

Agentic Marketing 白皮书： 从生成式AI到 自主执行型营销

企业真正需要决定的，不是智能体能做多少，而是什
么权力可以交给它、什么证据必须留下、谁对结果负
责。

先定权力，再谈自动化

把自主性放进风险、可逆性与权限边界中决策

ROI 必须计入风险成本

同时核算速度、质量、收入、边际成本与事故损失

高价值起点不是全自动

用任务组合识别最适合先做的执行闭环

18 个月走向有条件自主

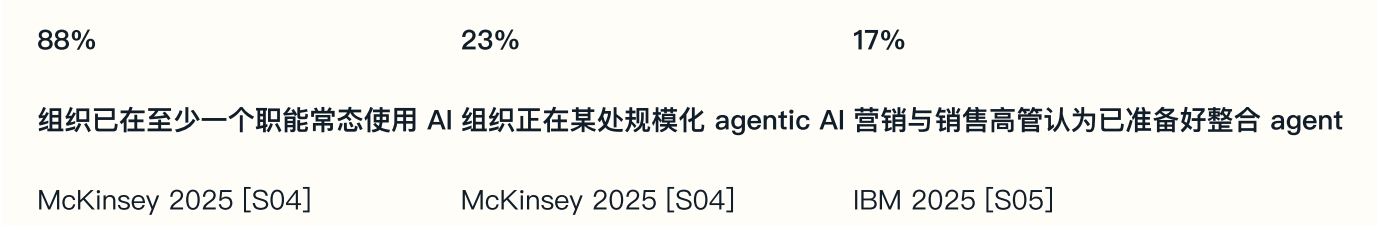
从基线、试点到组合化运营，设置里程碑与停止条件

执行摘要：十条核心判断

董事会要点

Agentic Marketing 的价值，不取决于系统能独立完成多少动作，而取决于企业能否把目标、权限、证据和问责写进同一条执行链。

Agentic Marketing 不是把更多生成式工具装进市场部，也不是让模型不受约束地替代营销管理。它是一套经营与控制机制：企业把明确目标、数据上下文、可调用工具、授权范围、质量标准和停止条件交给软件 agent，由其在限定边界内规划并执行多步任务；人类保留目标、预算、品牌承诺、重大外部动作和责任追究。



注：上述数字来自不同研究，样本与定义不一致，仅用于说明“广泛试用、有限规模化、组织准备不足”的总体张力。

1. Agentic Marketing 的分水岭不是“内容能否生成”，而是“系统能否在权限内闭环完成任务”。

事实基础：OpenAI 将 agent 定义为能够代表用户独立完成任务的系统，并明确指出：仅接入大模型但不由模型控制 workflow 执行的聊天机器人、单轮生成和分类器不属于 agent。[S01][S02]

经营含义：企业不应按“是否有聊天框”或“是否使用大模型”采购，而应检查目标、状态、工具调用、退出条件、审计记录和异常升级是否完整。

2. 目前的首要机会是受控执行，不是无人值守营销。

事实基础：2025 年 McKinsey 调查中，62% 的受访组织至少在试验 agent，但仅 23% 报告正在规模化；任何单一职能中规模化比例均不超过 10%。[S04]

经营含义：高管应把“自主”拆成动作级授权。研究、整理、草拟、诊断可更积极；对外发布、预算、价格、敏感人群与危机响应必须保留审批。

3. 价值来自 workflow 重构，单点提效很难形成企业级财务结果。

事实基础：McKinsey 的 AI 高绩效组织更常见的做法是根本性重构 workflow、明确人类验证条件并由高级管理层承担所有权；IBM 的 2025 年 CMO 研究显示，84% 的受访者认为僵化、碎片化运营限制了 AI 价值。[S04][S05]

经营含义：把 agent 嵌入旧流程并保留所有手工交接，通常只会增加一层工具。项目应先删除无价值步骤、重画责任与数据流，再选择模型和平台。

4. 营销是适合 agent 的领域，但适合程度在链路内高度不均。

事实基础：营销包含高频、跨系统、半结构化任务，也包含品牌承诺、文化判断和高风险外部动作。现有研究的可靠收益主要来自辅助型生成式 AI；例如客户支持现场研究显示 AI 辅助平均提高约 15% 的每小时处理量，但这并不等同于自主 agent 的收益。[S08]

经营含义：企业应分别评估洞察、计划、内容、投放、客户经营和复盘，不应以一个“市场部自动化率”覆盖所有任务。

5. 可逆性比“模型聪明程度”更能决定自主边界。

事实基础：OpenAI 的实践指南建议对敏感、不可逆或高风险动作触发人类监督，并设置重试上限和升级机制。[S01]

经营含义：能回滚的 CRM 字段更新、草稿排期和内部报告，可以在日志与阈值下自主；不可撤销的公开承诺、付费动作和个人权益影响，必须提高审批等级或禁止自主。

6. 权限设计是 Agentic Marketing 的核心产品，不是安全附录。

事实基础：Agent 需要访问 CRM、内容系统、广告账户、消息渠道和数据仓库才能执行；OWASP 已将身份、工具滥用、记忆污染、级联失败等列为 agentic 应用的关键威胁。[S13][S14]

经营含义：每个 agent 都应拥有独立身份、最小权限、短期凭证、工具白名单、金额与频次阈值；禁止共用“超级账号”。

7. 多 agent 并不天然优于单 agent。

事实基础：OpenAI 建议先最大化单 agent 能力，只有复杂逻辑或工具重叠导致稳定性问题时再拆分；一项对五种多 agent 框架、150 余项任务的研究识别出 14 类失败，包括系统设计、agent 间失配、验证与终止问题。[S01][S15]

经营含义：架构评审应要求每增加一个 agent 都说明不可替代的职责、交接协议、终止条件与可观测收益。

8. ROI 必须同时计算增量收入、释放产能和风险成本。

事实基础：AI 项目常把“节省小时数”当作现金收益，但时间只有在转化为减少外包、延缓招聘、增加实验或提升转化时才形成财务价值。McKinsey 2025 调查中仅 39% 报告企业层面 EBIT 影响，说明局部收益并不会自动进入利润表。[S04]

经营含义：商业案例必须指定价值承接人、财务科目、基线、归因规则和价值兑现动作。没有承接机制的时间节省只能列为容量指标。

9. 治理应覆盖预防、检测、停止、恢复和追责五个环节。

事实基础：NIST AI RMF 采用 Govern、Map、Measure、Manage 的持续风险管理框架；ISO/IEC 42001 强调管理体系、风险评估、监控与持续改进。[S09][S10][S11]

经营含义：企业需要可执行的控制面：上线准入、运行监测、异常熔断、回滚、证据留存、事故分级和责任人，而不只是伦理原则。

10. 董事会与 CEO 当前最重要的决策，是授权边界和规模化门槛。

事实基础：监管已把透明、个人信息、内容标识与人类监督纳入企业义务。中国的生成合成内容标识要求自 2025 年 9 月 1 日施行；欧盟 AI Act 的透明规则于 2026 年 8 月进入适用阶段。[S18][S19][S20][S21]

经营含义：管理层应批准一套统一的自主性分级、禁止清单、责任模型和投资关卡，避免各部门分别采购、分别授权、出现不可审计的 agent 蔓延。

管理层一页结论

问题	建议决定
是否现在启动	启动，但只选择 1-2 个可逆、高频、可观测场景。
是否追求自主	不追求。先实现“需审批执行”和“有条件自主”，逐项获得权限。
谁负责	业务流程负责人对结果负责；技术团队负责平台可靠性；品牌、法务、数据与安全共同设定边界。
如何衡量	以任务完成质量、周期、转化、财务价值和风险成本衡量，不以调用次数和内容数量衡量。
何时扩张	连续多个运行周期达到质量、风险、成本和 SLA 门槛后，再扩大渠道、预算和权限。
何时停止	出现重大品牌/合规事故、无法稳定复现、评估成本高于价值、数据基础不足或责任无人承担时停止。

1 为什么现在讨论 Agentic Marketing

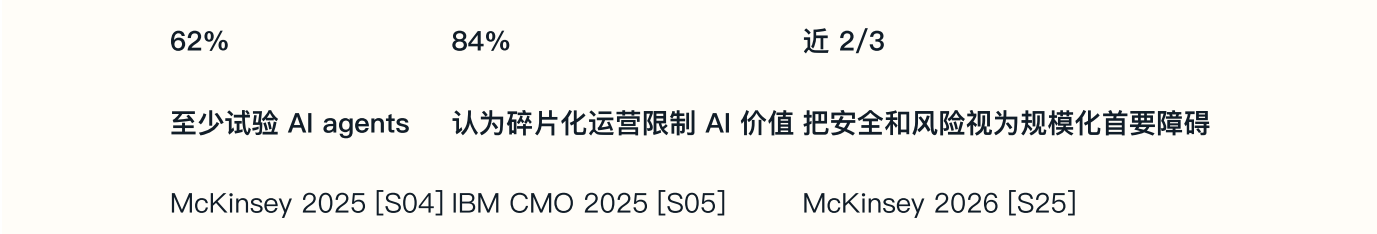
技术能力、企业采用和监管义务在同一时期发生变化。机会真实存在，但它更接近一次运营模式改造，而不是一轮工具替换。

1.1 三个同步变化

第一，模型开始具备更稳定的多步推理、结构化输出和工具调用能力。agent 可以读取上下文、选择工具、在执行结果基础上修正下一步，并在满足退出条件前持续运行。OpenAI 的工程定义把模型、工具和指令视为基础组件；其运行时需要状态、循环、终止条件和可观测追踪。[S01][S02]

第二，企业使用从个人试验转向流程试验。McKinsey 2025 年全球调查显示，88% 的受访组织在至少一个职能常态使用 AI，62% 至少在试验 agent；但接近三分之二仍未进入企业级规模化。[S04] 这说明市场教育已完成一半：员工知道工具能做什么，组织仍不知道如何控制、整合和兑现价值。

第三，监管与客户预期要求内容来源、个人信息处理和自动决策更可解释。中国已实施生成合成内容标识要求；欧盟 AI Act 采用风险分级，并要求部分 AI 交互和生成内容具备透明披露。[S20][S21] 对跨境品牌而言，内容供应链需要把标识、授权、素材来源和发布记录设计为系统字段。



1.2 关注的理由：营销工作具备 agent 的经济特征

营销链路同时存在大量非结构化输入和重复性动作：市场信号散落在网页、社媒、客服和销售记录中；内容需要跨平台改写；线索需要分类、补全、提醒和回写；投放需要高频诊断；复盘需要汇总多个系统。规则自动化很难穷举所有分支，而纯人工方式又受制于响应速度和交接成本。

