

出版说明

竹势 AI 营销智库出品

项目	内容
报告名称	《2030营销未来图景：人、AI员工与智能商业》
出品方	竹势 AI 营销智库
版本	2026年7月正式版
核心读者	企业董事长、CEO、CMO、市场负责人、数字化负责人、品牌负责人、代理商与媒体经营者
报告性质	未来情景与经营决策工具
引用原则	法规、政府与国际组织、同行评审研究、上市公司披露、平台官方材料和具名机构研究优先
重要边界	2030相关内容均按结构变化、关键不确定性、条件情景和触发信号表达，不构成精确预测或投资承诺。

版权与使用

本报告可用于企业内部战略讨论、经营规划和能力建设。引用时请注明“竹势 AI 营销智库，《2030营销未来图景：人、AI员工与智能商业》，2026年7月”。报告中的第三方商标、案例与数据归原权利人所有。

阅读方法

董事会可先阅读执行摘要、四情景矩阵和行动组合；经营团队再按AI员工、组织、客户入口、品牌信任、媒体、代理商与决策系统逐章落实。每章均按“判断—机制—证据—案例—边界—管理含义”组织。

执行摘要：十个可证伪判断

第一章 2030未来研究的边界与纪律

第二章 六股相互作用的结构力量

第三章 AI员工：执行权、授权半径与人工接管

第四章 人机组织：任务、岗位与决策权重构

第五章 客户入口：机器中介与客户关系资产

第六章 品牌信任：内容过剩时代的证据链

第七章 媒体与注意力：自动化交易和利润目标

第八章 代理商价值：执行商品化后的新合约

第九章 智能商业决策：速度、质量与责任接口

第十章 四种2030营销经营情景

第十一章 情景切换雷达与领先指标

第十二章 2026—2030行动路线

董事会与经营层工具箱

附录A 案例卷宗

附录B 专家与经典思想转译

附录C 来源索引

执行摘要：十个可证伪判断

竹势 AI 营销智库出品

BOARD FORESIGHT / 核心判断

2030 不是一个等待命中的预测，而是一组需要今天设置触发门、责任人和可逆投资的经营情景。

高确定性底座

关键不确定性

领先指标

无悔行动

到2030年，营销的核心竞争不会由“谁生成更多内容”决定，而由谁能在可验证、可逆和可追责的条件下，把AI接入真实经营动作，并保有客户关系、品牌证据和独立测量能力决定。

核心判断	可检验边界
1. AI执行权将扩大，但不会均匀扩大	高确定性变化是AI进入受约束执行链；可证伪条件是：到2030年企业关键营销动作仍主要停留在建议和草稿阶段，且没有可审计的工具调用、权限与回滚机制。
2. 岗位结构变化快于岗位总量消失	任务被拆分、重组和重新定价的速度，将高于完整岗位被端到端替代的速度；若多数营销职业在统计上出现持续、广泛、直接净消失且任务重组不显著，该判断被推翻。
3. 客户入口将多中心并存	传统搜索、生成式答案、推荐、电商、代理式采购与品牌自有触点会长期并存；若单一代理入口在主要品类持续占据绝对多数发现与交易，该判断需修正。
4. 品牌资产将增加“机器可读证据层”	产品事实、来源、库存、服务承诺、资质、评价与纠错记录将成为品牌基础设施；若机器中介不影响可见度、转化或信任，投入优先级应下降。
5. 透明只是信任的必要条件之一	披露AI使用不会自动产生信任；履约一致性、可验证来源、人类责任与纠错速度更关键。
6. 媒体自动化程度将上升，企业测量主权更稀缺	平台算法会扩大自动出价、创意组合与跨渠道分配，但平台最优目标与企业利润目标不会天然一致。
7. 代理商工时价值承压，责任价值上升	标准执行、基础制作和常规优化更易商品化；问题定义、跨域整合、专有数据、变革落地和结果责任将成为溢价来源。
8. 决策速度上升不等于决策质量上升	数据质量、目标函数、例外处理、反馈延迟和责任接口决定价值上限。
9. 监管将首先改变高风险动作的设计，而非阻止全部创新	可追踪、告知、内容标识、人工监督和风险管理将逐步进入产品与流程要求。
10. 最优战略是分层下注	企业应先做无悔能力，再以小规模期权验证不可逆投资；当领先指标越过决策门时再扩张。

表0-1 十个可证伪核心判断

董事会现在应决定的五件事

- 确定哪些营销动作可以由AI建议、起草、执行、承诺或交易，并为每一级定义责任人。
- 把客户入口资产、品牌证据链和独立测量能力列入核心资产清单。

- 停止以“工具上线数”和“内容产量”替代经营结果。
- 要求任何代理商或平台自动化方案公开目标函数、数据边界、异常控制和退出机制。
- 建立季度情景雷达：技术能力、监管、客户行为、平台集中度、组织采用与信任事件共同触发决策。

第一章 2030未来研究的边界与纪律

竹势 AI 营销智库出品

把未来判断拆成五层

不同类型的陈述，使用不同证据和决策方式。

01

已发生事实

可验证

02

结构变化

可外推但有限

03

关键不确定性

需持续观察

04

情景

内部一致

05

行动假设

可证伪、可退出

2030不是一个可被单点预测的技术终局，而是一个适合暴露假设、设置触发条件和管理战略期权的经营时间窗。

Shell情景规划的价值并非预测唯一未来，而是用多种内在一致的环境组合检验战略韧性。Philip Tetlock关于预测校准的研究提醒管理者：预测质量来自概率更新、分解问题、记录证据和纠错，而非宏大叙事。两者结合后，本报告采用“结构变化—关键不确定性—条件情景—领先指标—动作组合”的方法。

1.1 五类陈述必须分开

类型	定义	报告写法	管理用途
已发生事实	截至2026年7月可由公开材料验证	写明时间、主体、样本与口径	确认基线
高确定性结构变化	已有多重机制推动，方向稳定但速度不确定	给出反例与证伪条件	无悔投资
关键不确定性	方向、速度或治理结果存在显著分叉	不强行给精确概率	期权布局
条件性推演	只有在一组条件同时成立时出现	列必要条件和触发信号	条件行动
规范性建议	基于风险偏好和经营目标提出	明确价值判断与退出条件	治理与资源配置

