外可升级

结果可问责

AI 进入市场部后，最危险的组织误判是把它当作“更快的个人工具”。个人使用可以带来局部节省，但组织结果取决于任务如何被拆分、谁拥有何种权限、何时交接、什么结果算完成、异常由谁接管，以及最终由谁承担经营和法律责任。缺少这些机制，产量越高，返工、品牌偏差、数据风险和责任模糊越可能同步扩大。

表 1 | 十项可证伪核心判断

编号	核心判断
判断一	组织设计的基本单元正在由岗位扩展为任务契约。岗位回答“谁在这里”，任务契约回答“这项工作以什么目标、输入、权限、输出、验收、升级和责任运行”。
判断二	人、AI 与伙伴的边界不是按职位划线，而是按判断强度、可逆性、数据敏感度、规模收益、关系信任和责任后果配置。
判断三	授权不是开或关，而是从建议、草拟、受控执行、限额处置到条件自治的梯度。每一次升级都需要新的证据门和可撤销控制。
判断四	人机协作的大部分失误发生在交接处，而非生成处。没有状态、证据、未决项和下一责任人的交接，人工复核会退化为“重新做一遍”。
判断五	全量人工审批并不等于安全。它会制造积压、走过场与确认偏差。高后果事项逐项审批，低风险事项抽样，高异常信号即时升级，才是可持续的控制组合。
判断六	不存在唯一最优组织形态。中心化 CoE 适合建立标准，业务嵌入式适合贴近场景，联邦式适合规模化治理，平台型适合复用能力，轻量网络型适合资源有限企业。
判断七	外部伙伴必须进入同一任务契约。只签交付范围而不定义数据、验收、知识转移、审计和退出，会把短期专业性变成长期依赖。
判断八	AI 不能承担法律和管理责任。共同绩效可以分摊团队贡献，但必须保留最终责任人、审批记录和可追溯证据。
判断九	能力建设重点从“会使用工具”转向问题定义、证据判断、品味、关系、伦理和系统改进，同时保留基础技能，防止自动化偏差和能力退化。
判断十	人机边界必须按日、周、月、季度运行节奏持续重设。模型、数据、人员熟练度与监管环境变化都会移动最优边界。

董事会应把人机协同视为一种组织操作系统：以任务契约为底层对象，以授权梯度为控制面，以交接协议为接口，以共同绩效为反馈，以周期性边界重设为演进机制。

表 2 | 董事会一页决策框架

问题	董事会必须看到的证据	不能接受的替代品
为什么做	明确的经营任务、基线与责任人	“同行都在做”或工具演示
谁来做	人/AI/伙伴责任谱与最终责任人	笼统的人机协同口号
能做多深	授权等级、额度、暂停、回滚	一次性全开权限
如何控制	审批矩阵、异常升级、日志与抽样	所有事项都人工复核
如何扩展	90 天证据门与迁移条件	先铺平台再找场景

BOARD BRIEF 03

第一章 董事会命题：岗位之外，建立任务契约

竹势 AI 营销智库出品

任务契约七字段

先写清工作如何被接住，再决定由人、AI 或伙伴执行

01	目标	服务哪项经营结果
02	输入	上下文与证据
03	权限	可读取与可执行
04	输出	格式、状态、记录
05	验收	质量与完成定义
06	升级	异常与接管条件
07	责任	结果与事故归属

传统岗位说明适合相对稳定的工作：岗位拥有一组职责，管理者通过层级、流程与绩效指标完成协调。AI 员工进入后，同一项营销任务可能由人定义目标，AI 检索资料和生成方案，外部代理商完成制作，AI 再监测反馈，最终由业务负责人批准发布。岗位边界仍有价值，却无法完整描述跨主体、跨系统、跨时间的责任链。

任务契约提供更精确的组织语言。它不是法律合同，而是运行协议：目标说明经营结果；输入说明允许使用的数据与上下文；权限说明可调用工具和可执行动作；输出说明交付格式；验收说明质量阈值；升级说明异常条件；责任说明谁对最终结果负责。七项中任何一项缺失，协同都可能转化为隐性返工。

实证研究支持“系统设计比工具使用更重要”。Brynjolfsson、Li 与 Raymond 对 5,179 名客服人员的分阶段部署研究显示，AI 辅助使每小时解决问题数平均提高 14%，新手和低技能员工提高 34%，高经验员工影响很小，甚至出现部分质量下降。收益来自把高水平员工的沟通模式嵌入工作流，而不是简单增加一个聊天窗口。研究同时提醒，效果高度异质，不能把平均值当作所有岗位的承诺。

微软与 LinkedIn 2024 年覆盖 31 个国家、31,000 人的工作趋势调查发现，75% 的知识工作者已在工作中使用 AI，78% 的使用者把自选工具带入工作场景，但只有 39% 接受过企业培训。该口径说明个人采用速度快于组织治理速度；它不能证明企业绩效改善，却揭示了“影子 AI”形成的治理缺口。

低风险、低敏感、可逆的个人辅助任务不需要复杂设计。例如内部头脑风暴、无客户数据的标题草拟、公开资料摘要，可以使用简化契约。组织设计应与后果匹配，避免把每一次低风险使用都变成审批项目。

表 3 | 任务契约缺口与典型后果

缺失项	典型后果	董事会信号
目标	产量增加但经营结果无变化	内容量上升，线索与品牌指标不变
输入	引用过期、越权或不完整数据	团队不断补资料
权限	AI 能说不能做，或能做但失控	大量手工搬运或越权事故
输出	格式不一致，无法进入下一环节	下游重新整理
验收	人工只能凭感觉复核	审核时间持续上升
升级	异常被静默处理或无限搁置	高风险事项无人接管
责任	出错后归因给“系统”	无人拥有整改闭环

表 4 | 适用性分层

任务类型	契约深度	示例
个人低风险辅助	简版：目标、输入、输出、责任	内部构思、公开信息摘要
跨团队生产任务	标准版七要素	活动方案、内容生产、投放优化
高风险外部动作	增强版：七要素+限额+日志+回滚	对外发布、客户触达、预算调整

案例一 | Fortune 500 软件企业客服：AI 把隐性经验嵌入一线工作

案例卷宗

要素	内容
主体与时间	Fortune 500 软件企业客服组织，2020–2021 年分阶段部署
业务情境	文本客服，高频重复，质量与效率可观察
具体动作	AI 提供回复建议，客服选择采纳或忽略
结果证据	5,179 名客服；每小时解决问题数平均提高 14%，新手/低技能提高 34%
归因限制	单一企业与客服场景；高技能员工收益有限
管理启示	把高水平做法嵌入任务，不把工具平均效果外推到所有岗位

研究对象是某 Fortune 500 企业的软件客服组织，工具在 2020 至 2021 年间分阶段上线。系统向客服提供回复建议，员工可以采纳或忽略，最终仍由员工与客户互动。AI 因而处于“建议—人工决定”的授权级别。

结果显示总体生产率提高 14%，收益集中在新手和低技能员工；新员工两个月时的表现接近未使用工具、任职六个月以上的员工。研究还发现高技能员工收益有限，部分质量指标可能下降。机制是把高水平客服的表达模式扩散给新手，而不是无差别替代。

归因边界很清楚：研究发生在文本客服场景，任务高度重复、反馈快速、质量可度量；不能直接外推到品牌战略、危机公关或复杂 B2B 销售。管理启示是先寻找“经验可编码、结果可观察、员工可选择采纳”的任务，再讨论扩大授权。

第二章 任务分解：把工作拆成可委派、可审阅、可升级的单元

竹势 AI 营销智库出品

经营任务

完整客户与业务情境

拆解

可委派单元

目标、输入、权限、验收

组合

协作闭环

交接、例外、复盘、改进

过度切碎会丢失情境；拆解粒度以能独立验收、又保留关键判断为界。

岗位描述通常以“负责品牌、内容、投放、增长”概括工作，粒度过粗，无法决定哪些步骤可以委派给 AI。可操作的拆分需要沿着决策点展开：目标是否明确、输入是否结构化、结果能否验证、错误是否可逆、是否需要关系判断、责任后果是否可以被人承担。

