



027 / DECISION RIGHTS CIRCUIT

DECISION RIGHTS CONTROL CIRCUIT /
RACI-X · EVIDENCE · STOP

AI市场部RACI与 决策权设计指南

把机器执行、人类判断、经营批准与最终问责连接成可审计、
可撤销、可复盘的控制回路。

从 RACI 升级为 RACI-X

机器执行、专业复核与异常升级可观察

穿透八个责任事件

以行动、机制、结果与边界还原责任链

配置五类权力令牌

建议、执行、批准、停止与风险接受

完成 90 天权责改造

盘点、重配、运行与再认证

BOARD BRIEF 01

执行摘要

竹势 AI 营销智库出品

BOARD CONTROL THESIS / 董事会控制判断

治理 AI 不是给角色多加几个字母，而是让每项决策都能回答：谁能建议、谁能执行、谁能批准、谁能停止、谁接受残余风险。

核心判断：RACI 表本身不能治理 AI。传统 RACI 适合说明稳定任务中谁做、谁负责、谁被咨询、谁被告知；当 workflow 开始跨系统调用、自动触达客户、调整预算、调用个人数据并产生概率性输出时，单张角色矩阵无法表达授权额度、风险接受、证据包、例外升级、暂停、回滚和版本责任。企业需要的不是更复杂的字母表，而是一套可观察、可撤销、可复盘的决策权系统。

本指南提出 RACI-X 责任回路：保留 Responsible、Accountable、Consulted、Informed，并增加 eXecute（机器受限执行）、eXamine（专业复核）和 eScalate（异常升级）三类动作。任何具体决策对象仍只能有一个最终 Accountable 的人类岗位。复杂组织可以拆分批准权、否决权、风险接受权和系统授权权，但不能用“多人共同负责”掩盖事故后无人能够调资源、接受风险或纠偏。

表1 十大核心结论

序号	结论	董事会含义
1	RACI 仍是角色沟通底图，但不是 AI 控制系统	它不表达额度、期限、模型版本、数据范围、异常阈值和停止权。
2	机器可以拥有受限执行权，不能成为法律或经营责任主体	AI 不具备资源配置权、风险接受能力和组织处分能力。
3	最终 Accountable 必须绑定“可纠偏能力”	没有预算、人员、系统或供应商控制权的人，不应被写成最终 A。
4	决策权必须按对象拆分	内容准确性、品牌风险、预算、个人信息、系统权限是不同决策对象，不能笼统写“批准活动”。
5	低风险高频任务应默认减少审批	用白名单、额度、抽样复核和自动监控替代逐件盖章。
6	高风险任务采用双钥匙和停止权	业务批准与风险批准分离，任何一方可暂停，恢复需证据。
7	human in the loop 不是安全承诺	复核者可能疲劳、缺信息、缺能力或机械盖章，必须测量覆盖率和有效性。
8	日志不是问责本身	日志只有与身份、版本、决策理由、证据和纠偏动作连接时才有治理价值。
9	供应商合同不能替代运行控制	企业仍需定义数据边界、变更通知、事故协同、可迁移性和证据可得性。
10	权责设计要用经营指标验收	同时观察决策时延、返工、例外、越权、停机、事故恢复和经营结果。

第1章 传统 RACI 为什么在 AI 营销中失效

竹势 AI 营销智库出品

01 任务边界漂移	02 机器跨系统执行	03 输出概率化	04 责任链跨组织
--------------	---------------	-------------	--------------

判断。传统 RACI 的价值在于消除角色模糊，但它隐含四个稳定性假设：任务边界相对固定、执行者是可识别的人、输入输出较可预测、责任在一个流程节点内闭合。AI 营销 workflow 逐一破坏这些假设：一个内容发布动作可能同时调用模型、知识库、客户标签、社媒账号和审批接口；同一智能体可能被多个团队复用；模型和数据会变化；错误可能在数小时后通过投放、转发或 CRM 触达放大。

机制。治理对象应从“任务”扩展为“决策对象 + 风险级别 + 权力令牌 + 证据要求 + 授权期限 + 升级条件 + 停止权”。RACI 仍保留为组织沟通底图，但必须嵌入控制回路。NIST AI RMF 将治理、映射、测量、管理作为相互连接的四项功能；ISO/IEC 42001要求组织建立、实施、维护并持续改进 AI 管理体系。这些框架共同指向持续运行，而非一次性画表。

表2 传统 RACI 假设与 AI 工作流现实

传统假设	AI 营销现实	失效表现	修正控制
任务稳定	模型、数据、渠道、规则持续变化	角色表很快过期	授权版本化、定期再认证
执行者可追责	机器执行、多人配置、供应商参与	责任被推给模型或平台	人类最终 A + 事件链责任
输出可预测	概率性生成、上下文漂移	审批标准不一致	测试集、置信阈值、证据包
流程边界清晰	跨 CRM、广告、内容、数据系统	灰区与重复审批	按决策对象拆分权力
异常可人工发现	高速批量执行	监督疲劳、发现滞后	监控、熔断、抽样与红队