表1-1 未来研究的五类陈述

1.2 情景法的适用边界

情景法不能替代预算、统计预测或行业需求模型。对于短周期投放、销售漏斗和库存决策，应使用实验与预测模型；对于跨技术、监管、组织和客户行为的联动问题，情景法更适合。情景也不是四个故事的装饰，它必须改变资源分配、合同条款和退出条件。

1.3 经典思想转译

思想	原意	AI营销连接	修正与管理判断
Shell情景规划	用多种一致未来检验战略	检验AI自主权与治理基础的组合	不为情景分配伪精确概率；关注切换信号
Tetlock预测校准	分解、更新、记录概率与误差	把“AI将取代营销”拆成任务与时间窗	季度更新假设，保留预测日志
Simon有限理性	决策者受信息与计算限制	AI扩展搜索与生成，但目标仍有限	速度提升不消除目标偏差
实物期权	不确定环境下保留扩张、推迟与退出权	先试点、可逆集成、分阶段授权	为学习付费，而非提前锁定终局

表1-2 未来研究的理论底座

管理含义：任何2030战略提案都应附带四个字段——“成立条件、反证、领先指标、退出动作”。没有这些字段的预测不应进入重大资本决策。

第二章 六股相互作用的结构力量

竹势 AI 营销智库出品



AI 执行权

人机组织

客户入口

品牌信任

媒体交易

代理商价值

2030营销图景由六股力量共同形成：AI执行权、客户入口、品牌信任、媒体交易、组织重构和监管责任。单独观察任何一股力量都会高估技术、低估系统。

结构力量	主要驱动	经营机会	主要风险
AI执行权	模型推理、工具调用、工作流与记忆	更快、更便宜的动作闭环	错误扩散、权限滥用、目标错配
客户入口	生成式答案、推荐、电商与购买代理	决策摩擦下降	品牌失去可见接触与第一方关系
品牌信任	来源、履约、纠错与人类责任	可信信号更有溢价	合成内容稀释真实性
媒体交易	自动出价、创意组合、跨渠道优化	实时配置和长尾实验	平台目标替代企业利润目标
组织重构	任务拆分、管理跨度、技能与责任接口	小团队可管理更大动作量	隐形劳动、技能退化、责任真空
监管责任	风险分级、透明、数据与消费者保护	建立可信市场规则	合规碎片化和创新摩擦

表2-1 六股结构力量及其联动

2.1 反馈回路

执行权越大，企业越需要可观测性和责任接口；机器中介越强，品牌越依赖结构化证据；媒体交易越自动化，独立增量测量越重要；内容供给越丰富，可信来源和稳定履约越稀缺。治理因此不是创新之后的补丁，而是扩大授权半径的前置条件。

2.2 结构变化与关键不确定性分层

层级	项目	判断	观察重点
高确定性	生成和基础执行成本下降	方向明确，幅度和行业差异不确定	单位任务成本、人工复核率
高确定性	多入口并存	入口组合变化，但不会立即统一	发现、比较、购买路径份额
高确定性	可追责要求增强	监管与事故共同推动	日志、披露、人工接管
关键不确定性	通用购买代理普及速度	取决于支付、身份、授权和责任	代理交易占比、退款争议
关键不确定性	平台集中度	可能集中，也可能被互操作标准削弱	流量集中、数据可携带性
关键不确定性	企业是否形成AI原生组织	技术可用不等于组织采用	决策权、管理跨度、岗位重构
条件性	大规模自主媒体交易	需可验证测量与预算护栏成熟	无人干预交易比例、异常损失

表2-2 2030结构变化 / 关键不确定性分层表

第三章 AI员工：执行权、授权半径与人工接管

竹势 AI 营销智库出品

授权半径随风险逐级收紧

自主度不是模型属性，而是业务责任设计。

L0

建议

人类执行

L1

影子

只记录不行动

L2

审批

人类逐项确认

L3

受控自主

边界内执行

L4

禁止自主

不可逆高风险

AI员工的关键不是拟人化名称，而是企业是否把目标、工具、数据和动作权限交给软件系统。授权半径应由可观测性、可逆性、目标清晰度、责任接口和接管能力共同决定。

OpenAI官方把代理定义为代表用户独立完成任务的系统，并强调工具设计、护栏和评估；NIST AI风险管理框架则把治理、测量和管理风险置于持续循环。二者共同说明：能力增长只提供可执行性，组织控制决定可用性。[S01][S02]

3.1 授权半径五因子

因子	低风险特征	高风险特征	设计要求
可观测性	输入、推理依据、工具调用和结果可记录	黑箱动作、外部系统状态不可见	全链路日志、证据引用、状态监控
可逆性	可撤回、回滚、重放	付款、公开承诺、删除、不可逆发布	草稿模式、双重确认、额度与时间窗
目标函数	单一清晰、反馈及时	多目标冲突、代理指标易被操纵	目标分层、利润护栏、负面约束
责任接口	责任人和升级路径明确	供应商、平台、业务方互相推诿	RACI、事件分级、赔付与审计条款
人工接管	可暂停、转人工、恢复	无法中断或接管成本过高	熔断器、异常阈值、人工值守

表3-1 AI员工授权半径五因子

3.2 五级授权模型

级别	AI角色	允许动作	典型场景	决策门
L0 观察	记录与分析	不生成外部动作	舆情、账户诊断	数据质量达标
L1 建议	输出方案	人执行	预算建议、回复建议	建议可验证
L2 起草	生成待审批动作	人批准后执行	内容、邮件、投放变更草稿	审批与版本控制
L3 受约束执行	在额度和白名单内执行	可自动发布或调整	低风险FAQ、低额预算优化	异常可监控可回滚
L4 委托经营	围绕目标跨系统规划执行	多动作连续闭环	成熟场景的日常运营	高可观测、强治理、责任清晰

表3-2 AI员工五级授权模型

案例：Air Canada聊天机器人责任事件

主体与时间：加拿大航空，争议源于2022年客户查询，裁决于2024年2月作出。业务情境：网站机器人提供与正式政策不一致的丧亲票价信息。具体动作：客户依机器人答复购票并事后申请折扣。结果证据：不列颠哥伦比亚省民事解决法庭认定企业应对网站信息负责并判赔。作用机制：机器人输出成为企业对外承诺，但知识更新、校验和责任接口失效。边界：这是特定司法辖区的小额民事裁决，不代表所有司法辖区规则。启示：公开客户承诺属于高风险动作，必须绑定权威知识源、有效期、人工升级和责任主体。[S03]