任务越可结构化、反馈越快、错误越容易撤销，越适合较深授权。任务越依赖隐性知识、身份信任、跨情境权衡和价值判断，越需要人保持主导。Polanyi 所说“我们知道的多于能够说出的”，在营销中表现为品牌品味、组织政治、客户语气和时机判断。过度拆分会切断这些上下文，造成局部正确、整体失真。

因此任务分解必须设置“完整感保护”：每个任务单元都要指向上位经营目标；关键角色保留端到端可见性；定期把碎片重新组合为客户旅程、战役或品牌叙事；员工不仅审核单个输出，还参与改进整个系统。

表 5 | 任务契约画布

字段	必须回答的问题	责任人
目标	这项任务改变哪个经营结果？	业务负责人
输入	允许使用哪些数据、版本和上下文？	数据/任务所有者
权限	能读取、生成、调用、发送或修改什么？	系统与风险负责人
输出	以何种结构交付给谁？	任务设计者
验收	质量、事实、品牌、合规阈值是什么？	专业审核人
升级	什么异常必须暂停并交给谁？	运行负责人
责任	谁对最终业务与法律后果负责？	具名管理者

表 6 | 任务分解判定矩阵

判定维度	低分特征	高分特征	组织动作
可验证性	结果主观、反馈慢	结果可测、反馈快	高分可提高自动化
可逆性	发布后难撤回	可在沙箱回滚	低可逆保持人工批准
情境依赖	依赖关系与隐性信息	输入可结构化	高依赖保留人主导
规模收益	低频一次性	高频重复	高频优先产品化
责任后果	影响轻微	影响客户、资金、合规	高后果降低授权

表 7 | 过度切碎的五個预警

预警	表现	修复
上下文丢失	每一步都合格，整体不连贯	设置端到端任务所有者
员工失去完整感	只剩审核碎片	让员工参与目标与复盘
异常无人解释	系统只报错不说明影响	增加未决项与影响字段
优化局部指标	追求速度牺牲品牌	增加系统级质量指标
技能退化	员工不会从零完成关键任务	保留手工演练与轮岗

BOARD BRIEF 05

第三章 角色组合：配置人、AI 与外部伙伴的比较优势

竹势 AI 营销智库出品

HUMAN

人类判断

问题定义 · 权衡 · 品味 · 关系 · 责任

AI

机器执行

检索 · 生成 · 分类 · 编排 · 监测

PARTNER

外部专业

创意 · 媒介 · 实施 · 行业知识 · 峰值产能

角色边界随任务风险、证据和能力变化；责任归属不能随执行者漂移。

人、AI 与伙伴各自的优势并非固定。AI 的速度、规模和一致性适合检索、变体生成、监测和规则明确的执行；人更擅长目标选择、价值判断、关系、品味、伦理与承担责任；外部伙伴可以提供稀缺专业性、跨客户经验和弹性产能，但会增加协调、知识泄漏与依赖成本。

Simon 的有限理性提醒管理者，任何主体都在信息不完整、时间有限和注意力稀缺下决策。AI 可以扩展搜索和比较范围，却不能自动消除目标冲突；外部伙伴可以带来经验，却可能优化自身合同目标；人拥有最终责任，却可能受确认偏差影响。责任谱的价值在于让这些限制彼此制衡。

边界会移动。模型能力提高、企业知识库完善、员工熟练度增加后，原本需要人工完成的草拟和分析可以进入受控执行；当品牌进入危机、数据质量下降或监管变化时，同一任务应降级授权。组织必须把边界看作动态参数，而不是一次性岗位调整。

表 8 | 人 / AI / 外部伙伴责任谱

任务阶段	人	AI	外部伙伴	最终责任
定义目标	主导	提供分析	提供外部视角	业务负责人
收集证据	设定范围	规模检索与整理	补充专业来源	研究负责人
生成方案	判断与取舍	生成选项与模拟	提供行业经验	方案所有者
制作执行	审核关键创意	批量变体与编排	高专业制作	营销运营负责人
对外发布	批准高风险事项	按授权发布低风险项	按 SLA 执行	品牌/渠道负责人
监测复盘	解释因果与决策	持续监测与归纳	独立复核	业务负责人

表 9 | 责任分配的六项评分

维度	更适合 AI	更适合人	更适合伙伴
速度/规模	高频、大批量	低频关键决策	峰值弹性
判断	规则清晰	价值冲突、模糊目标	专业领域判断
情境	结构化输入	组织关系、客户语境	跨行业比较
责任	无独立责任	承担最终责任	合同范围责任
信任	内部可验证任务	高信任客户互动	已建立专业信誉
学习	从反馈中稳定迭代	重新定义问题	引入外部方法

案例二 | P&G 与哈佛商学院“赛博队友”现场实验

案例卷宗

要素	内容
主体与时间	Harvard Business School / P&G, 工作论文 2025 年更新
业务情境	真实产品创新任务, 跨职能协作
具体动作	791 名专业人员在个人、团队、AI 辅助条件下完成任务
结果证据	AI 辅助个体达到传统团队部分表现, 并提高顶尖方案概率
归因限制	特定任务、有限周期, 不能外推长期组织运行
管理启示	AI 可充当认知队友, 但目标与责任仍由人承担

2024–2025 年, 研究团队与 P&G 开展现场实验, 覆盖 791 名专业人员, 比较个人、团队以及使用生成式 AI 的不同工作方式。任务围绕真实产品创新问题, 评价包含方案质量、跨职能知识与参与体验。

研究显示, 使用 AI 的个人能够达到传统双人团队的部分表现, AI 也帮助不同职能人员更容易跨越专业边界, 并提高产生顶尖创意的概率。机制不是 AI 单独替代团队, 而是让个体获得类似“队友”的信息补充与反馈。

边界同样重要: 实验针对特定创新任务, 持续时间有限; 它不能证明长期团队关系、执行责任与组织学习可以被替代。管理启示是把 AI 视作可配置的认知协作者, 并保留人对问题定义、组合判断与后续执行的所有权。

第四章 授权梯度：把自主性建立在证据与可逆性上

竹势 AI 营销智库出品

L0

观察

只读与提示

L1

建议

人决定与执行

L2

草拟

人审后发布

L3

受限执行

额度内自动运行

L4

条件自治

持续监测与可接管

风险、可逆性、敏感数据与证据水平共同决定授权上限。

授权深度必须同时考虑风险、可逆性、数据敏感度与证据成熟度。NIST AI RMF 以 Govern、Map、Measure、Manage 四项功能组织风险管理，强调责任、情境、测量与持续管理。ISO/IEC 42001 要求组织建立、实施、维护并持续改进 AI 管理体系。两者都指向同一管理原则：授权不是一次配置，而是可证明、可监测、可纠正的管理过程。

欧盟 AI Act 对高风险系统强调人类监督、技术稳健、日志与可控性；中国《生成式人工智能服务管理暂行办法》以及个人信息、数据安全与内容标识规则要求提供和使用活动遵守安全、个人信息与内容治理义务。企业内部营销应用是否直接落入特定监管义务，要按服务对象、数据、行业与部署方式判断，不能把通用框架替代法律意见。

可撤销性是授权的关键工程变量。一个可以在沙箱中预演、限制预算、限定受众、保留草稿、设置停机开关并完整记录日志的动作，可以比同等业务影响但不可撤回的动作获得更深授权。

表 10 | 五级授权梯度

级别	AI 权限	人工控制	营销示例
L0 禁止	不得处理或调用	人工完成	高度敏感个人信息、未获授权数据
L1 建议	分析与建议	人从零决策	品牌战略、危机判断
L2 草拟	生成草稿/方案	逐项审核后发布	文章、邮件、活动方案
L3 受控执行	在白名单与额度内执行	关键节点审批+日志	定时发布、低额测试
L4 条件自治	满足条件时自主处置	抽样、异常升级、可停机	低风险客服、规则化优化