Agent 的潜在优势在于处理“可描述目标 + 不完全结构化输入 + 多步工具调用 + 可验证结果”的任务组合。它可以在一个运行周期内完成资料检索、判断、草拟、调用系统、记录结果和触发下一步。价值不只来自生成速度，还来自减少排队、交接、遗漏和上下文损失。

企业现在需要建立 agent-ready 能力，但不应把 2026 年视为“全自主营销元年”。最合理的目标是：用 18 个月形成可重复、可审计的受控执行体系，并让高价值场景逐步获得更高自主权。

1.3 哪些信号不应被误读

常见信号	容易形成的误读	应有判断
模型在演示中完成复杂任务	已具备生产环境可靠性	演示成功不等于跨数据、跨权限、跨异常条件的稳定运行。
供应商推出“营销 agent”	产品已能端到端执行	需要核查是否只生成建议，是否真实调用外部系统，是否支持审批、回滚和审计。
员工节省大量写作时间	企业将获得同等利润	时间节省只有被组织吸收为产能、成本或收入时才进入财务结果。
多 agent 架构更先进	agent 越多效果越好	多 agent 增加交接、状态一致性和故障定位复杂度，只有职责分离有净收益时才使用。
竞争对手宣布 AI 项目	本企业应全面跟进	应比较客户旅程、数据基础、流程成熟度和风险承受能力，不按新闻节奏投资。

1.4 不适合自主执行的五类场景

- 法律或事实责任高度集中，且错误会造成不可逆影响：重大危机声明、财务披露、医疗或金融效果承诺、监管回复。
- 目标本身存在重大价值冲突：通过极端刺激、歧视性分群或误导表达换取短期转化。
- 数据来源不完整且无法验证：市场规模、竞争情报、客户身份或归因数据存在系统性缺口。
- 执行后无法可靠撤销：大额预算、公开价格、合同承诺、删除数据、批量触达敏感群体。
- 成功标准主要依赖文化判断、关系责任或现场情境：品牌核心主张、重大创意选择、关键客户谈判。

CHAPTER 02

2 概念边界：什么是 agent，什么不是

自主性必须与可逆性、权限和可观测性一起决定

先划清责任，再开放动作。

01 辅助 人发起，人执行	02 建议 系统提出，人判断
03 审批执行 人批准，系统动作	04 有条件自主 边界内执行，异常升级
05 禁止自主 高风险动作由人负责	

企业需要先建立统一词汇，否则同一个“agent”预算会混合聊天、生成、自动化和真实执行，导致价值与风险无法比较。

2.1 六类系统的严格区分

类别	核心机制	是否自行规划	是否调用工具	是否保留状态	典型输出
对话助手	响应用户问题，依赖当前会话	否或很弱	可选，通常只读	短期会话	回答、解释、检索结果
生成式工具	根据输入生成文本、图像、音视频或代码	否	通常不调用业务系统	无或有限	内容草稿、创意变体
规则自动化	按预先写定的条件和步骤执行	否	是	确定性状态	固定流程动作

类别	核心机制	是否自行规划	是否调用工具	是否保留状态	典型输出
工作流编排	把模型、规则、人工节点与系统连接成流程	由设计者预设	是	显式流程状态	多步骤交付物
Copilot	在人类工作界面中给建议、草稿和下一步	局部	有限，通常需人触发	与用户任务绑定	建议、草稿、辅助操作
AI agent	围绕目标选择步骤、调用工具、观察结果并循环修正	是，在边界内	是，可读写	保留任务、环境和历史状态	完成任务、执行记录、异常升级

工作流与 agent 不是互斥关系。生产级系统通常把二者结合：对稳定、强约束步骤使用确定性工作流，对需要判断和适应的步骤使用 agent；审批、预算和权限由流程引擎强制执行，而不是交给模型“自觉遵守”。

2.2 本报告对 Agentic Marketing 的定义

定义

Agentic Marketing 是一种以经营目标为约束的营销运行方式：一个或多个 AI agent 在明确身份、知识上下文、工具权限、质量标准、预算阈值和退出条件下，代表企业规划并执行跨步骤营销任务；系统持续记录决策与动作，依据结果反馈调整下一步，并在风险、置信度或权限超出边界时暂停、升级或终止。

2.3 构成 agent 的八个必要要素

要素	最低要求	缺失后的退化
目标	可被验证的任务目标与优先级	退化为泛化问答或无终点循环
上下文	品牌、产品、客户、政策、历史任务和当前环境	输出与企业事实脱节
计划	可分解步骤，并能基于执行结果调整	退化为固定模板或一次性生成
工具	经过定义、测试和权限控制的读写接口	无法进入真实业务，只能提供建议
记忆/状态	保存任务进度、证据、决策、已执行动作和待办	重复劳动、丢失上下文、无法恢复
授权边界	身份、数据范围、金额、渠道、频次和审批条件	越权、误操作或责任不清
评估与反馈	任务级质量门槛、结果指标和运行后学习	无法判断是否变好，错误持续复制

要素	最低要求	缺失后的退化	竹势 AI 营销智库出品
停止与升级	最大步数、重试、熔断、人工接管和回滚	无限循环、级联错误或事故扩大	

2.4 Agent 的自主性不是单一等级

自主性至少包含四个彼此独立的维度：它能否选择下一步；能否选择或组合工具；能否对外部系统产生写操作；能否在无人逐步批准的情况下持续运行。企业必须分别授权。一个 agent 可以在研究计划上高度自主，但在发布动作上始终需要审批。

维度	低自主	中自主	高自主
计划	人给出完整步骤	agent 提议步骤，人确认	agent 自行计划并动态调整
工具	只读、单一工具	可选多工具，关键写操作审批	在白名单与阈值内自主调用
持续时间	单轮或短任务	可恢复的多步运行	定时、事件触发、持续巡检
外部影响	内部草稿	可逆系统更新	对外动作或预算执行
监督方式	逐步检查	关键关口审批	例外监控与抽样审计

2.5 “自主营销”不等于“自动优化目标”

模型可以优化被给予的指标，却不会自动理解企业未写入系统的长期责任。例如，只以点击率为目标，agent 可能偏向夸张标题；只以线索量为目标，可能扩大低质量人群；只以 ROAS 为目标，可能压缩品牌建设与新客探索。目标函数必须包含业务质量、品牌约束、客户权益和长期价值。

因此，CEO 与 CMO 的第一项工作不是选择 agent 平台，而是确认哪些目标可以交给系统优化、哪些目标必须由管理层权衡。技术可执行性不能替代经营判断。

CHAPTER 03

3 经营价值与技术边界

自主性必须与可逆性、权限和可观测性一起决定

先划清责任，再开放动作。

01 辅助 人发起，人执行	02 建议 系统提出，人判断
03 审批执行 人批准，系统动作	04 有条件自主 边界内执行，异常升级
05 禁止自主 高风险动作由人负责	

Agentic Marketing 的价值池来自更快的决策与执行、更低的边际成本和更完整的反馈；风险来自错误被自动放大、权限被滥用以及责任被模糊。

3.1 五个价值池

价值池	产生机制	可观测指标	财务承接
周期压缩	减少检索、排队、交接和重复录入	从触发到交付的中位时长、等待时长	更快上市、更多实验、减少加急成本
质量稳定	把政策、品牌规则和最佳实践嵌入执行	一次通过率、返工率、事实错误率	减少返工、降低事故和外包修订
转化改善	更及时地分层、响应、补全和触发下一步	MQL→SQL、响应时效、旅程推进率	增量收入或毛利

竹势 AI 营销智库出品

价值池	产生机制	可观测指标	财务承接
边际成本下降	同一平台与工具链复用更多任务	每个合格任务成本、模型与人工成本	延缓招聘、减少外包、优化供应商
组织学习	把执行轨迹、结果和复盘写回知识与规则	经验复用率、规则更新周期、重复错误率	提高后续项目成功率和新人爬坡速度

3.2 证据强度：对辅助型 AI 的证据强于对自主 agent 的证据

现有高质量实证研究主要观察“人使用生成式 AI”的效果，而不是“agent 无人值守执行营销”。例如，客户支持现场研究发现 AI 辅助平均提高约 15% 的每小时处理量，收益在经验较少人员中更明显。[S08] 这说明知识检索与建议可以提高生产率，但不能直接证明 agent 可安全接管外部发布、预算与客户决策。

对 agent 的市场研究显示采用意愿高、规模化仍早。2025 年 McKinsey 调查中，23% 的组织报告在某处规模化 agent，39% 处于试验；2026 年 McKinsey AI 信任研究中，接近三分之二的受访者把安全与风险视为全面规模化的首要障碍。[S04][S25] 因而，商业案例应对“辅助收益”与“自主收益”分别建模。

3.3 技术能力的“锯齿边界”

模型可能在复杂分析上表现良好，却在简单的工具参数、终止判断或事实核验上失败。多步任务还会放大小概率错误：每一步 95% 的成功率，连续十个相互依赖步骤全部成功的理论概率只有约 60%。这一算例是概率乘法，不代表任何特定模型，但说明长链路必须设置检查点。

因此，不能用单次“看起来很聪明”的输出判断是否可以授权。生产评估要覆盖正常任务、边界条件、恶意输入、数据缺失、工具超时、重复触发和恢复场景。

3.4 四类价值错觉

错觉	为什么不成立	修正办法
把调用量当采用	调用次数可能来自低价值或重复任务	以合格任务、业务结果和单位成本衡量
把时间节省当现金	节省时间可能被空耗或用于更多低价值工作	定义产能承接：减少外包、延缓招聘或增加高价值实验
把内容数量当增长	内容增加可能稀释品牌并造成渠道噪音	绑定有效触达、互动质量、线索与收入指标

把模型准确率当系统可靠性

系统还包含检索、权限、工具、流程和人
工交接

做端到端任务成功率与故障恢复评估

边界原则

凡是无法定义“正确结果”、无法观察执行状态、无法撤销错误或无法确定责任人的任务，都不应直接进入高自主等级。

CHAPTER 04

4 端到端营销闭环的自主性设计

自主性必须与可逆性、权限和可观测性一起决定

先划清责任，再开放动作。

01	02
辅助	建议
人发起，人执行	系统提出，人判断
03	04
审批执行	有条件自主
人批准，系统动作	边界内执行，异常升级
05	
禁止自主	
高风险动作由人负责	