案例：Klarna从效率高峰转向混合服务

2024年Klarna官方称AI助手首月处理230万次对话，覆盖约三分之二客服聊天，处理时长显著缩短；2025年公司又增加人类服务能力，强调复杂、敏感问题需要更高技能人工。供应商自报的效率数据不能单独证明长期客户价值。启示是按问题复杂度分层，而不是把统一自动化率作为目标。[S04][S05]

3.3 人工接管决策树

问题	是	否
动作是否对外承诺、支付、公开发布或改变客户权益？	进入高风险审查	继续
错误是否可在可接受时间和成本内回滚？	继续	必须人工批准
是否有权威数据源、权限白名单和完整日志？	继续	停留L1/L2
是否能检测异常并立即熔断？	允许小范围L3试点	不得自动执行
责任人、赔付和升级路径是否明确？	按额度扩展	停止授权

表3-3 AI员工授权与人工接管决策树

第四章 人机组织：任务、岗位与决策权重构

竹势 AI 营销智库出品

机器主导 高频 × 可验证 自动执行并抽检
人机共决 多约束 × 可解释 AI准备，人类批准
人类主导 关系 × 重大判断 AI提供证据
停止或重构 低价值 × 不可控 不把坏流程自动化

未来组织的主要变化单位是任务和决策权，而非岗位名称。企业需要重画“谁发现、谁判断、谁执行、谁批准、谁负责”的流程。

ILO 2025年研究使用任务级数据、专家输入和模型预测评估职业暴露，强调生成式AI更可能带来岗位转型而非简单整体消失；OECD研究也指出暴露、收益和风险在技能、收入和职业之间不均匀。[S06][S07]

4.1 人机任务四分法

任务类型	AI相对优势	人类相对优势	组织设计
高频规则任务	速度、一致性、并行	例外判断	AI执行，人监督
信息密集判断	检索、比较、模拟	目标澄清、因果质疑	AI准备证据，人决策
关系与信任任务	上下文提醒、记录	共情、谈判、责任承担	AI辅助，人面对客户
原创与高模糊任务	发散、变体、组合	品味、取舍、文化判断	人定方向，AI扩展探索

表4-1 人机任务与决策权重构画布

4.2 管理跨度会扩大，但监督劳动不会消失

AI可降低信息汇总、排期、跟踪和初稿成本，管理者可能覆盖更多项目与执行单元；同时，评估、纠错、权限设置、知识维护 and 异常处置形成新的“监督劳动”。若组织只计算减少的人时，不计算监督与风险成本，ROI会被高估。

案例：P&G / Harvard Business School“Cybernetic Teammate”现场实验

研究主体为哈佛商学院研究团队与P&G，样本为791名专业人员，比较个人、团队以及使用生成式AI的任务表现。公开研究显示，AI可让个人达到接近传统团队的部分效果，并改变跨职能协作方式。边界在于实验任务、参与者、工具与时间窗口特定，不能外推为所有创意工作或长期组织绩效。管理启示是重新设计工作单元与评审机制，而非只把AI席位加到原岗位。
[S08]

4.3 Coase与社会技术系统的修正

Coase指出企业边界与交易成本相关。AI降低搜寻、协调和生产部分成本，可能推动更小核心团队与外部网络；但治理、质量、责任和知识保护成本可能反向扩大企业内部化需求。社会技术系统理论则要求技术与岗位、流程、激励共同设计。管理判断：组织边界不会只因模型成本下降而收缩，它取决于“协调成本下降”与“责任成本上升”的净效应。

第五章 客户入口：机器中介与客户关系资产

竹势 AI 营销智库出品

品牌必须同时服务两类客户

人类客户

体验、关系、信任、服务

机器中介

实体、事实、条款、库存、可验证性

自有关系资产是共同底座

2030年前，客户入口将呈多中心结构。企业不能押注某一个搜索框或代理入口，而要让品牌事实、产品数据和服务能力在不同机器中介中可识别、可验证、可交易。

5.1 六类入口并存

入口	用户任务	机器中介角色	品牌风险	必备资产
传统搜索	主动查找	排序与摘要	点击减少、比较前置	权威内容、实体一致性
生成式答案	解释与方案	综合回答	来源被弱化	可引用证据、结构化事实
推荐流	发现与娱乐	预测兴趣	品牌被内容片段化	创意资产、反馈信号
电商平台	比较与购买	排序、定价、履约	客户关系归平台	商品数据、评价、履约
代理式采购	委托筛选与交易	读取偏好、比较、执行	品牌无法直接说服人	API、政策、可验证条款
品牌自有触点	咨询、服务、复购	个性化与关系维护	建设成本高	身份、同意、会员与服务数据

表5-1 客户入口与机器中介资产清单

5.2 信息觅食与Jobs to Be Done

信息觅食理论解释用户会选择“信息气味强、获取成本低”的路径；Jobs to Be Done提醒企业，客户不是为了使用某个渠道，而是为了完成进步。生成式答案和购买代理如果显著降低比较成本，会夺走部分入口价值。品牌要做的不是为每种界面重复生产内容，而是建立可被多入口调用的事实层、证据层和交易层。

5.3 机器中介资产清单

资产层	关键对象	质量标准	责任人
实体层	品牌、产品、型号、服务、地点、人员	命名一致、唯一标识、关系清晰	品牌/数据
事实层	参数、价格、库存、交付、保修、资质	可更新、有时间戳、有来源	产品/运营
证据层	测试、认证、案例、评价、方法	可核验、口径完整、边界明确	法务/品牌
交易层	支付、合同、退换、预约、授权	机器可读、权限可控、可回滚	商务/法务/IT
关系层	身份、偏好、同意、历史互动	客户可控、用途透明、可携带	CRM/隐私