表 11 | 授权与可逆性矩阵

后果/可逆性	高可逆	中可逆	低可逆
低后果	L3-L4	L2-L3	L2
中后果	L2-L3	L2	L1-L2
高后果	L1-L2	L1	L0-L1

表 12 | 授权升级证据门

证据门	最低证据	未通过时动作
质量	代表性样本达到阈值	维持当前级别
安全	权限、日志、隔离测试通过	禁止执行权限
可逆	暂停、回滚演练成功	仅保留草拟
运营	异常升级时限达标	缩小任务范围
责任	具名负责人接受责任	不得上线

案例三 | 平安集团：高频服务的规模化自动化与责任边界

案例卷宗

要素	内容
主体与时间	中国平安，2024 年及 2025 年上半年官方披露
业务情境	保险与综合金融客户服务、核保、理赔等高频流程
具体动作	AI 服务代表进入标准化服务，人类运营与风险体系保留控制
结果证据	2024 年约 18.4 亿次 AI 服务，占客户服务总量约 80%
归因限制	企业自报；未披露完整异常率、成本与岗位结构
管理启示	规模化必须建立在流程标准、监测和人工升级基础上

平安 2024 年可持续发展披露显示，AI 服务代表提供约 18.4 亿次服务，占集团客户服务总量约 80%；官方披露还给出寿险秒级核保等业务指标。2025 年上半年，AI 服务代表约 8.82 亿次，仍占客户服务总量约 80%。

其组织机制可以从披露中看到三点：把 AI 放入高频、规则化、可监测的服务流程；在人类运营中心、业务规则与风险体系下运行；用服务量、时效与风险指标持续监测。公开材料没有披露完整的人机岗位、异常率和单项 ROI，因此不能据此推断所有服务均由 AI 独立完成。

适用边界是大型金融集团拥有长期数据、流程和风险基础。中小企业不能复制规模，只能借鉴“先选择规则明确的高频任务、保留升级与人工服务、用运行数据逐步扩大授权”的路径。

BOARD BRIEF 07

第五章 交接协议：管理上下文、状态、证据与责任

竹势 AI 营销智库出品

HANDOFF PACKET

任务状态

已用证据

未决假设

异常与限制

建议下一步

接收责任人

→

接收确认

完整 · 可理解 · 可执行 · 可追溯

交接是人机系统的接口。航空和医疗的交接实践强调结构化信息、接收确认和未决风险，因为事故往往不是单点能力不足，而是信息在边界处丢失。营销工作同样存在交接风险：研究结论没有证据链，创意稿没有说明假设，AI 执行后没有返回状态，代理商提交物没有标记未解决问题。

有效交接必须让接收者能够继续工作，而不必重建全过程。交接包至少包含任务目标、当前状态、已完成动作、证据与版本、质量检查、未决项、风险、下一责任人和截止时间。接收者必须确认“收到且可继续”，否则交接不成立。

交接协议不应要求输出完整思维过程。它需要的是可审计的依据、引用、输入版本、规则命中、关键选择与异常，而不是不可验证的内部推理叙述。

表 13 | 人机交接协议

字段	交出方必须提供	接收方必须确认
目标	上位经营目标与本次任务范围	目标无歧义
状态	完成、待审、阻塞、失败	状态可识别
证据	来源、数据版本、测试结果	证据可访问
输出	文件、结构、渠道与版本	格式可继续处理
未决项	尚未解决的问题与影响	接管或升级
风险	合规、品牌、数据、时效风险	接受或拒绝
下一责任	具名责任人与截止时间	接收确认

表 14 | 完成定义 DoD

任务	最低完成定义
研究摘要	每个关键结论有来源、日期、口径与不确定性
内容草稿	事实核查、品牌语气、禁区、目标受众、CTA 完成
渠道发布	审批通过、账号正确、时间与回滚路径明确
线索跟进	客户同意、阶段、下一动作、负责人、记录完整
复盘	基线、结果、偏差、原因假设、下一轮动作完整

案例四 | Klarna：自动客服扩张后的品牌与服务再平衡

案例卷宗

要素	内容
主体与时间	Klarna，2024–2025 年公开披露
业务情境	全球消费者金融客服
具体动作	AI 助手承担大量文本客服，后续增加人工服务重心
结果证据	企业自报首月 230 万次对话、约三分之二聊天
归因限制	供应商与企业自报；长期客户价值口径有限
管理启示	授权应随质量证据回调，复杂事项必须顺畅转人工

Klarna 在 2024 年公开宣布 AI 助手在上线首月处理 230 万次对话，覆盖约三分之二客服聊天，并宣称平均处理时间和重复咨询下降。该披露来自企业与技术供应商，说明自动化可以快速吸收高频请求。

随后公司管理层在 2025 年公开强调需要重新增加人工客服投入，并承认过度强调成本可能影响服务质量。这个案例的价值不在于判断 AI 成功或失败，而在于展示授权边界会随客户体验证据变化而回调。

管理启示是：自动化比例不能成为目标本身。高频简单问题可由 AI 处理，复杂、情绪化或高价值客户事项要有清晰交接和人工接管；同时监测重复联系、投诉、升级率与客户信任，而不只看处理量。

BOARD BRIEF 08

第六章 审批与例外：在控制强度与运行速度之间取平衡

竹势 AI 营销智库出品



审批链的目标是拦截高后果决策，而不是让每一个动作都排队。全量人工复核会形成新的瓶颈，并诱发“点击通过”式形式控制。更有效的方式是按风险分层：高风险逐项批准，中风险规则校验加抽样，低风险自动执行但持续监测，异常信号触发升级。

Weick 关于高可靠组织的研究强调对失败保持敏感、避免过度简化、贴近运行现场、保持恢复能力和尊重专业知识。转译到 AI 市场部，就是不把“系统正常运行”当作安全证明；一线员工可以暂停；异常按专业性而非行政级别升级；复盘关注近失误而不仅是事故。

例外机制必须定义等级、时限、接管人和恢复条件。没有恢复条件的暂停会变成永久停摆；没有暂停权的员工会在风险可见时继续运行。

表 15 | 风险审批矩阵

风险级别	典型事项	控制方式	升级时限
R1 低	内部摘要、公开资料整理	自动+月度抽样	下一工作日
R2 中	普通内容草稿、低额 A/B 测试	规则校验+抽样	4 小时
R3 高	对外声明、敏感客户触达	逐项人工审批	1 小时
R4 重大	危机回应、价格承诺、大额预算	双人批准+法务/安全	立即

表 16 | 异常分级与处置

异常	触发条件	即时动作	恢复条件
数据异常	来源缺失、版本冲突	暂停输出	数据所有者确认
品牌异常	语气、禁区、事实偏差	撤回草稿/停止发布	品牌负责人复核
权限异常	越权调用或访问	熔断并保留日志	安全负责人解除
业务异常	转化/投诉显著偏离	降级授权	业务负责人批准
合规异常	个人信息、广告、行业规则风险	立即停止并升级	法务书面确认

表 17 | 抽样审查设计

任务稳定性	建议抽样	额外触发
新上线	20%–100%	连续错误即全量审查
稳定运行	5%–10%	模型/数据变更提高抽样
低频高后果	100%	不得仅抽样
高频低后果	1%–5%	异常信号自动扩大

第七章 组织形态：五种结构及其迁移条件

竹势 AI 营销智库出品

01

中心化 CoE

标准与稀缺专家集中

02

业务嵌入式

速度与情境优先

03

联邦式

中心规则、业务执行

04

平台型

共享底座、自助组合

05

轻量网络型

小团队与伙伴弹性协同

组织形态应匹配业务差异、风险、人才、平台成熟度和管理跨度。中心化 CoE 有利于统一标准、采购和风险，但可能远离业务；业务嵌入式响应快，却容易重复建设；联邦式在中心标准与业务自治间平衡，但需要成熟接口；平台型把数据、工具、日志与治理作为共享产品；轻量网络型适合中小企业，用少量核心人员和外部伙伴围绕关键任务运行。