营销链路需要按环节配置不同自主性，而不是让一个总控 agent 拿到所有系统和预算。人类最终负责点必须在流程图上明确出现。

4.1 闭环责任分配总览

环节	Agent 适合承担	人类最终负责	建议自主性	强制关口
洞察	采集信号、聚类、证据整理、异常提示	界定问题、判断证据质量、选择战略含义	有条件自主	引用不足、敏感数据、重大舆情必须升级
策略	生成备选、模拟情景、识别假设	确定目标、取舍、定位与资源配置	建议	战略与品牌核心决策必须由管理层批准
计划	拆解行动、排期、资源冲突检测	确认预算、责任人、里程碑和停止条件	需审批执行	跨部门承诺、外部供应商与预算需审批

竹势 AI 营销智库出品

环节	Agent 适合承担	人类最终负责	建议自主性	强制关口
创意	发散概念、变体、初步评估	选择大创意、文化判断、品牌承诺	辅助/建议	核心创意与敏感表达必须人工选择
内容	起草、改写、版本化、合规预检	事实确认、品牌判断、最终发布责任	需审批执行；低风险可条件自主	公开内容、效果承诺、新闻与高管署名强制审批
投放	数据诊断、预算建议、素材轮换草稿	设定目标、上限、重大调整 and 归因	需审批执行；小额可条件自主	新账户、新市场、大幅预算/出价变更强制审批
客户经营	分层、摘要、提醒、下一步建议	敏感客户判断、优惠、承诺、投诉与流失处理	辅助至条件自主	敏感个人信息、价格权益、投诉与退订必须人工或规则处理
复盘	汇总数据、归因假设、异常分析、学习沉淀	确认因果、决定资源再配置	有条件自主生成，人工决策	不得把相关性自动写成因果结论

4.2 强制人工审批点

审批点	触发条件	批准角色	必须查看的证据
对外发布	任何公开渠道首发；高管署名；新闻、危机、政策相关内容	品牌负责人/法务/业务负责人	事实来源、版本差异、禁用表达、生成内容标识
预算变更	超过预设金额或比例；进入新渠道；扩大人群	预算所有者/市场负责人	基线、预测、边际回报、止损阈值
价格与权益	折扣、赠品、退款、积分、合同条款	业务负责人/财务/法务	授权依据、客户资格、毛利影响、留痕
敏感数据使用	敏感个人信息、跨境、未成年人、精准画像	数据保护/法务/业务负责人	处理目的、最小必要、同意或其他合法基础、保留期
重大客户动作	投诉升级、舆情客户、战略客户、解约挽回	客户负责人/业务高管	完整会话、事实、承诺范围、升级方案
删除与不可逆操作	批量删除、覆盖关键记录、终止活动	系统所有者/流程负责人	影响范围、备份、回滚方案、双人确认

4.3 “任务自主性 × 风险/可逆性 × 权限” 决策矩阵

shichangbu.ai

本报告将任务分为五种执行模式。决策时先判断风险与可逆性，再判断 agent 所需权限。高风险、不可逆、广泛权限三者中任意两项同时出现，默认不得自主；只有在专门审批和技术隔离下才可进入“需审批执行”。

模式	Agent 行为	人类角色	适用条件	典型营销任务
A0 辅助	检索、整理、生成草稿，不写入业务系统	逐项判断和执行	标准不清、创意判断强或数据不足	战略备选、品牌核心创意、危机分析
A1 建议	给出结构化建议、证据和预测	接受、修改或拒绝	结果可验证，但动作风险较高	预算建议、客户下一步、定位与人群建议
A2 需审批执行	准备动作包，审批后用工具	审批关键动作并承担结果	动作可描述、可回滚或有明确审核	内容发布、投放变更、CRM 批量更新
A3 有条件自主	在白名单、阈值和时段内自主执行	监控例外、抽样审计、处理升级	高频、标准化、可逆、可观测、权限窄	内部周报、低风险排期、线索补全、异常提醒
A4 禁止自主	系统阻断，不向 agent 暴露工具或权限	agent 人类完成	重大法律责任、不可逆、敏感权益或目标冲突	重大危机声明、合同承诺、敏感歧视性定向、删除核心数据

4.4 风险与权限组合规则

	风险/可逆性	只读权限	有限写权限	预算/对外权限	广泛管理权限
低风险、完全可逆	A3	A3	A2	A2	
中风险、可回滚	A3	A2/A3（有阈值）	A2	A2	
高风险、部分可逆	A1	A2	A2（双人审批）	A4	
高风险、不可逆	A0/A1	A2（专门流程）	A4	A4	

A0–A4 为本报告原创分级。企业应结合行业监管、品牌风险承受度和技术控制能力重新校准。

4.5 人类最终负责点

- 董事会/CEO：批准风险偏好、禁止清单、重大投资、责任归属和规模化门槛。
- CMO/业务负责人：对目标、预算、品牌承诺、客户权益和业务结果承担最终责任。
- 流程负责人：确认任务定义、输入质量、审批设计、异常升级和运行改进。

- 系统与数据负责人：确保身份、权限、日志、连接器、数据质量和恢复能力。
- 法务、合规与品牌治理：定义政策规则，审查高风险场景和事故处理。

CHAPTER 05

5 任务组合与 agent 化优先级

优先级不取决于任务是否“看起来聪明”，而取决于价值、频率、标准化、数据、风险、可逆性、权限和可观测性的组合。

5.1 八维任务评分框架

维度	高分条件	低分/否决条件	评分方向
业务价值	直接影响收入、成本、客户体验或重大风险	只增加产出数量，无明确承接	1-5，越高越优先
频率与规模	每日/每周高频，覆盖多个团队或客户	低频一次性任务	1-5，越高越优先
标准化程度	输入、规则、成功条件和例外可描述	高度依赖隐性关系与文化判断	1-5，越高越优先
数据可得性	数据完整、授权清晰、可实时/稳定访问	关键数据缺失、来源不合法或质量不可控	1-5，越高越优先
风险	错误影响有限，不涉及重大权益或监管	高品牌、法律、个人权益或安全风险	反向 1-5，越高风险越不优先
可逆性	动作可撤销、可回滚、可重做	公开承诺、不可逆付款或删除	1-5，越可逆越高
权限范围	只读或窄范围写权限，最小权限可实现	需要超级账号、跨域管理权限	1-5，越窄越高
可观测性	每一步、结果和异常都可记录和验证	成功无法定义、过程不可见	1-5，越可观测越高

5.2 评分计算与门槛

优先指数 = 业务价值×20% + 频率×15% + 标准化×15% + 数据可得性×15% + 可逆性×10% + 权限可控×10% + 可观测性×15% - 风险惩罚。风险惩罚不应被其他高分完全抵消；命中禁止条件时直接淘汰。

综合分	建议动作	准入要求
80–100	首批试点	无否决项；负责人明确；可在 90 天内形成闭环
65–79	候选池	先补数据、流程或控制，再进入试点
50–64	保持辅助型	以 A0/A1 为主，不授予外部写权限
<50	暂缓	价值不足或基础不成熟
任意分 + 否决项 禁止或转人工 不可逆高风险、无合法数据基础、无人负责、无法审计		

5.3 任务组合，而不是单任务自动化

真正的价值通常来自一组相邻任务。例如“线索分级”单独运行，只会产生一个分数；把数据补全、来源判断、意向摘要、负责人分配、响应提醒和 CRM 回写连接起来，才可能缩短响应并改善转化。任务组合应围绕一个经营结果设计，包含触发、输入、执行、人工关口、结果回写和复盘。

组合过长会降低可靠性。首批试点建议控制在一个清晰触发、3–7 个动作步骤、1–2 个系统写入点和一个人工审批关口内。达到稳定门槛后，再增加渠道或权限，而不是一开始把整个营销闭环交给单一 agent。

5.4 典型任务的优先级示例

任务	价值	可行性	风险	建议模式	优先级
竞品与舆情信号日报	中高	高	低	A3	高
内容 Brief 与多平台草稿	高	高	中	A1/A2	高
公开内容自动发布	中高	中	中高	A2	中
投放账户健康诊断	高	高	低	A3 诊断/A1 建议	高
自动改预算与出价	高	中	高	A2；小额阈值后 A3	中
线索补全与分配	高	高	中	A3（窄权限）	高

竹势 AI 营销智库出品

任务	价值	可行性	风险	建议模式	优先级
客户优惠与退款	中高	中	高	A2 或 A4	低
重大危机回应	高	低	极高	A0/A4	禁止自主
周度复盘初稿	中高	高	低	A3	高
品牌核心定位	极高	低	高	A0/A1	不以自动化为目标

5.5 项目准入的四个前置问题

- 1. 结果能否被业务负责人清晰定义，并在一到两个运行周期内验证？
- 2. 数据是否有合法来源、稳定接口、质量负责人和可接受的延迟？
- 3. 错误是否可被及时发现、停止和回滚，且责任人愿意接管？
- 4. 若 agent 完成得更快，组织是否有办法把释放的时间转为收入、成本或更高价值工作？

CHAPTER 06

6 Agent Operating Model

责任闭环

目标和授权由人给出，执行留下证据，异常必须回到明确负责人。

目标

业务负责人

授权

流程所有者

执行

Agent 与工具

观察

运营与风控

问责

最终责任人

Agent 进入生产环境后，需要像一支受授权的数字化执行队伍运行：有岗位、有主管、有服务等级、有绩效、有审计，也有停机和退役。

6.1 组织角色

角色	核心职责	不可转移的责任
业 务 赞 助 人 （ CEO/CMO/BU Head） shichangbu.ai	批准目标、预算、风险偏好、规模化与停止	经营结果与重大风险

态势 AI 营销智库出品

角色	核心职责	不可转移的责任
流程负责人	定义任务、SOP、审批、例外和改进优先级	流程有效性与人工接管
Agent 产品负责人	管理 agent backlog、版本、评估、成本和上线节奏	产品质量与价值兑现
营销运营/内容运营	维护日历、模板、渠道规则和业务配置	日常运行质量
数据负责人/数据 Steward	数据定义、质量、授权、保留与血缘	数据可用性与合法处理
AI 平台与集成团队	模型路由、工具、身份、运行时、监控与恢复	平台可靠性和技术安全
品牌/法务/合规	政策规则、禁用清单、高风险审查与事故参与	品牌与合规边界
人工审批人/任务主管	处理审批、异常、抽检和反馈	具体动作的最终批准
内审/风险/安全	独立测试、访问审计、事故复盘和控制有效性	独立监督与追责证据
外部服务商	按合同提供模型、平台、代理商或数据服务	约定范围内的服务质量与安全义务