表5-2 客户入口资产五层模型

边界：代理式采购的普及仍受身份、支付、责任、欺诈、偏好表达与售后争议制约。对高情感、高风险、高复杂度购买，人类沟通和线下验证仍可能长期保留。

第六章 品牌信任：内容过剩时代的证据链

竹势 AI 营销智库出品

承诺
来源
审批
履约
纠错
复盘

生成成本趋近于零时，可追责的证据链成为稀缺资产。

生成供给越丰富，品牌越需要把信任建立在可验证来源、一致履约、透明纠错和明确人类责任上。披露AI使用只解决“是否告知”，不能替代真实能力与履约。

欧盟AI法案的透明义务要求特定AI交互和合成内容被识别，部分规则自2026年8月起适用；中国生成式AI服务规则强调内容、数据与服务责任；美国FTC已针对虚假AI能力、假评论和欺骗性商业承诺采取执法行动。[S09][S10][S11]

6.1 品牌信任证据链

环节	核心问题	可验证证据	常见失效
主张	品牌承诺了什么	明确范围、条件和时间	夸大、模糊、无边界
来源	依据来自哪里	原始数据、资质、测试、引用	二手转述、来源失真
生成	内容如何产生	AI披露、版本、审核记录	伪装真人、不可追溯
履约	实际交付是否一致	订单、服务、质量与响应记录	宣传与体验割裂
纠错	错误如何处理	更正时效、公开说明、补救	删除、拖延、推责
责任	谁承担后果	人类负责人、合同与赔付	责任真空、供应商甩锅

表6-1 品牌信任证据链审计表

6.2 Akerlof与信号理论

Akerlof的“柠檬市场”说明当买方无法区分质量，低质量供给会挤压高质量交易。AI降低内容生产成本后，表面专业更容易复制，质量差异更难观察。信号必须具有一定成本或约束力才可信，例如独立认证、可追溯测试、长期履约、公开纠错和可执行保证。单纯增加透明标签可能只增加信息量，不一定增加可信度。

案例：FTC“Operation AI Comply”

2024年9月，FTC宣布针对多家以AI为卖点的欺骗性方案采取行动，涉及假评论、自动化法律服务和“AI电商赚钱”等主张。关键不是AI本身违法，而是未经证实的能力和收益承诺触犯既有消费者保护规则。边界：案件适用美国法律且部分为和解，不代表所有地区相同。启示：AI品牌主张需要证据包、适用条件和禁用表达。[S11]

6.3 危机处置的2030升级

危机管理将从“删稿—声明—媒体沟通”扩展为“事实冻结—机器接口更新—多入口纠错—客户补救—模型与知识库修复”。若只在官网发布更正，而搜索摘要、聊天机器人、代理接口和经销商知识库继续传播旧信息，纠错并未完成。

第七章 媒体与注意力：自动化交易和利润目标

竹势 AI 营销智库出品

平台目标

可归因转化与平台收入

独立增量实验

预算上限

异常接管

企业目标

增量利润、客户价值与品牌资产

媒体交易自动化会扩大，但平台目标函数通常优化可观测转化、平台收入或短期行为，不会自动优化企业长期利润、品牌风险和渠道冲突。

Google在2025年继续为Performance Max增加控制与报告能力，反映市场需要在自动化效率和广告主可控性之间重新平衡。平台官方材料能证明功能方向，不能单独证明所有广告主的增量利润。[S12]

7.1 Simon注意力与Goodhart定律

Herbert Simon指出信息丰富会造成注意力稀缺。自动化系统通过预测和竞价配置注意力，但当点击、转化或ROAS成为硬目标时，系统可能优化可计量行为而损害不可观测价值。Goodhart定律提醒：指标一旦成为目标，就可能失去良好度量功能。企业必须把利润、增量、客户质量、品牌安全和长期价值写入边界。

7.2 平台目标与企业利润对齐检查表

检查项	必须回答的问题	可接受证据	红线
优化目标	系统实际优化什么事件？	事件定义、价值权重、延迟窗口	只说“AI自动优化”
增量性	转化是否本来就会发生？	随机实验、地域/人群留出	只看平台归因
利润	收入是否覆盖毛利、退货和服务成本？	贡献利润、客户质量	只看ROAS
边界	预算、品类、地域、频次如何限制？	硬上限、负面列表、异常规则	无限自动扩量
透明	哪些流量、创意和受众不可见？	可导出报告、审计权限	关键明细完全黑箱
异常	错误如何发现和停止？	阈值、熔断、回滚	只能事后申诉

表7-1 媒体平台目标与企业利润对齐检查表

7.3 独立测量的最低配置

- 建立平台外的转化事实表和成本口径；
- 对大预算或关键渠道设置留出、地域或时间实验；
- 把退款、毛利、复购、销售接受率纳入价值回传；
- 保留预算硬上限、负面关键词 / 品类、品牌安全和异常熔断；
- 对平台模型变更建立版本记录和前后对照。

边界：中小企业可能缺乏足够样本做复杂增量实验。此时应优先缩短反馈链、统一口径、设置预算边界，并避免把不稳定小样本包装为精确因果结论。

第八章 代理商价值：执行商品化后的新合约

竹势 AI 营销智库出品



生成、适配、排期和基础优化的边际成本下降，会压缩部分工时计价；代理商价值将更多来自问题定义、原创判断、跨域整合、专有资产、变革落地和结果责任。

大型代理集团持续投资统一数据与AI平台，并通过并购和能力整合增强规模效应。Omnicom对Omni平台的定位强调连接创意、媒体、数据与AI；2025年完成对IPG的收购也显示行业以规模、数据和技术整合应对结构变化。[S13][S14]

8.1 价值迁移六层

价值层	传统计价	AI压力	未来溢价来源
生产执行	工时、件数	高度承压	专有制作系统、极致效率
平台操作	账户管理费	承压	跨平台治理与透明
专业判断	资深人天	分化	原创洞察、品味与高风险判断
跨域整合	项目费	上升	连接品牌、媒体、销售、数据与法务
组织变革	咨询费	上升	流程重构、培训、采用与治理
结果责任	佣金/绩效费	上升但风险更高	可控因果范围、共享收益与损失