Galbraith 的组织信息处理观点认为，不确定性越高，组织越需要增加信息处理能力或减少信息需求。AI 可以增加信息处理能力，但也会制造更多输出和例外。组织形态选择的目标不是拥有更多 AI 团队，而是让信息、决策和责任以合理成本流动。

组织形态不是身份标签，而是阶段选择。企业可以从轻量网络或中心化试点开始，随着重复需求、共享数据和治理能力增加，迁移到联邦式或平台型。若平台团队无法提供可靠服务，业务单元会绕开平台，形成新的影子系统。

表 18 | 五种组织形态比较

形态	优势	主要风险	适用条件
中心化 CoE	标准统一、治理集中	离业务远、排队	高监管、早期建规
业务嵌入式	贴近场景、响应快	重复建设、标准分裂	业务差异大
联邦式	中心治理+业务自治	接口复杂	规模企业、治理成熟
平台型	能力复用、可观测	平台成为瓶颈	技术与产品能力较强
轻量网络型	低成本、灵活	依赖关键人和伙伴	中小企业、场景聚焦

表 19 | 组织形态选择器

问题	是	否
多个业务单元是否高度不同？	偏嵌入式/联邦式	偏中心化
是否存在统一高风险要求？	增强中心治理	可轻量
是否有稳定平台团队？	考虑平台型	避免先建大平台
关键人才是否稀缺？	建立共享 CoE	可业务自建
企业规模是否较小？	轻量网络型	评估联邦式

案例五 | WPP 与 NVIDIA：外部生态进入营销生产平台

案例卷宗

要素	内容
主体与时间	WPP 与 NVIDIA，2023 年合作披露
业务情境	全球广告内容生产与三维资产
具体动作	构建连接创意、制作和生成能力的内容引擎
结果证据	公开披露平台与合作范围；统一 ROI 未披露
归因限制	供应商合作声明，缺少跨客户独立评估
管理启示	平台能力、代理专业与客户责任必须分离并接口化

WPP 与 NVIDIA 在 2023 年宣布合作建设面向广告内容的生成式 AI 内容引擎，把设计、制作、三维资产与生成能力连接起来。此类合作体现代理集团、技术平台和客户品牌之间的新型任务分工。

公开材料说明了平台愿景与技术组合，但没有披露统一、可比较的客户 ROI，也不能证明所有创意任务适合自动化。组织机制上的关键是：代理商保留品牌与创意专业判断，技术供应商提供底层能力，客户继续拥有品牌责任与数据治理。

管理启示是，外部生态应通过共享资产标准、审批、权限、日志和知识转移进入企业系统，而不是让代理商在不可见的外部环境中独立生成并交付最终结果。

第八章 外部伙伴：把代理商与供应商纳入同一责任系统

竹势 AI 营销智库出品

任务接口
输入、交付物、完成定义
数据接口
权限、使用、留存、删除
知识接口
资产归属、文档、转移
责任接口
SLA、事故、审计、退出

Coase 与 Williamson 的交易成本理论提醒企业，外包与内制的边界取决于价格之外的搜寻、协调、监督、机会主义和退出成本。AI 降低部分生产成本，却可能提高数据治理、版本管理、知识产权和供应商锁定成本。

代理商、咨询和技术供应商必须被纳入任务契约：明确允许访问的数据、可使用的模型与分包方、交付格式、验收标准、日志保留、事故通知、知识产权、知识转移、服务水平和退出。企业不能把最终责任外包给“AI”或供应商。

外部伙伴最适合补充稀缺专业、短期峰值与跨市场经验；核心客户关系、品牌原则、关键数据、风险判断和能力架构应保留在企业内部。

表 20 | 外部伙伴任务契约附加条款

条款	必须写清
数据	范围、目的、存储、跨境、删除
模型与工具	允许清单、版本、分包
验收	事实、品牌、格式、性能阈值
SLA	响应、修复、升级、停机
审计	日志、抽样、证据访问
知识产权	输入、输出、训练使用、第三方权利
知识转移	文档、培训、可移交资产
退出	数据返还、删除证明、替代方案

表 21 | 内制、外包与混合决策

任务特征	内制	外包	混合
核心差异化	优先	谨慎	内部主导
专业稀缺	成本高	优先	共同交付
高峰产能	不经济	优先	弹性扩容
敏感数据	优先	限制	隔离环境
标准化执行	可平台化	可 SLA	平台+伙伴

案例六 | 可口可乐与 WPP Open X：把代理关系改造成统一网络

案例卷宗

要素	内容
主体与时间	The Coca-Cola Company 与 WPP，2021 年起
业务情境	全球品牌的多代理、多渠道营销
具体动作	建立 Open X 统一网络，整合创意、媒体、技术与数据
结果证据	官方披露合作结构；AI 单项 ROI 未披露
归因限制	缺少独立量化评估
管理启示	伙伴整合的核心是接口、治理与知识回流

可口可乐在 2021 年选择 WPP 建立 Open X，整合创意、媒体与营销技术能力，目标是减少碎片化并建立端到端体验。随后双方把生成式 AI、内容与数据能力纳入合作。

这不是单一 AI 项目，而是外部伙伴组织设计案例：品牌方通过统一的伙伴网络、共同流程与治理降低多代理协调成本，同时保留品牌目标和最终批准。公开资料没有提供可分离的 AI 投资回报。

管理启示是，企业引入 AI 代理商时，不应仅采购单个内容包；应明确统一 Brief、资产标准、数据接口、评审机制、学习回流与退出条件，让外部专业成为组织能力的一部分。

BOARD BRIEF 11

第九章 绩效与责任：共同结果、个人责任与系统质量

竹势 AI 营销智库出品

经营结果
团队共同承担

质量

速度

采用

学习

风险

可审计

AI 输出数量不计作绩效；远端收入与早期试点证据分层使用。

只看产量会诱发低质自动化：内容数量、线索触达数、自动回复量都可能上升，但客户价值、品牌质量和风险暴露恶化。Deming 的系统质量思想指出，大部分质量问题来自系统，而不是个体。人机团队绩效必须同时衡量经营结果、质量、速度、采用、学习、风险和可审计性。

共同绩效不等于责任平均化。AI 没有法律人格，不能承担管理责任。外部伙伴承担合同责任，员工承担岗位责任，具名管理者承担最终业务决策责任。发生时，组织要追踪任务设计、数据、权限、审核、系统版本与管理决策，而不是把错误归咎于某个员工“没有看仔细”。

绩效指标应防止 Goodhart 定律：一旦单一指标成为目标，它往往失去衡量价值。处理量必须与一次解决率、客户满意、投诉与升级率结合；内容产量必须与有效触达、品牌一致性和转化结合。

表 22 | 团队共同绩效模型

维度	示例指标	防止的失真
经营结果	收入贡献、合格线索、客户价值	只做活动不见结果
质量	事实准确、品牌一致、一次通过率	低质批量生产
速度	周期时间、响应时限	用加班掩盖流程问题
采用	关键任务覆盖、有效使用	账号开通即算采用
学习	规则改进、复盘转化、技能增长	重复犯错
风险	异常、投诉、越权、隐私事件	把风险外部化
可审计	日志完整、责任清晰、版本可追	无法复盘

表 23 | 责任模型 RACI-AI

角色	负责什么	不负责什么
业务负责人	目标、取舍、最终结果	把责任交给系统
专业审核人	事实、品牌、专业判断	替代业务目标决策
AI 系统所有者	性能、权限、监测、版本	承担业务法律责任
数据所有者	数据质量与授权	输出内容最终批准
外部伙伴	合同范围与专业交付	客户企业最终责任
法务/安全	规则与高风险审查	日常业务经营