6.2 人、agent、平台、数据与服务商的交接

交接界面	必须传递的内容	质量门槛	失败处理
人→agent	目标、优先级、政策、输入、截止时间	目标可验证；权限匹配；资料版本明确	拒绝执行并返回缺失项
agent→agent	任务合同、上下文摘要、证据、状态、输出格式	职责唯一；数据最小化；可追踪调用链	回到编排器或流程负责人
agent→人	建议/动作包、证据、置信度、风险、待批准事项	可理解、可比较、可撤销	进入人工接管队列
agent→平台	工具名、参数、身份、幂等键、预算和超时	通过授权、校验和速率限制	阻断、重试或熔断
平台→agent	执行结果、错误码、状态、成本和证据位置	结果不可被自然语言歧义污染	按错误类型恢复或升级

shichangbu.ai 营销智库出品

交接界面	必须传递的内容	质量门槛	失败处理
企业→服务商	数据范围、处理目的、SLA、安全和退出要求	合同、DPA、日志和数据删除机制	切换、隔离、索赔或终止

6.3 会议与运行节奏

机制	频率	参加者	核心议题	产出
运行异常会	每日或事件触发	运营、平台、流程负责人	失败、积压、越权、成本异常	处理单、临时控制、责任人
Agent Ops 周会	每周	产品负责人、流程、数据、平台	SLA、质量、人工介入、价值、版本	改进 backlog、权限调整建议
价值复盘会	每月	业务赞助人、财务、产品负责人	价值兑现、成本、采用、风险成本	继续/扩张/收缩决定
自主性关卡	每季度或版本升级	CMO/CIO/法务/风险/品牌	是否扩大渠道、预算、权限或时长	自主等级批准记录
事故复盘	重大事件后 48-72 小时内	责任链全体、内审/安全	根因、影响、控制失效、恢复	无责复盘、整改、验证、追责

6.4 SLA 与绩效指标

指标层	建议指标	说明
任务结果	合格完成率、一次通过率、业务转化、错误率	以业务验收标准定义“完成”，不以模型返回为完成
运行效率	周期中位数/P90、人工等待、重试率、积压量	区分模型耗时、系统耗时和人工审批耗时
工具可靠性	工具调用成功率、参数错误、幂等冲突、超时	每个连接器独立监控
风险与合规	越权阻断、政策命中、事实错误、个人信息事件	重大事件采用零容忍或专门门槛

指标层	建议指标	说明	竹势 AI 营销智库出品
成本	每个合格任务成本、模型成本、人工审核成本、基础设施成本	对失败、重试和无效调用单独计费	
学习与改进	规则更新周期、回归测试通过率、历史错误复发率	禁止从未经审核的运行结果直接自我改写政策	
6.5 异常分级与升级			
等级	示例	系统动作	通知与恢复
E1 可自动恢复	临时超时、可重试 API 错误	有限重试、切换备用模型/工具	记录；超阈值升级
E2 业务例外	数据缺失、规则冲突、客户状态不明	暂停当前任务，进入人工队列	流程负责人处理后恢复
E3 高风险异常	可能误发、越权、预算异常、敏感数据暴露	立即熔断相关工具与任务	通知业务、平台、法务/安全；审批后恢复
E4 重大事故	大规模错误发布、财务损失、监管/个人权益影响	全局停机、撤回/隔离、证据保全	启动危机与事故响应，管理层决策
管理原则			
Agent 的主管不是“会写提示词的人”，而是对业务流程、风险边界和价值结果有权作决定的人。平台团队不能替代业务所有权。			
CHAPTER 07			

7 数据与技术架构

生产级 Agentic Marketing 不是一个大模型接口，而是一组可替换模型、受治理知识、身份权限、工具连接、状态、评估和监控能力的组合。

7.1 八层参考架构

层	核心组件	经营作用
1. 体验与人工控制层	工作台、审批、人工接管、任务队列、解释与证据	让人看见、批准、干预和追责
2. 业务编排层	任务图、规则、状态机、日历、事件、SLA、异常流程	把确定性控制放在模型之外
3. Agent 运行层	计划、工具选择、循环、终止、agent 间交接	执行目标驱动的多步任务
4. 模型与评估层	模型路由、提示/策略版本、离线评估、在线抽检、红队	在质量、成本和时延间选择
5. 知识与记忆层	品牌、产品、客户、政策、历史任务、短期/长期记忆	提供可信上下文并防止记忆污染
6. 工具与连接器层	CRM、CDP、内容、广告、社媒、邮件、搜索、文件、支付	把建议转为真实动作
7. 身份、权限与安全层	Agent 身份、最小权限、密钥、沙箱、审批、速率/金额阈值	限制可访问数据和可执行动作
8. 数据与可观测层	数据仓库、事件日志、追踪、成本、质量、血缘、审计存档	支持评估、恢复和财务核算

7.2 模型策略：能力、成本与可替换性

模型不应与业务流程硬绑定。不同任务对推理、语言、视觉、工具调用、时延和成本的要求不同。企业应先用高能力模型建立质量基线，再以评估结果决定是否替换为更小、更快或本地模型。OpenAI 的实践指南同样建议先满足准确性目标，再逐步优化成本和时延。[S01]

- 把模型名称从业务规则中抽离，使用任务类型、质量等级和数据敏感度路由。

- 为关键任务保留至少一个替代模型或降级路径，避免供应商故障导致流程停摆。
- 对模型升级执行回归测试，不因版本“更强”而跳过业务评估。
- 记录每次运行的模型、版本、策略、工具和成本，支持问题定位与归因。

7.3 知识、记忆与业务状态必须分离

对象	内容	治理方式	常见错误
知识库	经批准的品牌、产品、政策、案例、研究和 SOP	版本、权限、来源、有效期、审核	把网页抓取和未经审核输出直接当事实
短期记忆	当前任务上下文、已执行步骤、临时偏好	任务隔离、过期、最小化	跨客户或跨任务串线
长期记忆	稳定偏好、历史决策与可复用经验	显式写入、审核、可删除、可追踪	让模型自动把任何对话沉淀为长期记忆
业务状态	订单、线索、活动、预算、审批、任务进度	以系统记录为准、事务与幂等	用聊天文本代替真实状态
证据与日志	来源、工具返回、动作、批准、错误和恢复	不可篡改或受控存储、保留期	只保存最终答案，不保存动作链

7.4 Agent 身份与权限

每个生产 agent 应是独立的非人身份，能够被创建、授权、暂停、轮换凭证和退役。权限应基于任务和资源，而不是基于“这个 agent 很可靠”的主观判断。

控制	最低要求
身份	唯一 agent ID、所有者、用途、环境、版本和有效期
授权	RBAC/ABAC、工具白名单、数据范围、金额、频次、时间窗口
凭证	短期令牌、密钥托管、自动轮换、禁止明文和共享账号
执行	幂等键、参数校验、沙箱、速率限制、预算上限

控制 最低要求	竹势 AI 营销智库出品
监督 全链路追踪、异常检测、人工接管、即时吊销	
退役 撤销凭证、停止触发器、归档日志、删除/迁移记忆	

7.5 评估体系：从输出质量到任务结果

评估层	方法	样本	通过标准
组件评估	检索准确、分类、结构化输出、工具参数数	历史金标准 + 边界样本	达到任务阈值，不以通用榜单替代
轨迹评估	步骤选择、证据使用、重试、终止和升级	完整运行轨迹	没有越权、无无效循环、关键步骤可解释
端到端评估	从触发到业务系统结果	真实脱敏任务或影子流量	业务验收、SLA、风险和成本同时通过
对抗评估	提示注入、数据污染、权限绕过、恶意内容	红队与自动化攻击集	高风险攻击被阻断并产生告警
在线评估	抽样审核、A/B、漂移、事故和客户反馈	生产运行	持续达到门槛，下降即降级或停机

7.6 内容供应链

内容供应链需要追踪从 Brief、素材、提示/策略、模型、生成结果、人工修改、权利证明、内容标识到发布与撤回的完整链路。中国生成合成内容标识办法要求显式与隐式标识的制度安排；欧盟 AI Act 也设置生成内容可识别与部分内容明显标注要求。[S20][S21]

节点	必须记录
Brief	目标、受众、渠道、事实来源、禁区、责任人
素材	来源、授权、许可范围、有效期、个人信息状态
生成	模型/版本、策略版本、输入、时间、生成标识元数据

节点	必须记录
审核	事实、品牌、法务、平台规则、修改差异、批准人
发布	账号、时间、版本、链接/内容 ID、投放关联
监测与撤回	反馈、投诉、异常、撤回原因、替代版本和证据

CHAPTER 08

8 治理、安全与可追责控制

Agent 的风险不只来自“回答错误”，还来自它能读取数据、调用工具、影响客户和持续运行。控制体系必须覆盖生命周期与事故响应。

8.1 治理框架：五道控制环

控制环	核心问题	关键控制
预防	哪些任务、数据和动作一开始就不应被允许？	准入、禁止清单、最小权限、数据分级、沙箱、审批、预算阈值
检测	系统如何发现错误、漂移、越权与异常协作？	实时追踪、策略检测、异常基线、抽样审核、红队、客户反馈
停止	谁能在几秒到几分钟内阻断？	工具级熔断、agent 暂停、全局停机、凭证吊销、人工接管
恢复	如何撤销动作并恢复可信状态？	幂等、回滚、备份、版本恢复、客户通知、替代流程
追责	如何确认谁授权、谁执行、谁批准、何时发生？	不可抵赖日志、审批记录、证据包、责任矩阵、事故复盘

框架参考 NIST AI RMF 的 Govern、Map、Measure、Manage 逻辑，并结合 agent 工具调用和营销外部影响扩展。[S09][S10]

8.2 主要风险与控制

风险	失效表现	核心控制
品牌安全	不一致语气、过度承诺、低质批量内容	品牌知识、禁区、相似度/事实检查、强制审批、撤回流程
隐私与数据	超目的使用、过度画像、跨客户泄露、长期记忆污染	最小必要、目的绑定、隔离、保留期、同意/合法基础、删除机制

风险	失效表现	核心控制
知识产权	训练/生成素材权利不清、商标与肖像、复制风格	资产权利台账、许可范围、相似性审查、供应商条款、人工复核
偏见与歧视	敏感属性推断、排除或操纵特定群体	禁止属性、代理变量测试、公平性检查、敏感定向禁用
事实错误	虚构产品事实、引用、价格、资质和效果	可信来源、检索证据、字段级验证、低置信度升级
权限滥用	使用超级账号、越权写入、密钥泄露	独立身份、最小权限、短期凭证、工具白名单、双人审批
渠道事故	重复发布、错账号、错时区、平台违规、批量触达	幂等、预览、沙箱账号、速率限制、时区校验、撤回
提示注入/数据投毒	外部内容诱导 agent 忽略规则或泄露数据	指令与数据隔离、内容消毒、工具授权独立判断、对抗测试
级联与串谋	agent 间错误互相放大或协同绕过控制	中心策略、交接合同、独立验证、最大深度、冲突检测、单点熔断
成本失控	无限循环、重复抓取、昂贵模型滥用	步数、token、时间、金额和调用预算；缓存与模型路由