表8-1 代理商价值与计价模式迁移图

8.2 新合同应写清什么

合同模块	关键条款
AI使用与知识产权	训练、生成、第三方模型、素材来源与权利保证
数据与权限	可访问数据、用途、留存、删除和跨境
责任分配	错误发布、预算损失、虚假承诺、侵权与安全事件
透明与审计	工具链、人工审核、日志、分包商和模型变更
计价机制	基础能力费、使用量、里程碑、绩效口径与上限
退出与可迁移	数据导出、模型/工作流迁移、账户归属、知识资产交付

表8-2 AI时代代理商合同重构要点

8.3 反例与边界

小客户和标准化任务仍可能需要低价、快速、打包式代理服务，工时计价不会立即消失。结果付费也不是普遍答案：若代理商无法控制产品、价格、销售和履约，要求其对最终收入承担完整责任会造成逆向选择。更可行的是按“可控结果”分层，例如合格素材通过率、实验速度、增量线索或特定渠道贡献。

第九章 智能商业决策：速度、质量与责任接口

竹势 AI 营销智库出品

目标
证据
方案
授权
执行
例外
复盘

速度只有进入质量与责任闭环，才会成为经营优势。

AI会显著降低信息准备与方案生成成本，但决策质量的瓶颈会转向数据、目标函数、例外处理和反馈设计。企业需要管理“决策系统”，而非追求更多自动建议。

9.1 决策闭环七步

步骤	AI可承担	人类必须保留	质量门
感知	采集、监控、异常发现	定义重要信号	数据完整与时效
解释	归因候选、模式比较	因果质疑、情境判断	证据层级
目标	拆解目标、模拟权衡	价值取舍、红线	目标一致性
方案	生成与压力测试	选择和责任承担	备选充分
执行	工具调用、排期、调整	授权与例外批准	权限和可逆性
反馈	汇总结果、更新模型	判断外部变化	反馈延迟
学习	沉淀规则、知识和案例	废止错误规则	版本与反证

表9-1 智能商业决策闭环

9.2 Deming系统与质量思想

Deming强调大多数质量问题来自系统而非个体。AI把执行速度提升后，坏数据、错目标和弱流程会更快放大。管理者应把模型错误视为系统事件：追查输入、知识源、权限、反馈和组织激励，而不只要求“换更强模型”。

案例：阿里巴巴客服生成式AI现场实验

2026年公开论文报告了阿里巴巴电商售后客服的大规模现场实验：随机分配AI助手，提供问题诊断和解决建议，员工可采用、修改或忽略。结果显示服务速度和部分主观质量改善，但客观重试率无显著改善；低绩效员工收益更大，顶尖员工在多任务环境中出现部分质量下降。边界：论文为特定业务、组织和工具的研究，且发表与复核状态需持续关注。启示：平均提升会掩盖人群异质性，部署策略应按能力与任务分层。[S15]

9.3 决策系统的四项审计

- 目标审计：代理指标是否代表利润、客户价值和风险；
- 数据审计：来源、延迟、缺失、偏差和可追溯性；
- 例外审计：哪些边缘案例会突破规则，如何升级；
- 责任审计：建议、批准、执行和后果分别归谁。

第十章 四种2030营销经营情景

竹势 AI 营销智库出品



四情景由两条主轴构成：AI自主执行权高 / 低，与信任和治理基础强 / 弱。情景用于检验战略，不代表概率总和，也不构成精确预测。

情景	坐标	形成机制	组织图景
A 可信自主商业	高自主 × 强治理	AI在明确责任、可审计和可回滚条件下跨系统执行；客户允许代理参与比较与交易。	人类聚焦目标、关系、原创判断和例外；AI员工承担大部分常规闭环。
B 平台化黑箱增长	高自主 × 弱治理	平台和供应商掌握执行权，企业缺乏数据、测量和责任接口。	效率高但依赖强；异常、虚假优化和责任争议频发。
C 增强型专业组织	低自主 × 强治理	监管、风险偏好或能力边界限制自主执行，但AI广泛增强分析与生产。	岗位重组显著，人保留批准与外部承诺。
D 合成噪声与信任衰退	低自主 × 弱治理	内容供给泛滥、客户不信任、组织无法建立治理，AI停留在碎片工具。	成本下降但经营结果弱，品牌和代理商陷入低价竞争。

表10-1 四情景全景矩阵

A 可信自主商业

维度	情景描述
必要条件	代理能力稳定；身份、支付、权限、审计和赔付机制成熟；企业数据和流程可调用。
组织与岗位	小型高判断核心团队管理多个AI执行单元；出现AI运营控制、代理审计、知识产品与责任设计岗位。
客户入口	生成式答案、购买代理和品牌自有代理可互操作；客户可查看证据与授权范围。
品牌信任	机器可读证据、履约记录和纠错机制成为品牌资产。
媒体交易	预算在护栏内自动配置，独立增量测量常态化。
品牌—代理商	代理商提供经营系统、专有资产和结果责任；计价转向能力费+可控绩效。
风险	目标错配被快速放大；高集中度基础设施形成系统性依赖。
企业动作	扩大L3/L4授权；投资可迁移架构、测量主权和责任保险。

B 平台化黑箱增长

维度	情景描述
必要条件	平台自动化速度快于监管和企业治理；数据与接口高度集中。
组织与岗位	企业内部执行岗位减少，但平台协调、争议处理和风险岗位增长。
客户入口	平台代理控制发现、比较、交易和售后；品牌直接关系弱化。
品牌信任	平台评分替代部分品牌信号，企业难以解释和纠错。
媒体交易	跨渠道自动交易扩大，平台归因成为事实标准。
品牌—代理商	代理商依附平台能力，差异化和议价受限。
风险	锁定、不可解释预算损失、虚假增量、集中性事故。
企业动作	限制不可逆接入；建立第二供应商、数据导出、预算熔断和合同赔付。

C 增强型专业组织

维度	情景描述
必要条件	AI能力进步但高风险执行监管严格；客户偏好人类责任；企业治理成熟。
组织与岗位	人类审批与专业责任保留，AI承担研究、生成、模拟和草稿。
客户入口	多入口并存，机器提供建议，人类完成高风险成交。
品牌信任	专家背书、权威证据和人类服务成为差异化。
媒体交易	自动化受预算、行业与合规限制，实验与人工审核并重。
品牌—代理商	代理商作为专业判断与合规整合者，工时减少但高端溢价上升。
风险	过度谨慎导致效率落后；审批拥堵。
企业动作	优化L2/L3流程；投资证据、专业人才和快速审批。