案例七 | 淘宝天猫商家 AI 工具：规模使用与平台自报边界

案例卷宗

要素	内容
主体与时间	淘宝天猫，2023–2025 年官方披露
业务情境	大量商家的内容、客服与经营分析
具体动作	向商家开放十项 AI 工具并持续升级
结果证据	2023 双 11 使用超过 15 亿次；其他效果为平台自报
归因限制	使用量不等于增量；存在选择偏差
管理启示	平台降低门槛，但企业仍需自有验收与责任机制

淘宝天猫在 2023 年 10 月向部分商家推出十项 AI 工具，并在 2024 年初扩大开放，覆盖营销内容、客服与经营分析。阿里官方披露这些工具在 2023 年双 11 期间被使用超过 15 亿次；2025 年又披露商家经营分析和客服工具的使用结果。

该案例显示平台可以把生成、分析与客服嵌入商家日常任务，降低小企业采用门槛。使用次数说明规模，不等于经营增量；后续转化率与效率数字来自平台自报，受商家选择、活动环境和其他变量影响。

管理启示是：中小企业可以借助平台化能力快速进入 L2–L3 授权，但必须把平台数据、品牌审核、客户同意和经营结果纳入自己的绩效模型，不能把平台使用量当作成功。

BOARD BRIEF 12

第十章 能力与职业：岗位任务迁移与人的稀缺价值

竹势 AI 营销智库出品

保留在人

问题定义 · 权衡 · 关系 · 伦理 · 最终责任

人机共作

洞察 · 策略 · 创意 · 复盘 · 系统改进

逐步委派

检索 · 适配 · 编排 · 监测 · 记录

保留基础技能、真实任务训练与定期脱机演练，降低自动化偏差和失能。

人的稀缺价值会向问题定义、判断、品味、关系、伦理和系统改进移动，但这不意味着基础技能可以放弃。自动化偏差会让员工过度相信系统；长期不练习会导致在系统失效时无法接管。能力设计要同时提高 AI 协作能力和独立完成关键任务的能力。

劳动经济研究显示收益分布不均。客服研究中，新手受益最大；生成式 AI 的总体工作采用率也随职业、教育和年龄显著不同。组织若只奖励早期熟练用户，可能扩大内部差距。CHRO 需要把培训、岗位迁移、绩效和心理安全放在同一方案中。

Argyris 的双环学习区分“纠正动作”与“反思目标和规则”。AI 可以帮助第一类学习，通过日志和反馈修正模板；人必须主导第二类学习，判断是否追错了指标、是否需要改变客户策略或授权边界。

表 24 | 岗位任务迁移地图

岗位	减少	增强	新增责任
品牌经理	重复整理与初稿	定位、品味、跨部门取舍	品牌规则与审核设计
内容人员	从零起草与多平台搬运	选题、编辑、叙事	训练素材与质量评测
投放人员	手工拉数与常规调整	实验设计、因果判断	授权额度与异常复盘
销售运营	记录与提醒	关系判断、商机策略	人机交接与数据质量
营销分析	重复报表	解释、预测、决策支持	指标治理与证据链
管理者	逐项催办	目标、边界、系统改进	任务契约所有权

表 25 | 能力发展路径

层级	能力目标	验证方式
基础	理解局限、数据与安全规则	情景测验
任务	能设计输入、验收与升级	完成任务契约
专业	能判断质量与品牌适配	盲评与案例复盘
系统	能改进流程与指标	周期时间/质量改善
治理	能设授权、审计与退出	演练与审计

案例八 | 招商银行：高监管行业的 AI 客服与数字产品边界

案例卷宗

要素	内容
主体与时间	招商银行，2025 年可持续发展披露
业务情境	银行客户服务与数字产品
具体动作	整合 AI 智能客服等技术
结果证据	官方披露持续建设；详细 ROI 未披露
归因限制	缺少完整组织和运行指标
管理启示	信息服务可自动化，高后果决定保留具名人工责任

招商银行 2025 年可持续发展报告披露，公司把 AI 技术整合进智能客服等数字产品。银行场景涉及个人信息、金融消费者权益与高后果决策，组织上不能只追求自动化比例。

公开披露能够证明其持续建设 AI 客服与数字能力，但未公开完整的人机岗位分工、异常升级率、模型版本和单项 ROI。因此本案例的价值是边界设计：常见问答、流程引导可以自动化；产品适当性、投诉、争议和重大客户事项必须由具备资质和责任的人员接管。

管理启示是，高监管企业要把数据目的限制、身份验证、记录、投诉、人工渠道与审计前置到任务契约。技术能力越强，越需要清晰区分“提供信息”与“作出受监管决定”。

BOARD BRIEF 13

第十一章 运行节奏：把协同变成可持续的管理循环

竹势 AI 营销智库出品

DAILY	日常编排
	任务、负载、接管
WEEKLY	例外复盘
	失败、返工、边界
MONTHLY	能力评测
	质量、成本、风险
QUARTERLY	边界重设
	角色、授权、组织形态

人机协同不是上线项目，而是运行纪律。日常任务编排保证工作不断链；周度例外复盘发现异常与近失误；月度能力评测检验质量、模型、数据与员工熟练度；季度边界重设决定扩大、收缩或退出授权。

运行节奏必须记录任务、版本、输入、输出、审批、异常和责任人。日志不是为了无限保留所有信息，而是为了支持复盘、审计和纠错。个人信息和敏感数据的保留期限应符合法律与企业政策。

稳定、低频、低风险场景可以降低复评频率；模型更新、数据源变化、重大人员调整、投诉上升或监管变化必须触发临时复评。

表 26 | 日 / 周 / 月 / 季度运行节奏

周期	会议/机制	输入	输出	负责人
每日	任务编排与异常站会	待办、状态、阻塞	优先级与接管	营销运营
每周	例外与近失误复盘	异常、抽样、投诉	修复、规则调整	业务+风险
每月	能力与质量评测	样本、指标、版本	授权建议、培训	系统所有者
每季度	边界重设	经营结果、风险、能力	扩大/收缩/退出	董事会/经营层

表 27 | 最小运行日志

字段	用途
任务 ID 与契约版本	定位责任与规则
模型/工具/数据版本	解释性能变化
输入与权限记录	验证合法与最小必要
输出与验收结果	质量审计
审批与接管	证明人工控制
异常与回滚	恢复与学习
责任人与时间	追踪闭环

案例九 | 微软内部与全球调查：个人采用快于组织对齐

案例卷宗

要素	内容
主体与时间	Microsoft / LinkedIn, 2024 与 2026 年工作趋势调查
业务情境	全球知识工作者 AI 使用
具体动作	跨国调查与产品遥测结合
结果证据	2024 年 31 国 31,000 人；75% 使用 AI, 78% 自带工具
归因限制	调查与厂商研究，不能证明企业财务因果
管理启示	个人采用需要组织目标、培训与治理同步

微软与 LinkedIn 2024 年调查显示，全球 75% 的知识工作者使用 AI，78% 的用户把自己的 AI 工具带入工作场景；59% 的领导担心难以量化生产率，60% 担心组织缺乏实施愿景。2026 年工作趋势报告又指出，仅约四分之一 AI 使用者认为领导层在 AI 上清晰一致。

这些数字来自企业研究机构的跨国调查，反映认知与采用，不直接证明财务结果。它们揭示了组织设计问题：个人自发采用可以快速发生，但目标、规则、培训、平台和责任若不同步，组织难以把个人节省转化为稳定结果。

管理启示是建立固定运行节奏，让一线使用、管理目标、风险规则和能力建设持续对齐，而不是依靠一次培训或年度政策。

第十二章 90 天组织导入路线

竹势 AI 营销智库出品

0-30 设计任务契约 选闭环 · 定基线 · 划权限 · 写退出
31-60 运行协作单元 真实交接 · 异常注入 · 人工接管
61-90 审议扩展条件 复核证据 · 调整形态 · 扩展或退出