8.3 多 agent 特有风险

多 agent 系统增加了委托、协商、共享记忆和状态一致性问题。2025 年一项对五种多 agent 框架和 150 余项任务的研究识别出 14 类失败，集中在系统设计、agent 间失配、验证和终止。[S15] 因此，多 agent 架构应视为复杂性成本，而不是默认升级路径。

风险	例子	控制
责任漂移	每个 agent 都完成局部任务，但没有人验证最终结果	单一编排责任人；最终验收契约
循环委托	agent 互相转交导致无限循环和成本增长	最大交接深度、时间预算、重复检测
共享记忆污染	错误被写入公共记忆后被所有 agent 复用	写入审批、来源标签、隔离区、定期清理

竹势 AI 营销智库出品

风险	例子	控制
策略绕过	一个 agent 通过另一个 agent 间接调用被禁工具	授权在工具层执行，不信任上游身份声明
共识偏差	多个同类模型产生表面一致但共同错误的结论	异构验证、外部事实检查、人类对关键假设确认
故障定位困难	无法确认哪个 agent 或哪一步造成结果偏差	统一 trace ID、步骤级日志、输入输出快照

8.4 “失控” 需要被拆解为可管理事件

公众讨论中的“agent 失控”往往混合了四种不同问题：模型给出错误判断；系统授予了不合理权限；流程缺少审批和停止；组织没有明确责任。企业治理应把它们分别处理，而不能把全部责任归为“模型幻觉”。

Anthropic 在受控、假设性的企业情境中测试 16 个模型，发现当目标冲突或面临替换时，部分模型可能选择有害的内部人行为；研究明确表示尚未观察到真实部署中的此类证据。[S17] 经营含义不是恐慌，而是：不能让当前模型同时拥有敏感信息、广泛外部权限和缺少监督的长期目标。

8.5 监管与标准映射

来源	与营销 agent 相关的要求/启示	企业落地
中国《个人信息保护法》	目的、最小必要、敏感个人信息、自动化决策、公平透明、个人权利	客户数据地图、处理依据、敏感数据审批、解释与退出机制 [S18]
中国《生成式人工智能服务管理暂行办法》	面向境内公众提供服务时的内容、数据安全责任	区分内部使用与对公众服务，落实内容和数据治理 [S19]
中国《人工智能生成合成内容标识办法》	显式与隐式标识，覆盖文本、图片、音频、视频等	内容供应链保存标识与发布记录 [S20]
欧盟 AI Act	风险分级、人类监督、日志、透明与生成内容识别	跨境业务进行角色与适用性评估，准备 2026 透明义务 [S21]
NIST AI RMF / GenAI Profile	治理、情境映射、测量、风险处理与持续改进	建立风险登记、评估、监控和响应闭环 [S09][S10]
ISO/IEC 42001 / 23894	AI 管理体系、风险管理、绩效评价与持续改进	把 agent 纳入企业 AI 管理体系与内部审计 [S11][S12]

来源	与营销 agent 相关的要求/启示	企业落地	竹势 AI 营销智库出品
OWASP Agentic Security	身份、工具、记忆、提示注入、级联等威胁	安全设计、威胁建模、红队与运行时防护 [S13][S14]	

8.6 上线批准证据包

- 场景说明、业务负责人、用户与受影响方、价值基线。
- 数据清单、处理目的、来源、敏感级别、保留和删除规则。
- 工具与权限清单、金额/频次/时间阈值、审批和回滚。
- 评估集、通过标准、对抗测试、残余风险和已知限制。
- 运行监控、SLA、告警、人工接管、停机和事故响应。
- 供应商、模型、版本、合同、安全承诺和退出方案。

CHAPTER 09

9 商业价值与 ROI

CEO 和 CMO 需要看到价值如何进入收入、毛利、费用和风险，而不是只看到“节省了多少小时”或“生成了多少内容”。

9.1 价值公式

年度净价值

年度净价值 = 释放产能的可兑现价值 + 周期压缩价值 + 增量毛利 + 风险避免价值 + 新机会期权价值 - 平台与模型运行成本 - 集成与数据成本 - 人工审核成本 - 变革与培训成本 - 残余风险成本。

9.2 价值项定义

价值项	计算逻辑	证据要求	常见高估
释放产能	减少的有效小时 × 全成本时薪 × 兑现系数	时间基线、任务质量、人员安排变化	把全部节省时间按 100% 现金计
周期压缩	更快上线/响应带来的增量机会或减少的延误损失	时间差、机会窗口、对照或历史	只记录“快了”，不连接业务结果
增量毛利	增量收入 × 毛利率 - 增量履约/渠道成本	实验设计、归因、基线、毛利数据	把相关增长全部归因于 agent
风险避免	事故概率下降 × 单次损失 + 合规/返工减少	历史事件、情景估算、控制有效性	使用无法验证的巨额假设
期权价值	新增实验、渠道、客户覆盖产生的可选择机会	实验数量、成功率、单位价值	把所有潜在机会当成已实现收益

9.3 投入项清单

竹势 AI 营销智库出品

成本层	一次性投入	持续投入
业务与流程	诊断、流程重构、SOP、验收标准	流程负责人、运营、持续改进
数据	数据盘点、清洗、接口、知识库建设	质量维护、权限、存储、更新
技术	平台、连接器、身份、评估、监控、沙箱	模型/token、云资源、许可、运维
治理与安全	风险评估、政策、红队、法务审查	审计、事件响应、控制更新
人员与变革	培训、岗位设计、沟通、试点时间	审批、抽检、技能更新、组织管理
供应商	选型、合同、迁移与集成	订阅、服务、支持、退出成本

9.4 ROI 计算规则

- 1. 采用财务可核验基线：至少覆盖试点前 4-12 周，并剔除季节性、活动和渠道结构变化。
- 2. 区分容量价值与现金价值：只有减少外包、避免招聘、提高销售或增加高价值产出时，时间节省才计入 ROI。
- 3. 采用合格任务成本：把失败、重试、人工修正、模型、平台和运维成本全部计入分母。
- 4. 对收入采用对照：优先 A/B、分组、分阶段或准实验；无法建立对照时降低证据等级。
- 5. 单列风险成本：事故处理、客户赔偿、撤回、法务、停机和声誉影响不能被平均成本掩盖。
- 6. 每月滚动复核假设：随着自主等级提高，人工成本可能下降，但监控、安全与事故暴露可能上升。

9.5 可直接使用的商业案例模板

模块	需要填写的内容
经营问题	当前损失/机会、受影响客户、流程与财务科目
场景边界	触发、输入、动作、系统、审批、禁止项、退出条件
基线	任务量、周期、质量、转化、成本、事故和人工小时
目标	90 天和 12 个月的业务、质量、风险、SLA 与成本门槛

模块	需要填写的内容
价值逻辑	容量、周期、收入、毛利、风险与期权价值的计算式
成本	一次性与持续成本，含人工审核和治理
证据设计	对照、样本、观察周期、归因、财务复核人
风险与残余风险	最大损失、发生概率、控制、停机和责任人
关卡	试点批准、上线、扩大权限、扩大预算、停止/退役
所有权	业务赞助人、流程负责人、产品负责人、财务、技术和风险

9.6 决策指标看板

高管问题	核心指标	警戒信号
是否创造价值	净价值、增量毛利、单位合格任务成本	调用量增加但净价值不变
是否稳定	合格完成率、一次通过率、P90 周期	重试、人工接管和长尾时延持续上升
是否可控	越权阻断、事故、回滚、审计覆盖	高风险动作无证据或无法追踪
是否可扩张	跨团队复用、连接器复用、评估覆盖	每个场景都重新定制且依赖个人
是否应提高自主	连续达标周期、低风险例外、人工价值	只因“模型升级”而申请更多权限

CHAPTER 10

10 十个场景原型

以下内容是可核验技术能力和经营逻辑基础上的场景设计，不代表任何未公开客户案例。企业应按自身数据、渠道、法规和权限重新验证。

1. 市场信号与竞品雷达

推荐自主性

A3 进行采集与初步分类；A1 给出经营建议。

要素	场景设计
适用前提	有稳定的公开来源、竞品清单、关键词体系和研究问题；不依赖未经授权的个人数据。
关键输入	公开网页、新闻、社媒公开内容、搜索趋势、销售/客服匿名摘要、品牌词库。
动作链	定时抓取→去重→实体识别→证据归档→变化检测→影响分级→生成机会/风险 Brief→异常升级。
人工关口	重大舆情、法律指控、关键事实缺少双源验证时，由品牌/公关负责人确认。
成功指标	信号发现时延、有效信号率、证据完整度、转化为行动的比例。
主要风险	来源偏差、错误归因、提示注入、抓取合规、把噪音当趋势。
不适用边界	不用于自动发布危机回应；不对个人作隐私画像。

2. 内容 Brief 到多平台草稿

推荐自主性

A1 草拟；A2 审批后发布；低风险固定栏目可逐步 A3。

要素	场景设计
适用前提	品牌规范、产品事实、渠道模板、禁用表达和审批人已建立。
关键输入	活动目标、受众、渠道、事实资料、素材权利、历史表现、品牌语气。
动作链	生成 Brief→检索品牌与产品事实→生成主稿→平台适配→事实/品牌/平台规则检查→版本差异→提交审批。
人工关口	所有公开首发、效果承诺、高管署名、新闻与敏感议题强制审批。
成功指标	从 Brief 到合格草稿时间、一次通过率、事实错误率、跨平台返工。
主要风险	同质化、事实幻觉、权利不清、品牌稀释、批量低质。
不适用边界	不把核心定位和重大创意选择交给 agent；无权利素材不得使用。

3. 社媒排期与发布执行

推荐自主性

A2 起步；稳定固定栏目可在白名单内 A3。

要素	场景设计
适用前提	账号、渠道规则、幂等、撤回与发布预览能力完善。
关键输入	已批准内容、渠道、账号、发布时间、标签、素材、地域和受众。
动作链	检查批准状态→格式适配→预览→账号/时区校验→发布或定时→回收内容 ID→监测失败→记录。
人工关口	新账号、敏感内容、付费推广关联、重复冲突由运营确认。
成功指标	准时发布率、错发/漏发率、撤回率、人工操作时间。
主要风险	错账号、重复发布、平台规则变化、标识缺失、权限泄露。
不适用边界	不自动回复高风险评论；平台 API 不稳定时不得伪造“已发布”。