D 合成噪声与信任衰退

维度	情景描述
必要条件	低质生成泛滥，欺诈和事故频发；组织缺乏数据与治理；客户信任下降。
组织与岗位	工具碎片化，员工反复核验和返工，生产率红利被监督成本抵消。
客户入口	客户回到熟人、社群、线下与强平台担保。
品牌信任	声量失效，履约、真实体验和第三方验证成为稀缺信号。
媒体交易	作弊、无效流量与归因争议增加，预算趋向封闭生态。
品牌—代理商	低价内容代理大量淘汰，少数信誉与行业专家存活。
风险	信任成本、监管成本和客户获取成本同时上升。
企业动作	停止规模化合成内容；收缩渠道，优先真实客户证据与人工服务。

第十一章 情景切换雷达与领先指标

竹势 AI 营销智库出品

技术 长任务成功率与接管率
监管 责任、标识与审计要求
客户 代理入口采用与信任事件
平台 集中度、费率与接口控制
组织 岗位、审批与监督成本
经营 增量利润与风险损失

情景切换不由单一模型发布触发。董事会应同时观察技术、监管、客户、资本、组织和信任六类信号，并记录方向、持续性与反证。

雷达维度	领先指标	证据来源	解释原则
技术能力	真实工具调用成功率、长任务完成率、异常恢复与成本	官方基准、企业生产日志	能力是否跨越可用门槛，而非演示
治理成熟	审计、身份、权限、回滚、保险和责任条款采用	标准、监管、合同	是否支持扩大授权半径
客户行为	生成式答案使用、代理比较 / 交易、人工服务偏好	平台披露、客户研究、一方数据	入口迁移是否持续
媒体结构	自动化预算占比、平台集中度、独立实验使用	平台与上市公司披露	平台权力是否上升
组织采用	任务重构、管理跨度、AI负责人、培训与事故	企业披露、岗位数据	是否从工具试用进入系统运营
信任事件	深伪、错误承诺、数据泄露、假评论执法与品牌危机	监管与法院材料	弱治理情景是否加速

表11-1 情景切换雷达

11.1 领先指标台账模板

指标	当前状态	方向	证据日期/口径	触发动作	反证/退出
AI执行可靠性	待填	↑ /→/ ↓	生产任务样本	扩大或收缩授权	连续异常/成本上升
客户入口迁移	待填	↑ /→/ ↓	按品类的一方数据	调整内容与接口投资	转化质量下降
品牌信任事件	待填	↑ /→/ ↓	投诉、监管、舆情	强化证据与人工责任	事件消退且机制修复
平台集中度	待填	↑ /→/ ↓	流量、预算、交易	第二供应商与可迁移	开放接口改善
监管义务	待填	↑ /→/ ↓	法规实施节点	合规产品化	适用范围变化

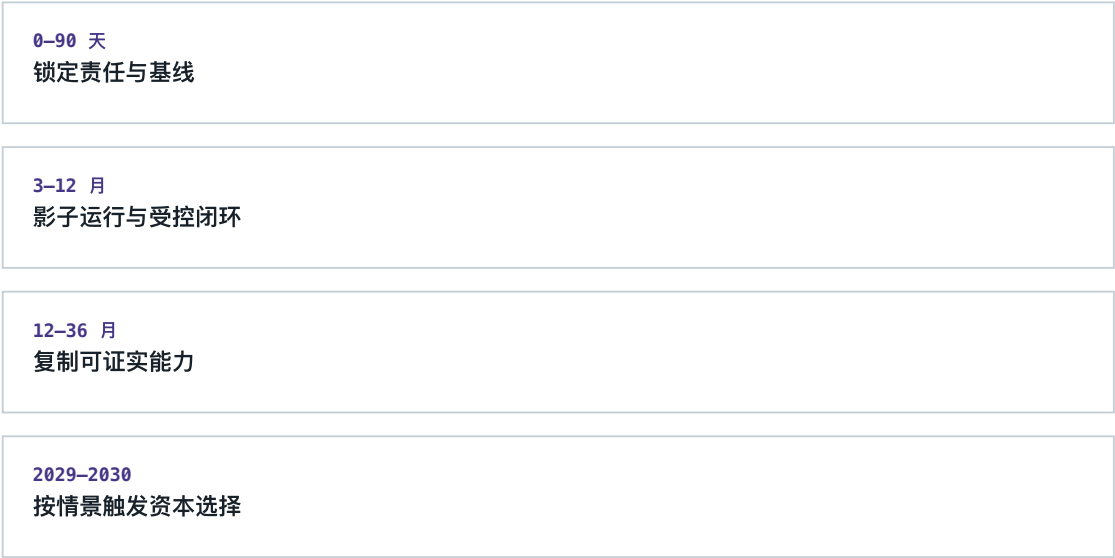
表11-2 情景切换雷达与领先指标台账

11.2 触发式决策原则

只有当信号跨越“能力、治理、经济性”三道门，企业才应扩大不可逆投资。能力门回答能否稳定完成；治理门回答能否控制和负责；经济性门回答全成本是否优于替代方案。任何一门未通过，都应保留期权而非全面扩张。

第十二章 2026—2030行动路线

竹势 AI 营销智库出品



企业应先建设无悔能力，再按场景信号扩大授权和资本投入。路线图的核心交付物不是“上线一个AI平台”，而是形成可运营的责任、数据、流程、测量和学习系统。

阶段	核心目标	责任与动作	关键交付物	退出条件
90天	建立基线与一项受控试点	CEO定风险偏好；CMO选闭环；CIO建日志与权限；法务定红线；HR做任务盘点	授权矩阵、数据清单、试点评估、事故预案	没有明确业务负责人或基线则不启动
12个月	形成3—5个可重复场景	统一知识、工作流、测量和供应商条款；重构岗位与审批	场景组合、AI员工台账、证据链、增量测量	全成本高于替代方案且无学习价值则退出
36个月	跨部门扩张与客户入口建设	连接营销、销售、服务和数据；建立机器中介资产层	统一实体/事实/证据/交易接口；多供应商架构	锁定过强、责任不清、事故率上升则暂停
2030	按情景切换组织与资本配置	根据雷达选择可信自主、增强专业或防御收缩路径	动态授权、情景组合、董事会年度重估	不以既有投资沉没成本阻止退出