90 天导入的目标不是完成全面改造，而是建立一个能经受审查的协作单元。前 30 天定义任务契约和风险边界；第 31-60 天在真实业务中运行一个单元；第 61-90 天根据证据决定扩展、重构或退出。

每一阶段都必须有终止条件。若无法取得合法数据、无法明确最终责任、质量达不到基线、异常无法在时限内接管，项目应暂停或缩小范围。受监管、劳动关系复杂、跨境数据或多代理环境可能需要更长周期。

导入路线应由经营负责人主导，技术、HR、法务和安全共同参与。把项目交给单一创新团队会导致业务不拥有结果；把全部责任交给 IT 会让任务设计与客户价值脱节。

表 28 | 90 天组织导入路线

阶段	核心任务	证据门	终止条件	交付物
0-30 天 设计	选场景、基线、任务契约、授权与责任	目标/数据/责任清晰	数据不合法、无人负责	契约画布、风险矩阵
31-60 天 运行	真实任务、交接、审批、日志、培训	质量不低于基线，异常可接管	持续返工、无法回滚	运行单元、周报、异常记录
61-90 天 决策	评估经营、质量、风险与学习	达到扩展阈值	无增量或风险不可接受	扩展/重构/退出决议

表 29 | 阶段角色分工

角色	30 天	60 天	90 天
CEO/业务负责人	确定经营任务	解决跨部门阻塞	批准扩展或退出
CMO	设计流程与验收	运行协作单元	评估经营结果
CHRO	岗位与能力基线	培训与心理安全	岗位迁移决策
CIO/CDO	平台、数据、日志	稳定性与版本	平台扩展
法务/安全	风险与权限	异常与审计	边界重设
伙伴负责人	合同与接口	SLA 与知识转移	续约或退出

表 30 | 扩展、重构、退出决策

决策	触发条件	动作
扩展	质量达标、经营增量、风险受控	增加任务或提高授权
重构	有价值但返工/接口问题突出	重做契约与交接
降级	异常上升或环境变化	降低授权、提高抽样
退出	无增量、风险不可控、依赖不可接受	停止并迁移资产

管理工具总览与使用说明

竹势 AI 营销智库出品

表 31 | 十项原创管理工具的衔接

工具	主要使用者	输入	输出	不适用情形
任务契约画布	业务负责人/运营	目标、流程、数据	七要素契约	纯个人低风险构思可简化
责任谱	CMO/CHRO	任务阶段、能力	人/AI/伙伴分工	无明确任务时不可使用
授权梯度矩阵	CIO/法务/业务	风险、可逆、敏感度	授权级别与控制	不能替代行业法律判断
人机交接协议	运营/审核人	状态、证据、未决项	可继续工作的交接包	一次性无下游任务可简化
审批与例外矩阵	法务/安全/运营	风险分级、时限	审批、抽样、升级	低风险内部任务可轻量
组织形态选择器	CEO/CHRO	规模、差异、风险、人才	结构与迁移条件	不能脱离战略单独决定
共同绩效模型	CEO/CMO/HR	基线、指标、责任	平衡绩效卡	数据不可信时先修数据
岗位迁移地图	CHRO/职能负责人	任务清单、技能基线	减少/增强/新增任务	不等于裁员名单
运行节奏	运营/风险	日志、指标、异常	日周月季机制	低频稳定场景可降频
90 天路线	经营层	目标、场景、资源	扩展/重构/退出证据	高监管项目需延长

BOARD BRIEF 16

董事会与职能负责人审议清单

竹势 AI 营销智库出品

表 32 | CEO 审议清单

问题	通过标准
是否有明确经营任务？	目标、基线、负责人、停止条件齐备
是否改变责任而非只加工具？	任务契约与最终责任清晰
是否能回滚？	暂停、回滚和人工接管已演练
是否可扩展？	能力、平台、伙伴和成本边界清晰

表 33 | CMO 审议清单

问题	通过标准
任务是否端到端？	从目标到复盘无断点
质量如何验收？	事实、品牌、客户、渠道标准明确
是否制造新瓶颈？	审批与返工周期可接受
是否沉淀学习？	复盘能更新规则和资产

表 34 | CHRO 审议清单

问题	通过标准
岗位变化是否透明？	减少、增强、新增任务明确
员工能否安全质疑？	有暂停权和心理安全
基础技能是否保留？	定期手工演练与接管
收益是否公平？	培训覆盖与差异监测

表 35 | CIO/CDO 审议清单

问题	通过标准
数据是否合法最小必要？	目的、范围、保留和访问清晰
系统是否可观测？	日志、版本、指标、告警齐备
接口是否稳定？	输入输出与错误状态标准化
是否可退出？	数据与能力可迁移

表 36 | 法务与安全负责人审议清单

问题	通过标准
责任主体是否明确？	AI 不承担责任，具名人员负责
高风险是否逐项控制？	风险分级与审批匹配
个人信息与内容规则是否满足？	同意、目的、标识、投诉机制明确
事故是否可处置？	通知、暂停、回滚、证据保全

表 37 | 伙伴管理者审议清单

问题	通过标准
伙伴是否进入任务契约?	数据、权限、验收、SLA 清晰
是否完成知识转移?	文档、培训、资产可接管
是否存在锁定?	替代方案与退出条件明确
最终责任是否留在企业?	企业具名负责人批准关键事项

附录 A | 经典思想的当代转译

竹势 AI 营销智库出品

表 38 | 七项经典思想与组织设计含义

思想	原始问题	当代转译	边界
Drucker：贡献与责任	知识工作者如何对成果负责	任务契约从贡献结果而非活动定义工作	不能忽视团队互赖
Simon：有限理性	决策受信息与注意力限制	AI 扩展搜索，人仍处理目标冲突	更多信息不等于更好决策
Licklider：人机共生	人与计算机如何互补	把 AI 设计为协作伙伴与执行组件	不等于赋予人格或责任
Galbraith：信息处理	不确定性如何影响组织	用平台、接口和横向关系处理更多信息	AI 也会增加信息噪音
Weick：高可靠组织	如何在复杂系统避免事故	近失误、暂停权、专业升级与恢复力	不要求所有场景同等重控
Coase/Williamson：交易成本	何时内制或外包	比较协调、监督、锁定与退出成本	不能只比较软件价格
Deming：系统质量	质量为何是系统结果	绩效追踪流程、数据与管理而非只责个人	个人专业责任仍需保留

附录 B | 当代专家与机构观点

竹势 AI 营销智库出品

表 39 | 具名观点及适用边界

专家/机构	实质观点	本报告转译	适用边界
Erik Brynjolfsson、Danielle Li、Lindsey Raymond	AI 对客服生产率影响高度异质，新手收益大	按人员与任务设边界，不能承诺平均收益	文本客服场景
Karim Lakhani、Raffaella Sadun 等	AI 可改变团队知识边界与创意表现	AI 可作为认知队友，责任仍在人	特定创新实验
David Autor	AI 可扩展中等技能劳动者的能力与任务范围	能力建设不应只围绕替代	宏观与劳动市场推论
Ethan Mollick	管理者需要在真实任务中实验并建立组织学习	小范围真实运行优于工具演示	观点需结合企业风险
NIST	以 Govern/Map/Measure/Manage 管理 AI 风险	授权必须有治理、测量与持续管理	自愿框架，不替代法律
ISO/IEC 42001	建立并持续改进 AI 管理体系	把人机协同纳入管理系统	认证不保证单一输出正确
欧盟立法机构	高风险 AI 需要人类监督、日志、稳健与透明	高后果营销动作提高控制强度	具体适用需法律判断
中国网信等主管部门	生成式 AI 服务需兼顾发展、安全、个人信息与内容治理	数据、内容、投诉与处置进入任务契约	内部应用适用性按场景判断
Microsoft / LinkedIn	个人采用快于组织愿景与培训	建立组织对齐与运行节奏	厂商调查，非财务因果
Amy Edmondson	心理安全使员工能够报告错误与提出异议	赋予一线暂停与升级权	需要与问责并存