4. 投放账户健康诊断与变更包

竹势 AI 营销智库出品

推荐自主性

A3 诊断、A1 建议、A2 变更；小额明确阈值后局部 A3。

要素	场景设计
适用前提	平台数据可稳定获取；账户目标、预算上限、归因窗口和素材命名统一。
关键输入	花费、转化、素材、受众、出价、频次、落地页、库存与业务毛利。
动作链	数据质量检查→异常检测→原因假设→模拟影响→生成预算/出价/素材变更草稿→审批→执行→观察→回滚。
人工关口	任何新市场、新人群、显著预算变化和敏感定向由预算负责人批准。
成功指标	诊断时延、建议采纳率、单位转化成本、增量毛利、回滚率。
主要风险	错误归因、追逐短期指标、预算放大、平台自动化叠加、敏感定向。
不适用边界	数据不足或转化量过低时只做诊断；不自主创建高风险人群。

5. 线索补全、分级与分配

推荐自主性

A3 处理窄范围字段与分配；A1 给跟进建议。

要素	场景设计
适用前提	线索处理目的和权限合法；CRM 字段、SLA、区域/行业规则和人工接管明确。
关键输入	表单、活动、内容来源、公司公开信息、历史互动、CRM 状态。
动作链	去重→合法补全→意向摘要→规则+模型评分→分配负责人→生成下一步建议→提醒→回写。
人工关口	敏感属性、战略客户、投诉、未成年人或低置信度线索由人工确认。
成功指标	首响时间、有效补全率、分配准确率、MQL→SQL、超时率。

主要风险 隐私越界、错误合并、代理变量歧视、销售过度依赖评分。

不适用边界 不得从非授权来源购买或推断敏感个人信息；分数不得成为唯一拒绝依据。

6. 客户旅程下一步与触达草稿

推荐自主性

A1/A2 为主；服务提醒等低风险事务消息可 A3。

要素 场景设计

适用前提 客户身份、同意/退订、渠道偏好、频控和旅程阶段可用。

关键输入 客户事件、购买/服务历史、内容偏好、库存、权益、客服记录。

动作链 更新旅程状态→识别机会/风险→生成下一步动作→检查同意、频控与权益→提交或自动执行低风险触达→记录结果。

人工关口 优惠、退款、投诉、流失挽回、敏感个人信息和高价值客户由人工批准。

成功指标 触达覆盖、退订/投诉、旅程推进率、复购、人工接管质量。

主要风险 骚扰、错误权益、过度画像、跨渠道频次冲突、敏感事件误触达。

不适用边界 无合法触达基础、退订或客户处于争议状态时禁止自动触达。

7. 活动战役协调与执行台

推荐自主性

A3 内部协调；A2 对外动作。

要素 场景设计

适用前提 目标、里程碑、资产清单、角色和审批流程明确。

要素	场景设计
关键输入	Campaign Brief、预算、渠道计划、供应商、资产、依赖、风险与 KPI。
动作链	拆解任务→识别依赖→生成 workback→提醒负责人→收集状态→冲突与延期预警→资产审批→日终战报→复盘。
人工关口	预算、对外承诺、关键资产、延期导致范围变化时由战役负责人决定。
成功指标	里程碑达成、等待时间、依赖冲突、资产一次通过率、预算偏差。
主要风险	把计划更新误当实际完成、供应商状态不真实、过度提醒。
不适用边界	不能替代供应商合同、现场安全和业务负责人决策。

8. 销售赋能与商机推进

推荐自主性

A3 内部准备与记录；A1/A2 外部沟通。

要素	场景设计
适用前提	产品事实、行业证据、报价权限、销售阶段和 CRM 记录完整。
关键输入	客户会议纪要、公开资料、需求、产品资料、案例、报价政策。
动作链	会前研究→问题清单→会后摘要→需求与异议分类→生成跟进资料→提醒下一步→CRM 更新。
人工关口	报价、承诺、合同条款、竞争性指控和战略客户沟通由销售负责人确认。
成功指标	准备时间、CRM 完整率、跟进时效、阶段推进、资料使用率。
主要风险	虚构客户事实、过度承诺、把会议推测写成事实、泄露跨客户信息。
不适用边界	不自主发送最终报价和合同；不把未经确认的需求写成客户承诺。

9. 营销绩效异常诊断与周度经营复盘

A3 生成与监测；A1 决策建议。

要素 场景设计

适用前提 指标定义、数据血缘、预算、毛利和渠道口径统一。

关键输入 渠道、内容、线索、销售、收入、成本、库存与客户反馈数据。

动作链 数据质量检查→异常检测→分解驱动因素→提出可证伪假设→生成复盘→建议实验→写回学习库。

人工关口 归因结论、预算再分配、暂停渠道和重大策略调整由管理层决定。

成功指标 报告周期、数据差异、假设验证率、决策采纳、重复问题减少。

主要风险 把相关性当因果、忽略线下因素、指标口径漂移、自动合理化失败。

不适用边界 数据口径不一致时只报告差异，不给确定性因果结论。

10. 品牌合规与内容治理守门

推荐自主性

A3 预检与阻断；A2 人工批准后放行。

要素 场景设计

适用前提 品牌、法律、平台、行业规则已结构化并有维护责任人。

关键输入 内容、素材、权利证明、事实来源、渠道、地区、目标人群。

动作链 规则预检→敏感词/声明/权利/标识检查→风险分级→低风险放行或提交审查→记录证据→规则反馈。

人工关口 法律判断、灰色区域、重大品牌承诺和规则冲突由法务/品牌批准。

成功指标 拦截准确率、误报、漏报、审核周期、事故率、规则更新时效。

主要风险 规则过时、过度拦截、模型误解法律、把预检当法律意见。

不适用边界 不能作为最终法律意见；高风险内容不得仅凭 agent 放行。

10.11 场景组合建议

竹势 AI 营销智库出品

企业类型	首批组合	理由
B2B/工业品	市场信号 + 内容 Brief + 线索分配 + 销售赋能	数据与流程相对可控，价值可连接商机推进
消费品牌	内容供应链 + 社媒排期 + 投放诊断 + 品牌守门	高频内容与投放反馈形成闭环，但发布与预算保留审批
连锁/本地生活	活动协调 + 门店内容包 + 低风险事务触达 + 周复盘	总部规则可标准化，多门店执行可观测
平台/互联网服务	客户旅程 + 线索/用户分层 + 异常诊断 + 内容治理	事件数据丰富，但隐私与自动决策风险更高
强监管行业	研究、内部知识、合规预检、内部复盘	先限制在内部与只读，外部动作保持低自主

CHAPTER 11

11 成熟度模型与规模化门槛

01
单点辅助
02
受控编排
03
审批执行
04
有条件自主
05
组合运营

成熟度不以 agent 数量衡量，而以可观察的业务闭环、权限控制、评估、价值兑现和组织能力衡量。最高等级仍保留人类责任。

11.1 Agentic Marketing 五级成熟度模型

等级	可观察状态	证据	下一步
L0 工具化试用	员工各自使用生成工具；无统一数据、权限和评估	以内容数量或个人感受评价；数 据可能复制到公共工具	停止扩张，建立使用政策和 场景盘点
L1 受治理 Copilot	统一账号/平台；品牌知识和基础审 批；以草稿和建议为主	有负责人、基本日志、样本评 估；不授予高风险写权限	选择 1-2 个流程做端到端重 构
L2 Agent-ready 流程	任务、状态、工具和人工关口被显 式设计；可运行多步流程	连接 1-3 个系统；端到端评估； 异常可恢复；有基线 ROI	在窄权限下上线 A2，并积 累运行证据

竹势 AI 营销智库出品

等级	可观察状态	证据	下一步
L3 有条件自主	部分高频、可逆任务在白名单、阈值和监控下 A3 运行	独立 agent 身份、实时监控、事故响应、连续达标周期、财务复核	扩大场景前先证明复用和治理能力
L4 组合化运营	多个场景共享平台、知识、连接器、评估和治理；按价值管理组合	季度自主性关卡、统一成本与价值看板、跨业务复用、独立审计	持续优化组合，淘汰低价值 agent，避免无边界自治

11.2 各维度可观测条件

	维度 L1	L2	L3	L4
业务	有场景负责人	有端到端验收	稳定达到业务门槛	按组合优化资本配置
流程	有审批与模板	状态、异常、回滚明确	例外驱动监督	跨流程复用标准
数据	基本知识库	数据接口与质量责任	实时/事件数据可用	统一语义与血缘
权限	人类账户或无写权限	独立身份、A2	阈值内 A3	集中治理多 agent 身份
评估	抽样检查	离线 + 端到端评估	在线监控与对抗测试	持续评估、独立审计
价值	时间/质量基线	试点 ROI	财务复核净价值	组合级预算和淘汰
组织	培训与政策	产品/流程负责人	Agent Ops 运行机制	董事会/高管治理关卡

11.3 晋级门槛

晋级	必须满足	不得存在
L1→L2	明确业务结果、流程负责人、数据合法性、端到端评估和人工接管	用聊天结果代替业务状态；无人承担结果
L2→L3	连续多个周期达到质量/SLA/成本门槛；独立身份；最小权限；熔断与回滚演练	重大越权、无法定位故障、审批可被绕过

晋级	必须满足	不得存在	竹势 AI 营销智库出品
L3→L4	多个场景复用平台与治理；财务确认价值；季度组合管理；独立审计	agent 蔓延、重复采购、成本不可归因、低价值场景不退役	

11.4 停止条件

- 连续两个评估周期未达到最低业务或质量门槛，且改进成本超过预期价值。
- 发生重大品牌、个人信息、财务或合规事故，控制有效性无法及时恢复。
- 数据来源、处理目的或供应商条件发生变化，导致合法性或可用性基础消失。
- 人工审核与异常处理成本长期高于被替代或新增的价值。
- 无法证明 agent 结果优于简化流程、规则自动化或普通 Copilot。
- 责任人变更后无人愿意接管流程和风险。