表12-1 90天、12个月、36个月、2030路线图

12.1 角色责任矩阵

角色	90天	12个月	36个月及以后
CEO	风险偏好、场景赞助、决策门	资本组合与责任制度	按情景调整组织边界
CMO	闭环选择、指标基线	入口、品牌证据与代理商重构	经营级AI营销系统
CIO/数据	权限、日志、数据源	平台架构、身份、互操作	可迁移、多供应商和实时决策
法务风控	红线、披露、事件响应	合同、审计、监管映射	责任保险与跨域治理
HR	任务盘点、培训与沟通	岗位、绩效、管理跨度	技能组合与劳动力转型
代理商管理	AI披露与成本透明	新计价、资产归属、绩效口径	组合生态与退出机制

表12-2 跨职能责任矩阵

12.2 行动组合

行动类型	现在做什么	何时扩大	何时停止
无悔行动	数据与知识质量、权限日志、品牌证据、独立测量、任务盘点	持续投入	仅在无业务价值时缩减范围
期权行动	购买代理接口、L3自主执行、AI原生组织单元	能力+治理+经济性三门通过	验证失败或锁定过强
条件行动	大规模自动交易、结果付费、跨系统L4授权	领先指标持续触发	事故、监管或客户信任恶化
禁止行动	无日志自动发布、无上限预算、虚假AI主张、不可导出锁定	不适用	立即停止

表12-3 无悔行动—期权行动—条件行动—禁止行动组合

董事会与经营层工具箱

竹势 AI 营销智库出品

以下工具可直接用于季度经营会、年度战略会、AI治理委员会和代理商评审。

工具1：董事会季度提问清单

1. 本季度AI获得了哪些新执行权限？每项权限的责任人、额度和回滚机制是什么？
2. 哪些经营结果有因果或增量证据，哪些只是平台归因或供应商自报？
3. 客户入口份额是否发生变化？我们拥有的一方关系资产是在增加还是减少？
4. 品牌主张是否都有来源、适用条件、更新时间和纠错责任人？
5. 发生了哪些近失事件、错误承诺、异常预算或数据泄露？系统如何修复？
6. AI带来的节省是否扣除了监督、返工、模型、集成与风险成本？
7. 岗位任务、决策权和管理跨度发生了什么变化？谁承担了新增隐形劳动？
8. 代理商交付中，哪些是商品化执行，哪些是专有资产和结果责任？
9. 我们对单一平台、模型或代理商的锁定程度是否上升？退出测试是否通过？
10. 四情景中，哪一个正在获得更多证据？哪些动作需要切换？

工具2： AI员工台账字段

字段组	必填字段
身份	名称、业务负责人、技术负责人、供应商、版本
目的	任务、目标、禁止目标、适用客户/地区
权限	数据、工具、额度、时间窗、外部承诺
控制	审批、日志、异常阈值、熔断、回滚、人工接管
质量	评估样本、通过标准、已知失效、监控频率
责任	RACI、合同、赔付、事件分级、监管适用
经济性	模型、集成、监督、返工、事故与机会成本

工具3： 品牌信任证据链评分

评分	标准	董事会判断
0 缺失	无来源、无责任、无纠错	禁止对外自动化
1 初始	有主张和人工审核，证据分散	仅L1/L2
2 可控	来源、版本、审核和纠错可追溯	可小范围L3
3 可信	跨入口一致、履约验证、客户可申诉	扩大授权
4 韧性	主动发现错误、快速纠错、独立审计	可支持高自主场景

工具4：代理商价值评审问题

- 代理商拥有何种不可替代的专有数据、方法、人才或分发资产？
- AI降低的成本是否传递给客户，还是只改变代理商毛利？
- 哪些结果由代理商可控，绩效费如何排除产品、价格与销售因素？
- 生成内容、模型配置、工作流、提示资产和学习数据归谁？
- 发生错误发布、侵权、预算异常或数据事件时谁赔付？
- 终止合作后，账户、数据、工作流 and 知识能否迁移？

附录A 案例卷宗

竹势 AI 营销智库出品

案例	主体/时间/情境	结果证据	机制与启示	边界
Air Canada聊天机器人	2022—2024，加拿大，航空/高责任服务	错误政策答复造成赔付	企业对机器人公开信息承担责任	特定司法辖区裁决
Klarna客服AI	2024—2025，瑞典/全球，金融科技	高自动化率与效率后补充人工	按复杂度分层与混合服务	早期结果多为公司自报
P&G/HBS实验	公开研究，791名专业人员，消费品/知识工作	AI改变个人与团队绩效边界	重构工作单元而非只减员	实验任务不可直接外推
Alibaba客服实验	2026论文，中国，电商售后	平均提速、异质性明显	按员工与任务分层部署	论文场景与状态需持续复核
FTC Operation AI Comply	2024，美国，多行业	打击虚假AI能力、假评论与收益承诺	AI主张必须有证据与边界	美国执法语境
Google Performance Max	2024—2025，全球，广告平台	自动化同时增加控制与透明功能	平台效率需独立测量和护栏	官方功能不等于普遍增量
Omnicom Omni与整合	2024—2025，全球，代理集团	数据、媒体、创意和AI平台化	代理商价值向整合与资产迁移	并购协同结果仍需观察
欧盟AI法案透明义务	2024—2026，欧盟，监管	合成内容与AI交互透明要求分阶段适用	披露进入产品设计	适用范围与细则持续演进