附录 C | 纯文本来源索引

竹势 AI 营销智库出品

- [1] NIST, Artificial Intelligence Risk Management Framework (AI RMF 1.0), 2023 年 1 月。
- [2] NIST, Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile, 2024 年 7 月。
- [3] ISO, ISO/IEC 42001:2023 Artificial intelligence management systems, 2023 年。
- [4] ISO/IEC, ISO/IEC FDIS 42105 Guidance for human oversight of AI systems, 2026 年最终草案阶段。
- [5] European Union, Regulation (EU) 2024/1689 (Artificial Intelligence Act), 2024 年 6 月。
- [6] 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 7 月。
- [7] 全国网络安全标准化技术委员会, 生成式人工智能服务安全基本要求, 2024 年 2 月。
- [8] 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。
- [9] 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。
- [10] Brynjolfsson, Erik; Li, Danielle; Raymond, Lindsey R., Generative AI at Work, NBER Working Paper 31161, 2023 年; 5,179 名客服人员。
- [11] Dell'Acqua, Fabrizio 等, The Cybernetic Teammate: A Field Experiment on Generative AI Reshaping Teamwork and Expertise, Harvard Business School Working Paper 25-043, 2025 年; 791 名专业人员。
- [12] Dell'Acqua, Fabrizio 等, Navigating the Jagged Technological Frontier, Harvard Business School Working Paper, 2023 年; 咨询任务实验。
- [13] Bick, Alexander; Blandin, Adam; Deming, David, The Rapid Adoption of Generative AI, NBER Working Paper 32966, 2024 年。
- [14] Autor, David, Applying AI to Rebuild Middle Class Jobs, NBER Working Paper 32140, 2024 年。
- [15] Microsoft and LinkedIn, 2024 Work Trend Index Annual Report, 2024 年 5 月; 31 个国家、31,000 人。
- [16] Microsoft, 2026 Work Trend Index: Agents, human agency, and the opportunity for every organization, 2026 年 5 月。

- [17] Ping An Insurance (Group) Company of China, 2024 Sustainability Report, 2025 年 3 月发布。
- [18] Ping An Insurance (Group) Company of China, 2025 Interim Results, 2025 年 8 月。
- [19] Ping An Insurance (Group) Company of China, 2025 Annual Report, 2026 年 3 月。
- [20] China Merchants Bank, 2025 Sustainability Report, 2026 年 3 月。
- [21] Alibaba Group, Taobao and Tmall Upgrades Consumer Shopping Experience with AI, 2024 年 6 月。
- [22] Alibaba Group, AI Powers Large-scale Applications at the World's Largest Shopping Festival, 2025 年 10 月。
- [23] Alibaba Group, Tmall Records Strongest 11.11 Growth in Four Years, 2025 年 11 月；平台自报口径。
- [24] Klarna, AI assistant handles two-thirds of customer service chats in its first month, 2024 年 2 月；企业与供应商披露。
- [25] Klarna 管理层关于增加人工客服投入的公开访谈与公司声明, 2025 年。
- [26] WPP and NVIDIA, 合作建设面向广告内容的生成式 AI 内容引擎, 2023 年 5 月官方发布。
- [27] The Coca-Cola Company and WPP, Open X 全球营销网络合作披露, 2021 年及后续更新。
- [28] Drucker, Peter F., The Effective Executive, 1967 年。
- [29] Simon, Herbert A., Administrative Behavior, 1947 年及后续版本。
- [30] Licklider, J. C. R., Man-Computer Symbiosis, 1960 年。
- [31] Galbraith, Jay R., Organization Design, 1977 年。
- [32] Weick, Karl E.; Sutcliffe, Kathleen M., Managing the Unexpected, 2001 年及后续版本。
- [33] Coase, Ronald H., The Nature of the Firm, 1937 年。
- [34] Williamson, Oliver E., The Economic Institutions of Capitalism, 1985 年。
- [35] Deming, W. Edwards, Out of the Crisis, 1986 年。
- [36] Polanyi, Michael, The Tacit Dimension, 1966 年。
- [37] Argyris, Chris; Schön, Donald, Organizational Learning, 1978 年。
- [38] Edmondson, Amy C., The Fearless Organization, 2018 年。

结语：AI 市场部的组织能力不来自“拥有多少 AI 员工”，而来自企业能否把经营目标转化为可执行任务，把权限与风险匹配，把交接与责任写清，把异常纳入日常运行，并依据证据持续移动人机边界。

附录 D | 典型营销任务契约示例

竹势 AI 营销智库出品

以下示例用于把本报告的原则转化为可直接讨论的组织对象。示例不是通用标准答案，企业应根据行业监管、客户承诺、数据分类、渠道规则和人员能力修改。每份契约均要求具名业务负责人批准，任何技术或供应商配置都不能替代这一责任。

D.1 品牌内容发布任务契约

表 D-1 | 品牌内容发布任务契约

字段	约定
目标	在既定传播主题下形成一篇对外文章并按批准时间发布；经营目标是提高目标客户对核心能力的理解，而非单纯增加发布量。
输入	已批准的品牌定位、产品事实、目标受众、禁用表达、历史高质量内容和公开可引用资料。不得使用未授权客户资料、内部报价和未发布经营数据。
权限	AI 可检索白名单知识、生成结构和草稿、执行格式检查；不得自行新增事实、作出效果承诺或直接发布。编辑拥有改写权，品牌负责人拥有最终批准权。
输出	结构化草稿、事实清单、引用清单、风险标记、平台版本和发布建议。
验收	事实可追溯；品牌语气符合规范；广告与行业规则通过；没有敏感信息；CTA 与目标一致。
升级	出现事实冲突、监管敏感表述、客户可识别信息、重大负面舆情或模型无法确认的关键数字时，立即暂停并升级品牌负责人及法务。
责任	品牌负责人对最终发布负责；系统所有者对权限、日志和版本负责；编辑对专业复核负责。

使用方法：业务负责人先确认目标和责任，再由专业、数据、技术与风险角色共同补齐输入、权限和升级条件。试运行期间，每次异常都要回写契约；连续两轮稳定后才讨论提高授权。

不适用边界：若任务目标本身仍有重大争议，或数据取得缺乏合法基础，不应通过细化技术步骤来掩盖上游决策缺口。此时应回到经营决策或合规评估。

D.2 广告投放优化任务契约

表 D-2 | 广告投放优化任务契约

字段	约定
目标	在不改变总预算上限和品牌安全边界的前提下，发现低效单元并形成可执行优化建议。
输入	广告平台授权数据、归因口径、预算上限、历史实验、受众排除规则、品牌安全清单。
权限	AI 可读取数据、诊断异常、生成调整草稿和模拟影响；L2 阶段不得直接修改账户。L3 阶段仅可在单日限额和白名单活动内执行，且保留回滚。
输出	账户健康报告、异常证据、建议清单、预期影响区间、风险、草稿变更和回滚方案。
验收	数据完整；建议与证据对应；没有突破预算、受众和品牌限制；结果可在平台中验证。
升级	归因数据中断、支出异常、转化骤降、受众敏感或平台政策变化时，暂停自动执行并通知投放负责人。
责任	投放负责人对预算和业务结果负责；财务对预算规则负责；系统所有者对权限和日志负责。

使用方法：业务负责人先确认目标和责任，再由专业、数据、技术与风险角色共同补齐输入、权限和升级条件。试运行期间，每次异常都要回写契约；连续两轮稳定后才讨论提高授权。