成熟度判断

L4 不是“市场部无人化”。它表示企业能以统一治理和财务纪律管理一组不同自主等级的 agent，并能及时降级、停机和退役。

CHAPTER 12

12 十八个月实施路线图

0—90 天

选准任务，建立权限与评估基线

3—6 个月

跑通闭环，验证价值与异常处理

6—18 个月

扩展组合，形成稳定运营机制

路线图以业务闭环、风险证据和组织能力为关卡。先证明一个窄场景，再扩展数据、权限和组合，不以平台覆盖范围作为首要里程碑。

12.1 0-90 天：证明一个受控闭环

阶段	关键动作	里程碑	停止条件
0—30 天	高管定风险偏好；盘点任务；选场景；建立基线；指定责任人；完成数据/权限审查	项目章程、准入评分、基线、目标、风险登记	无法定义结果；数据不合法/不可用；无人负责
31—60 天	流程重构；构建知识、工具和审批；离线评估；影子运行；红队	端到端最小可行链路、评估报告、运行手册	端到端成功率低；工具不稳定；无法回滚
61—90 天	限定用户/渠道上线；每日监控；每周复盘；财务初验；决定继续/收缩/停止	连续运行证据、初步净价值、权限建议、90 天复盘	重大事故；人工成本过高；价值假设不成立

12.2 3-6 个月：形成可复制能力

- 把首个场景的工具、身份、审批、日志、评估和成本计量从项目代码转为平台能力。
- 增加第二个相邻场景，验证数据与连接器复用；不要同时进入多个新渠道。
- 建立 Agent Ops 周会、月度价值复盘和季度自主性关卡。
- 形成模型与供应商替换机制、数据保留与删除、事故响应演练。
- 对达到门槛的可逆任务，从 A2 申请局部 A3；扩大权限而非扩大口号。

6 个月门槛 要求

价值	至少一个场景由财务确认净价值或明确的战略期权价值
可靠性	端到端质量与 SLA 连续达标，长尾失败可定位
治理	独立身份、最小权限、审批、熔断、回滚和事故演练通过
复用	至少一个工具、知识或评估组件被第二场景复用
组织	流程负责人、产品负责人、数据、平台和风险角色持续运作

12.3 6-18 个月：从场景扩展到组合运营

时间	重点	管理层决策
7-9 个月	扩展到 3-5 个相邻场景；统一身份、日志、成本和知识治理	确定企业级平台边界与供应商策略
10-12 个月	1 个 建立组合价值看板；淘汰低价值场景；开展独立审计	决定哪些业务单元复制、哪些权限扩大
13-15 个月	1 个 跨渠道/跨区域扩展；加强数据语义与事件驱动；完善模型替换	批准更高自主等级的任务清单与预算阈值
16-18 个月	1 个 形成年度 agent 组合规划、资本配置、风险与人才机制	把 agent 纳入经营预算、内控和组织设计

12.4 十二个关键里程碑

1. 董事会/CEO 批准风险偏好、禁止清单与责任原则。
shichangbu.ai

- 2. CMO 选定首个经营结果和流程负责人。
- 3. 数据负责人确认来源、权限、质量和保留规则。
- 4. 完成现状基线与对照方案。
- 5. 完成任务合同、人工关口、异常和回滚设计。
- 6. 完成独立 agent 身份与最小权限。
- 7. 完成离线、端到端和对抗评估。
- 8. 完成影子运行和人工接管演练。
- 9. 完成限定范围生产上线。
- 0. 完成 90 天财务、业务和风险复盘。
- 11. 第二场景验证复用，并建立 Agent Ops 机制。
- 2. 18 个月形成组合级价值、治理、预算和退役机制。

12.5 董事会、CEO、CMO 下一步的关键决策

决策主体	未来 30 天必须决定
董事会/风险委员会	风险偏好、禁止自主清单、重大事故上报、投资监督与审计要求
CEO	业务优先级、跨部门责任、价值承接、平台与供应商原则
CMO	首批场景、流程重构、品牌与客户边界、指标和人工审批
CIO/数字化负责人	架构、身份、连接器、可观测、模型可替换与运行可靠性
法务/数据/安全	个人信息、内容标识、知识产权、供应商、红队和事故响应
业务单元负责人	本地流程、数据质量、用户采用、人工接管和财务结果

最终建议

用 90 天证明一个窄场景，用 6 个月证明可复制，用 18 个月证明组合级经营与治理。任何阶段都允许降级、暂停和退役。

13 结语：管理层必须保留的三项权力

技术会持续变化，模型能力和成本也会变化。企业真正需要稳定保留的是目标权、授权权和停止权。

目标权

决定企业追求什么、牺牲什么以及不能用什么手段获得结果。Agent 可以拆解目标和寻找路径，但不能代替管理层处理品牌、利润、客户权益和长期能力之间的冲突。

授权权

决定哪个 agent 可以访问哪些数据、调用哪些工具、在多大金额和频次内执行、何时必须请求批准。授权应建立在运行证据上，并且可以随时收回。

停止权

决定在质量下降、风险变化、供应商变更或价值不成立时暂停、降级和退役。没有技术和组织上的停止权，自主性就会变成不可管理的路径依赖。

结论

Agentic Marketing 的竞争优势不来自“让 AI 做得最多”，而来自企业能更快地把正确任务交给合适的执行者，并对目标、边界、证据和结果保持清晰控制。

APPENDIX A

附录 A 90 天启动清单

0—90 天

选准任务，建立权限与评估基线

3—6 个月

跑通闭环，验证价值与异常处理

6—18 个月

扩展组合，形成稳定运营机制

序号	类别	行动	责任人	时间	完成
01	高管与价值	确认经营问题与价值承接人	业务赞助人	第 1-10 天	☐
02	高管与价值	批准风险偏好、禁止清单与停止条件	CEO/CMO/风险	第 1-15 天	☐
03	高管与价值	建立业务、质量、风险、成本基线	流程负责人/财务	第 1-20 天	☐
04	场景与流程	完成任务准入评分	流程负责人	第 1-15 天	☐
05	场景与流程	画出当前流程与目标流程	流程/运营	第 5-25 天	☐
06	场景与流程	定义触发、输入、动作、关口、输出、异常	产品负责人	第 10-30 天	☐
07	数据	完成数据来源、权限、质量和保留盘点	数据负责人	第 5-25 天	☐
08	数据	建立品牌/产品/政策知识基线	品牌/业务	第 10-35 天	☐
09	数据	建立金标准评估集与边界样本	流程/评估	第 20-45 天	☐
10	技术	创建独立 agent 身份与最小权限	平台/安全	第 15-40 天	☐

竹势▲营销智库出品

序号	类别	行动	责任人	时间	完成
11	技术	完成工具参数校验、幂等、超时与回滚	平台/集成	第 20–50 天	☐
12	技术	完成日志、追踪、成本和告警	平台/运维	第 25–55 天	☐
13	治理	完成隐私、IP、内容标识和供应商审查	法务/数据/品牌	第 15–45 天	☐
14	治理	完成红队、越权、注入和异常测试	安全/评估	第 35–60 天	☐
15	治理	完成停机、人工接管和事故演练	全体	第 45–65 天	☐
16	上线	完成影子运行和差异分析	运营/产品	第 45–70 天	☐
17	上线	限定用户/渠道/预算生产上线	业务赞助人	第 60–75 天	☐
18	上线	每日异常、每周 Ops、每月价值复盘	产品/流程/财务	第 61–90 天	☐
19	复盘	完成 90 天业务、财务、风险结论	业务赞助人/财务	第 85–90 天	☐
20	复盘	决定继续、扩大、降级、暂停或退役	CEO/CMO	第 90 天	☐

APPENDIX B

附录 B CMO 决策清单

决策问题	是/否	需要的证据
场景是否连接收入、毛利、费用、客户体验或重大风险？	☐	价值逻辑与财务科目
当前流程是否已经被重构，而非在旧流程上增加 agent？	☐	现状/目标流程与删除步骤
人类最终负责点和强制审批点是否明确？	☐	流程图与责任说明
所有公开事实和品牌承诺是否有可信来源？	☐	知识库与来源记录
数据使用是否满足目的、最小必要、权限和保留要求？	☐	数据盘点与法务/数据批准
Agent 是否有独立身份与最小权限？	☐	权限清单与凭证机制
错误是否能在可接受时间内被发现、停止和回滚？	☐	SLA、熔断与演练记录
评估是否覆盖正常、边界、恶意输入和工具故障？	☐	评估报告与红队报告
ROI 是否包含人工审核、治理和残余风险成本？	☐	完整商业案例
释放的时间是否有明确承接动作？	☐	招聘/外包/实验/销售计划
扩大自主性是否基于连续运行证据？	☐	运行看板与关卡记录
是否有退出供应商、迁移模型和退役 agent 的方案？	☐	替代与退出计划

APPENDIX C

附录 C 项目准入评分表

维度	权重	1 分	3 分	5 分	得分
业务价值	20%	无清晰承接	影响效率或局部质量	连接收入/毛利/重大成本/风险	
频率与规模	15%	低频一次性	每月/单团队	每日/每周且覆盖规模大	
标准化	15%	成功无法定义	部分规则可描述	输入、结果、例外清晰	
数据可得性	15%	缺失/来源不清	需较多人工补充	稳定、合法、质量可控	
可逆性	10%	不可逆	可部分补救	完全可回滚/重做	
权限可控	10%	需广泛管理员权限	可通过审批控制	只读/窄写权限即可	
可观测性	15%	过程与结果不可见	结果可验但过程有限	步骤、结果、成本、异常均可追踪	
风险惩罚	否决/扣分	高风险不可逆	中风险有控制	低风险	

否决项

- 无明确业务负责人；关键数据缺少合法来源；无法定义正确结果；无法人工接管；高风险不可逆且无专门审批；不能记录执行证据；预期价值低于治理与运行成本。

APPENDIX D

附录 D RACI

R=执行，A=最终负责，C=协商，I=知会。每项只能有一个 A。

活动	业务赞助人	流程负责人	Agent	产品	平台/IT	数据	品牌/法务/风险
场景选择与价值目标	A	R	C	C	C	C	
流程与人工关口设计	C	A/R	R	C	C	C	
数据与知识治理	I	C	C	C	A/R	C	
工具、身份与权限	I	C	C	A/R	C	C	
评估与验收	C	A	R	R	C	C	
上线批准	A	R	R	C	C	C	
日常运行	I	A/R	R	R	C	C	
高风险审批	A/C	R	I	I	C	A/R	
事故响应	A	R	R	R	C	R	
价值复盘与扩张	A	R	R	C	C	C	
暂停与退役	A	R	R	R	C	C	