表A-1 八个关键案例总览

附录B 专家与经典思想转译

竹势 AI 营销智库出品

专家/来源	原意或观点	本题连接	管理判断
Herbert Simon	有限理性与注意力稀缺	AI扩展计算和信息供给，但目标设定与注意力仍稀缺	不能用更多信息替代更好问题
George Akerlof	信息不对称与柠檬市场	生成供给降低可观察质量差异	投资高成本可信信号
Ronald Coase	企业边界与交易成本	AI同时降低协调成本、提高治理责任成本	按净交易成本重画组织边界
W. Edwards Deming	系统与质量管理	AI会放大系统缺陷	先修目标、数据和流程
Philip Tetlock	预测校准	拆解预测、持续更新、记录误差	建立情景预测日志
Peter Schwartz/Shell传统	情景规划	多种一致未来检验韧性	情景必须改变资源配置
NIST AI RMF团队	治理—测量—管理风险	把AI风险管理嵌入生命周期	治理是扩权前提
ILO/OECD研究团队	任务暴露与劳动转型	岗位结果取决于任务、技能和制度	避免把暴露等同替代
Lina Khan/FTC执法语境	既有消费者保护法适用于AI	没有“AI免责”	主张、证据和责任一致
平台产品负责人视角	自动化需更多控制与透明	广告主需要边界与可解释操作	平台最优不等于利润最优

表B-1 具名视角与经典思想的经营转译

说明：上述观点用于构建分析框架，不替代事实证据；引用时应回到原著、论文、官方材料或正式发布语境。

附录C 来源索引

竹势 AI 营销智库出品

以下索引保留机构 / 作者、标题、日期和口径说明，不设置第三方可点击超链接。

编号	机构/作者	标题	日期	样本 / 口径与限制
S01	OpenAI	A Practical Guide to Building AI Agents	2025	面向产品与工程团队的代理设计、工具、护栏与部署指南；供应商官方材料。
S02	NIST	AI Risk Management Framework; Generative AI Profile NIST-AI-600-1	2023—2024	自愿风险管理框架与生成式AI风险画像。
S03	British Columbia Civil Resolution Tribunal / Moffatt v. Air Canada	2024 BCCRT 149	2024-02-14	聊天机器人错误信息责任的具体裁决。
S04	Klarna	AI assistant handles two-thirds of customer service chats in its first month	2024-02-27	公司自报：230万次对话、覆盖约三分之二客服聊天等。
S05	公开采访与后续报道	Klarna reinvests in human customer service	2025	公司转向混合服务的后续信息；不同报道口径需区分。
S06	ILO	Generative AI and Jobs: A Refined Global Index of Occupational Exposure	2025-05-20	任务级数据、专家输入与模型预测；评估职业暴露而非直接预测失业。
S07	OECD	Artificial intelligence and the changing demand for skills in the labour market	2024	AI对非AI专业岗位技能需求的影响。
S08	Harvard Business School researchers / Procter & Gamble	The Cybernetic Teammate	2025	791名专业人员现场实验；任务、样本与外推边界需保留。
S09	European Union	Regulation (EU) 2024/1689 — Artificial Intelligence Act	2024-07-12	风险分级、透明与通用AI等规则；分阶段实施。
S10	中国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； 2023-08-15施行	面向境内公众提供生成式AI服务的管理规则。

S11	U.S. Federal Trade Commission	Operation AI Comply enforcement announcements	2024-09-25	针对虚假AI主张、假评论及欺骗性商业方案的执法行动。
S12	Google	Kick off 2025 with new Performance Max features	2025-01-23	平台官方功能材料，涉及控制、报告与自动化。
S13	Omnicom	Omni platform corporate materials	2024-2026	集团对连接创意、媒体、数据与AI能力的官方定位。
S14	Omnicom	Completion of acquisition of Interpublic	2025-11	集团官方并购完成公告；协同结果仍待后续披露。
S15	Ni, Wang, Feng, Lu et al.	Generative AI in Action: Field Experimental Evidence from Alibaba's Customer Service Operations	2026	随机现场实验论文；速度、主观质量、客观重试与员工异质性。
S16	OECD	Employment Outlook 2023: Artificial Intelligence and the Labour Market	2023	七国制造与金融调查及劳动质量讨论。
S17	OpenAI	New tools for building agents	2025-03-11	代理作为代表用户独立完成任务的官方定义与工具发布。
S18	European Commission	AI Act policy and implementation timeline	持续更新	透明规则与分阶段实施说明。
S19	FTC	Rulemaking and enforcement on impersonation and AI claims	2024	消费者保护、冒充与AI欺骗风险。
S20	Philip E. Tetlock & Dan Gardner	Superforecasting: The Art and Science of Prediction	2015	预测分解、校准、更新与团队实践。
S21	Pierre Wack / Shell scenario tradition	Scenarios: Uncharted Waters Ahead	1985	情景规划用于战略学习而非单点预测。

S22	Herbert A. Simon	Administrative Behavior; Designing Organizations for an Information-Rich World	1947/1971	有限理性与注意力稀缺。
S23	George A. Akerlof	The Market for “Lemons”	1970	质量不确定性与市场机制。
S24	Ronald H. Coase	The Nature of the Firm	1937	交易成本与企业边界。
S25	W. Edwards Deming	Out of the Crisis	1982/1986	系统、质量与管理责任。
S26	William H. Starbuck / High Reliability Organization literature	高可靠组织研究	1980s—	异常识别、冗余、升级和安全文化。
S27	John D. Sterman / control and systems traditions	Business Dynamics	2000	反馈、延迟、非线性与系统决策。
S28	Clayton Christensen et al.	Competing Against Luck	2016	Jobs to Be Done与客户 进步任务。
S29	Peter Pirolli & Stuart Card	Information Foraging Theory	1990s	信息气味、搜索成本与路径选择。
S30	Daniel Kahneman, Olivier Sibony, Cass Sunstein	Noise	2021	决策噪声、判断一致性与 流程设计。

结语：把未来写进决策门，而不是口号

竹势 AI 营销智库出品

到2030年，企业很可能拥有更多会生成、会调用工具、会协调流程的软件执行单元。真正稀缺的能力将是：知道该把什么权力交给它们，知道如何观察、限制和纠错，知道如何保有客户关系和品牌证据，并能在情景发生变化时及时切换。

董事会最重要的工作不是挑选“最像未来”的单一情景，而是建设跨情景都有效的能力：高质量数据与知识、可追责权限、品牌信任证据、独立增量测量、可迁移架构、人机任务设计和动态退出机制。完成这些基础后，企业才有资格在自主执行、购买代理和智能交易上扩大下注。

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。

本报告以关键不确定性、四种经营情景、领先指标与触发式决策门，帮助管理团队跨情景配置人、AI 员工、客户关系、品牌信任与商业系统。

shichangbu.ai

专题中心 · 返回顶部 ↑