不适用边界：若任务目标本身仍有重大争议，或数据取得缺乏合法基础，不应通过细化技术步骤来掩盖上游决策缺口。此时应回到经营决策或合规评估。

D.3 B2B 线索跟进任务契约

表 D-3 | B2B 线索跟进任务契约

字段	约定
目标	在客户授权和销售策略范围内，提高线索首次响应及时性与下一步动作完成率。
输入	CRM 阶段、客户同意状态、来源、行业、公开企业信息、已批准话术、销售负责人规则。
权限	AI 可生成个性化草稿、总结历史互动、提醒下一动作；不得虚构关系、承诺价格、替代销售作出合同或信用判断。低风险提醒可自动发送，首次外联与关键客户内容需人工批准。
输出	跟进草稿、客户摘要、下一动作、责任人、截止时间和 CRM 更新建议。
验收	客户身份与同意状态明确；信息准确；语气适当；下一动作具体；记录可追溯。
升级	客户投诉、退订、敏感个人信息、价格争议、合同问题或高价值客户情绪异常时，立即转人工。
责任	销售负责人对客户关系和承诺负责；营销运营对线索流程负责；数据负责人对同意和记录负责。

使用方法：业务负责人先确认目标和责任，再由专业、数据、技术与风险角色共同补齐输入、权限和升级条件。试运行期间，每次异常都要回写契约；连续两轮稳定后才讨论提高授权。

不适用边界：若任务目标本身仍有重大争议，或数据取得缺乏合法基础，不应通过细化技术步骤来掩盖上游决策缺口。此时应回到经营决策或合规评估。

D.4 舆情监测与危机升级任务契约

表 D-4 | 舆情监测与危机升级任务契约

字段	约定
目标	持续识别与企业、品牌和高管相关的重大信号，缩短从发现到人工判断的时间。
输入	公开媒体与平台信号、品牌实体词典、历史事件、风险分类、值班表和危机预案。
权限	AI 可抓取、聚类、去重、摘要和初步分级；不得自行公开回应、联系媒体或判断法律责任。
输出	事件卡片、来源、传播速度、情绪与主题、潜在影响、可信度、建议等级和下一责任人。
验收	来源可访问；重复内容已合并；事实与观点区分；高风险信号在时限内送达。
升级	涉及人身安全、监管调查、重大产品事故、财务披露、诉讼、数据泄露或高管事件时立即 R4 升级。
责任	公关负责人对回应策略负责；法务对法律判断负责；系统团队对监测连续性负责。

使用方法：业务负责人先确认目标和责任，再由专业、数据、技术与风险角色共同补齐输入、权限和升级条件。试运行期间，每次异常都要回写契约；连续两轮稳定后才讨论提高授权。

不适用边界：若任务目标本身仍有重大争议，或数据取得缺乏合法基础，不应通过细化技术步骤来掩盖上游决策缺口。此时应回到经营决策或合规评估。

附录 E | 组织设计常见失败模式与纠偏

竹势 AI 营销智库出品

表 E-1 | 八类失败模式

失败模式	识别信号	纠偏动作
把个人节省当作组织生产率	员工报告节省时间，但流程周期、客户结果和成本没有改善。节省的时间可能被更多修改、会议和低价值产出吸收。	建立端到端基线，把个人时间、流程周期、一次通过率与经营结果放在一起。若节省不能释放瓶颈，就重新设计流程。
先买平台，再寻找任务	技术团队建设功能丰富的平台，业务团队没有具名任务和责任，最终只剩零散试用。	要求每项平台能力绑定任务契约、业务所有者和证据门；没有任务需求的能力不进入优先开发。
用人工全审掩盖低质量	所有输出都要求人工复核，表面安全，实际审核者重新完成任务，成本没有下降。	测量审核时间与修改率。对稳定低风险任务采用规则校验和抽样，对高后果任务保留逐项批准。
把 AI 员工拟人化而忽略系统责任	组织给 AI 起岗位名称，却没有明确数据、权限、版本和管理者，出错时归因给“AI 员工”。	保留角色化入口，但后台必须映射到系统所有者、任务契约、权限策略和最终责任人。
伙伴交付黑箱化	代理商提交成品，企业不知道使用了哪些数据、工具、分包和版权材料。	合同中要求工具与数据清单、证据包、版权保证、审计权和知识转移；高风险资产不得黑箱交付。
只培训工具，不改变管理机制	员工参加课程后仍使用旧流程，审批、数据和绩效不变。	培训围绕真实任务契约，管理者同步修改目标、权限、质量标准和运行节奏。
过度自动化导致客户疏离	响应更快但语气机械，复杂客户被反复转圈，投诉与重复联系增加。	按情绪、价值、复杂度和争议设置转人工条件；把客户信任与一次解决率纳入绩效。
忽视能力退化	员工长期只做审核，系统故障时无法独立完成任务，也难以识别细微错误。	保留关键任务的周期性手工演练、盲测、轮岗和无 AI 应急运行。

失败复盘应区分四类根因。第一类是任务设计错误，例如目标不清或验收标准错误；第二类是系统错误，例如模型、数据、工具或接口失效；第三类是运行错误，例如审批超时、交接遗漏或人员能力不足；第四类是管理错误，例如为了追求数量而提高授权、忽略一线警告或不愿暂停。只有把根因放回系统，组织才能避免把责任简单推给一名员工。

纠偏顺序也应固定：先保护客户、品牌、数据与资金；再暂停或降级授权；随后保全证据并恢复人工流程；最后才分析是否更新规则、数据、人员培训或供应商安排。恢复速度本身是组织能力，应定期演练，而不是事故发生后临时讨论。

附录 F | 董事会季度复评模板

竹势 AI 营销智库出品

表 F-1 | 季度边界重设会议

议题	必须提交的材料	可能决策
经营结果	基线、增量、成本、客户影响	继续、扩大、停止
质量	抽样、修改率、事实和品牌偏差	提高或降低授权
风险	异常、投诉、越权、近失误	增加控制或熔断
人员	使用差异、培训、技能退化、负荷	调整岗位与培训
技术	模型、数据、工具、稳定性与成本	升级、替换或平台化
伙伴	SLA、知识转移、依赖与退出准备	续约、重谈或退出
法规	新规则、行业要求、跨境与内容义务	修改契约和审批

季度复评不应成为例行汇报。会议需要对至少一个边界作出明确决定：提高授权、维持、降低、重构或退出。若连续两个季度所有任务都“维持不变”，董事会应询问指标是否足够敏感，或组织是否在回避艰难选择。

复评材料必须区分事实、解释和建议。事实包括日志、样本、客户反馈与财务结果；解释是对原因的假设；建议是下一步动作。把三者混在一起会让管理层无法判断证据强度，也会让供应商自报结果被误当成独立验证。

对于稳定低风险任务，季度复评可以合并到半年；对模型大版本更新、重大渠道政策变化、业务重组、数据源变化、投诉集中出现或监管要求变化，应立即召开临时复评，不等待固定周期。

附录 G | 最小可行人机协作单元

竹势 AI 营销智库出品

一个最小可行协作单元至少包含五个角色：具名业务负责人、专业审核人、AI 系统所有者、数据所有者与运行协调人。小企业可以由同一人兼任多个角色，但不能让责任消失。兼任时应在任务契约中明确该人员以何种身份作出决定，避免“自己设计、自己审核、自己批准”成为常态。

该单元只承接一个端到端经营任务，例如“每周形成并发布一组面向目标客户的内容，收集并交接有效线索”。它不以接入多少模型、建立多少智能体或生成多少素材作为成功标准，而以周期时间、一次通过率、有效线索、客户反馈、异常率和可审计性判断。

运行的第一原则是先保持边界窄而完整：宁可覆盖一个真实闭环，也不要同时铺开十个互不连接的辅助功能。第二原则是先证明接管与恢复：上线前至少完成一次权限撤销、人工接管、输出撤回和日志追踪演练。第三原则是让一线拥有暂停权：任何发现事实、品牌、数据或客户风险的人都可以停止任务，并在规定时限内获得具名负责人响应。

当协作单元连续达到质量阈值、异常可被及时处置、经营指标出现可解释改善且员工能够独立接管时，企业可以扩大任务范围或提高授权。若结果改善仅来自额外人工投入，或审核者持续重做大部分工作，应判定为尚未形成可扩展机制。

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 ↑