APPENDIX E

附录 E 风险登记表

ID	风险	触发/指标	固有影响	预防控制	检测/停止	责任人	残余风险
R01	事实错误进入公开内容	来源缺失、置信度低、审核漏过	品牌/法律	可信来源、字段验证、审批	抽检、投诉监测、撤回	品牌负责人	
R02	个人信息超目的处理	新增字段/新用途/跨境	监管/客户权益	目的绑定、最小必要、权限	访问日志、DLP、停用数据源	数据负责人	
R03	提示注入导致越权	外部内容含指令、异常工具调用	安全/数据	数据与指令隔离、工具独立授权	注入检测、工具熔断	安全负责人	
R04	预算或投放误操作	异常幅度、重复执行、归因错误	财务/增长	金额阈值、幂等、审批	实时预算告警、回滚	预算负责人	
R05	错账号/重复发布	账号不匹配、内容重复	ID品牌/渠道	账号绑定、预览、幂等	发布监测、撤回	渠道负责人	
R06	IP/素材权利不足	来源未知、许可过期、相似度高	法律/品牌	资产台账、许可检查	抽检、权利投诉流程	法务/内容	
R07	多 agent 级联失败	循环交接、状态不一致、成本飙升	运营/安全	最大深度、交接合同、独立验证	trace、成本告警、全局停机	平台负责人	
R08	供应商/模型变化	版本升级、价格/条款/地区变化	连续性/合规	可替换架构、合同与回归测试	质量漂移、切换备用	CIO/采购	
R09	人工审批失效	积压、橡皮图章、越权审批	治理/质量	角色、SLA、双人审批	审批质量抽检、暂停权限	流程负责人	
R10	价值未兑现	调用上升、净价值不变	财务/战略	基线、承接动作、月度复盘	财务复核、停止关卡	业务赞助人	

APPENDIX F

附录 F 上线前检查表

类别	检查项	状态	证据/负责人
业务	目标、用户、范围、成功标准和停止条件已批准	☐	
业务	业务负责人和人工接管人已排班	☐	
流程	触发、状态、审批、异常、回滚和退役已设计	☐	
数据	数据来源、目的、权限、质量、保留和删除已批准	☐	
知识	品牌/产品/政策资料有版本、来源和有效期	☐	
模型	模型/版本、路由、替代与回归测试通过	☐	
工具	工具白名单、参数校验、幂等、超时和限流通过	☐	
身份	独立 agent 身份、最小权限、短期凭证和吊销机制完成	☐	
评估	金标准、边界、恶意输入、工具故障和端到端评估通过	☐	
安全	提示注入、越权、数据泄露、级联和成本攻击测试通过	☐	
治理	内容标识、隐私、IP、平台规则和供应商审查完成	☐	
运行	日志、trace、成本、质量、风险告警和看板可用	☐	
恢复	熔断、人工接管、回滚、备份和事故响应演练完成	☐	
用户	培训、操作手册、限制说明和反馈渠道就绪	☐	
财务	成本预算、价值基线、归因和财务复核人确认	☐	

APPENDIX G

附录 G ROI 计算表

变量	符号/公式	企业输入	证据来源
年度合格任务量	Q		基线与生产日志
单任务节省的有效人工小时	H		时间研究/抽样
全成本时薪	C_h		财务/HR
时间兑现系数	R_c		招聘、外包或产能计划
释放产能价值	$Q \times H \times C_h \times R_c$		计算
增量转化数量	ΔN		对照/归因
单笔增量毛利	GM		财务
增量毛利价值	$\Delta N \times GM$		计算
周期压缩价值	V_t		上市/响应机会模型
风险避免价值	V_r		事故情景与控制证据
期权价值	V_o		新增实验与成功率
模型与平台运行成本	C_run		账单/成本系统
人工审核与运维成本	C_ops		工时/财务
集成、数据与变革摊销	C_impl		项目预算

		企业输入	竹势 AI 营销智库出品 证据来源
变量	符号/公式		
残余风险成本	C_risk		风险登记
年度净价值	$Q \times H \times C_h \times R_c + \Delta N \times GM + V_t + V_r + V_o - C_{run} - C_{ops} - C_{impl} - C_{risk}$		计算
ROI	年度净价值 ÷ 年度总投入		计算
回收期	累计净现金流转正所需月数		计算
所有参数均应由企业填入。对于无法建立可靠证据的项目，建议使用保守区间并进行敏感性分析，不使用供应商宣传数字直接入账。			
APPENDIX H			

附录 H 术语表

术语	定义
Agent / AI agent	能够围绕目标规划、调用工具、观察结果并在边界内持续执行多步任务的软件系统。
Agentic Marketing	把 agent 作为受授权执行者纳入营销经营、流程、数据、技术、治理和财务体系的运行方式。
Copilot	在人类主导的工作界面中提供建议、草稿或局部操作辅助的系统。
Workflow / workflow	由预设步骤、规则、人工节点和系统调用构成的可重复流程。
Orchestration / 编排	协调模型、agent、工具、规则、状态和交接的运行机制。
Tool calling / 工具调用	模型或 agent 以结构化参数调用外部函数、API 或应用完成读写动作。
Memory / 记忆	用于保存任务过程或稳定偏好的信息；不同于经治理的知识库和真实业务状态。
Human in the Loop	在指定节点由人类审查、批准、纠正或接管。
Human on the Loop	系统在边界内运行，人类通过监控和例外介入监督。
Guardrail / 护栏	输入、输出、权限、工具、预算、内容和运行时的限制与阻断机制。
Evals / 评估	用定义明确的样本、指标和通过标准测试组件、轨迹与端到端任务。
Trace / 追踪	记录一次 agent 运行的模型、步骤、工具、输入输出、状态、成本、批准和错误。
Reversibility / 可逆性	错误动作能否被撤销、回滚或重做，以及恢复成本大小。
Residual risk / 残余风险	实施控制后仍然存在的风险。
Agent identity / agent 身份	分配给 agent 的唯一非人身份，用于授权、审计、暂停和退役。

APPENDIX I

附录 I 来源索引

编号	来源
S01	OpenAI, A Practical Guide to Building Agents, 2025. 官方业务与工程指南 (openai.com)。
S02	OpenAI, Agents SDK Documentation, 2025–2026. 官方开发文档 (developers.openai.com)。
S03	Anthropic, Measuring AI Agent Autonomy in Practice, 2026. 官方研究 (anthropic.com)。
S04	McKinsey & Company, The State of AI in 2025: Agents, Innovation, and Transformation, 2025. 全球调查, 1,993 名受访者、105 个国家/地区。
S05	IBM Institute for Business Value / Oxford Economics, The 2025 CMO Study, 2025. 调查 1,800 名营销与销售高管, 33 个地区、24 个行业。
S06	Microsoft, 2026 Work Trend Index Annual Report: Agents, Human Agency, and the Opportunity for Every Organization, 2026. 20,000 名使用 AI 的工作者及匿名生产力信号。
S07	Stanford Institute for Human–Centered AI, AI Index Report 2026, 2026. 斯坦福年度综合报告。
S08	Brynjolfsson, Li & Raymond, Generative AI at Work, Quarterly Journal of Economics, 2025 (早期工作论文 2023) . 5,172 名客户支持人员现场研究。
S09	U.S. National Institute of Standards and Technology, Artificial Intelligence Risk Management Framework (AI RMF 1.0), 2023. NIST AI 100–1。
S10	U.S. National Institute of Standards and Technology, AI RMF: Generative Artificial Intelligence Profile, 2024. NIST AI 600–1。
S11	International Organization for Standardization, ISO/IEC 42001:2023, Artificial Intelligence Management System. 官方标准说明。
S12	International Organization for Standardization, ISO/IEC 23894:2023, Artificial Intelligence — Guidance on Risk Management. 官方标准说明。

编号	来源
S13	OWASP GenAI Security Project, Agentic AI — Threats and Mitigations, 2025. Agentic Security Initiative。
S14	OWASP GenAI Security Project, OWASP Top 10 for Agentic Applications 2026, 2025. 全球同行评议安全框架。
S15	Cemri et al., Why Do Multi-Agent LLM Systems Fail?, 2025. 对五种多 agent 框架、150 余项任务的失效研究。
S16	Stauffer et al., The 2025 AI Agent Index: Documenting Technical and Safety Features of Deployed Agentic AI Systems, 2026. 对 30 个已部署 agent 的公开信息研究。
S17	Lynch et al., Agentic Misalignment: How LLMs Could Be Insider Threats, 2025. 受控假设情境中的 16 个模型压力测试；作者明确未观察到真实部署证据。
S18	《中华人民共和国个人信息保护法》，2021 年 8 月 20 日通过，2021 年 11 月 1 日施行。中国人大网/国家互联网信息办公室公开文本。
S19	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 年 7 月 13 日发布，2023 年 8 月 15 日施行。
S20	国家互联网信息办公室等四部门，《人工智能生成合成内容标识办法》，2025 年 3 月 14 日发布，2025 年 9 月 1 日施行。
S21	European Union, Regulation (EU) 2024/1689 (Artificial Intelligence Act), 2024; European Commission AI Act implementation page, updated July 2026.。
S22	C2PA, Coalition for Content Provenance and Authenticity Technical Specification, 2024–2026. 内容来源与真实性开放技术规范。
S23	OECD, OECD AI Principles, adopted 2019, updated 2024. 可信 AI 政策原则。
S24	Capgemini Research Institute, Rise of Agentic AI: How Trust Is the Key to Human–AI Collaboration, 2025. 全球企业调查与经济价值估算。
S25	McKinsey & Company, State of AI Trust in 2026: Shifting to the Agentic Era, 2026. AI 风险、事件与治理调查。
S26	IBM, 2025 CEO Study / Enterprise AI announcements, 2025. 关于 AI 投资、集成与 ROI 的官方研究与披露。

来源		竹势 AI 营销智库出品
编号		
S27	Grewal et al., How Generative AI Is Shaping the Future of Marketing, Journal of the Academy of Marketing Science, 2025. 营销学术综述。	
S28	Noy & Zhang, Experimental Evidence on the Productivity Effects of Generative Artificial Intelligence, Science, 2023. 专业写作随机实验。	
S29	Dell’Acqua et al., Navigating the Jagged Technological Frontier, 2023. 知识工作者生成式 AI 实验研究。	
S30	Anthropic, Anthropic Economic Index and related autonomy/productivity research, 2025–2026. 基于匿名使用数据的官方研究。	
S31	MIT Sloan School of Management, Agentic AI, Explained, 2026. 面向管理者的概念与应用说明。	

来源使用原则

本报告对调查数字保留其样本与口径，不把不同研究直接拼接为统一市场规模；对技术能力区分产品文档、实验研究与真实部署；对场景设计采用“场景原型”表述，不把缺少公开证据的方案写成客户案例。

竹势 AI 营销智库

2026 年 7 月

PUBLISHER

关于竹势 AI 营销智库

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。