表3 RACI 中仍然有效与必须扩展的部分

元素	仍然有效	必须扩展
R	明确人类任务执行责任	与机器 eXecute 分开，避免把自动执行伪装成人工执行
A	要求最终归责唯一	必须验证其有资源、风险接受与纠偏能力
C	识别专业咨询关系	规定何时必须咨询、意见冲突如何裁决
I	保证受影响者知情	定义通知时点、内容和事故通报
新增动作	—	eXecute、eXamine、eScalate

第2章 RACI-X：把责任角色转成可观察动作

竹势 AI 营销智库出品

R 人类执行	A 最终问责	C 事前咨询	I 约定知情	X 受限执行	X 专业复核	X 异常升级
------------------	------------------	------------------	------------------	------------------	------------------	------------------

RACI-X 的设计原则是“角色不等于动作”。同一岗位可以在不同对象上承担不同权力：CMO 对品牌表达可能是 A，对个人信息合法性只能是 C；法务对违法风险可以拥有停止权，却不应替代业务负责人承担增长结果。机器 eXecute 只描述被授权动作，不承担责任。专业复核 eXamine 必须说明复核标准、证据和 SLA。异常升级 eScalate 必须由阈值触发，不能依赖执行者临场勇气。

表4 RACI-X 七角色定义

符号	名称	可观察行为	不得替代
R	Responsible	完成人类任务、整理证据、处理例外	不得把“监督机器”写成空泛负责
A	Accountable	对具体决策对象作最终取舍并承担结果与纠偏	不得多人并列；不得由 AI 承担
C	Consulted	在决定前提供专业意见并留下意见记录	不得以咨询名义拥有隐形否决
I	Informed	在约定时点收到结果、风险或事故信息	不得被动承担事后责任
X-eXecute	机器受限执行	在白名单、额度、期限和数据范围内调用工具	不得批准自身输出或扩大自身权限
X-eXamine	专业复核	按标准检验证据、事实、品牌、法律或数据风险	不得只点击“通过”而无能力与时间
X-eScalate	异常升级	阈值触发暂停、转人工、扩大审查或通知	不得将所有异常都升级到最高层

表5 “唯一最终 A”的适用边界

情形	正确拆分	错误做法
跨部门战役	对收入目标由业务负责人 A；对品牌表达由 CMO 决策；对隐私风险由法务/隐私负责人拥有否决与风险接受流程	CMO、销售 VP、法务三人同时 A
集团与区域	集团定义不可突破红线；区域总经理对本地执行结果 A	总部与区域互相写 A，事故后互相等待
高监管内容	业务负责人对是否发布 A；医学/合规负责人拥有强制否决；董事会批准风险偏好	合规成为所有经营结果 A
供应商交付	内部项目负责人对采用与上线 A；供应商对合同交付和缺陷修复负责	把“平台负责”写进矩阵后不设内部 A

第3章 七类决策对象与五类权力令牌

竹势 AI 营销智库出品

01 建议权 提出方案	02 执行权 限定动作	03 批准权 允许生效	04 停止权 阻断暂停	05 风险接受权 接受残余风险
-------------------	-------------------	-------------------	-------------------	-----------------------

权责冲突通常不是角色问题，而是对象没有被拆开。一个“自动发布”动作至少包含七类对象：目标、内容事实、品牌表达、客户触达、数据使用、系统权限和预算。企业应分别配置建议权、执行权、批准权、否决/停止权、风险接受权。

表6 七类决策对象

对象	典型问题	默认所有者	关键证据
目标与优先级	为什么做、服务哪个经营结果	业务单元负责人/CMO	目标、客户、单位经济性
专业判断	事实、洞察、创意是否成立	领域专家/品牌负责人	来源、测试、品牌规范
经营批准	是否投入资源并对外行动	预算或业务 P&L 负责人	预算、预测、替代方案
风险接受	在残余风险下是否继续	被授权风险所有者	风险评估、补偿控制
系统授权	AI 可调用哪些账户、工具和数据	CIO/CDO 或系统所有者	权限清单、最小权限、期限
异常升级	何时暂停、转人工、通知谁	流程 A + 风险职能	阈值、告警、SLA
最终问责	谁对结果和纠偏负责	具有资源与纠偏能力的人类岗位	责任契约、复盘记录

表7 五类权力令牌

令牌	对象	额度/范围	期限	证据	撤销条件
建议权	方案、内容、预算建议	不直接产生外部影响	按任务或版本	依据、置信度、替代方案	持续低质量、偏离规则
执行权	发布、触达、查询、调价等动作	白名单账户、金额、频次、客户群	短期授权并再认证	执行日志、前置检查	越权、异常率、规则变更
批准权	允许输出或行动生效	限定决策对象	任期或场景	证据包、审批理由	岗位变化、能力不足、利益冲突
否决/停止权	阻断或暂停流程	明确红线和触发条件	持续有效	告警、违规证据	只能由更高等级治理程序恢复
风险接受权	接受残余风险继续运行	风险类别与上限	有到期日	风险评估、补偿控制、退出方案	风险超限、条件变化、事故

表8 权力令牌配置检查

检查项	合格标准	红旗
对象	精确到动作与数据对象	“负责 AI”
额度	金额、频次、受众、数据范围可量化	“适当额度”
期限	有起止日期或再认证周期	永久授权
证据	列明上线前和运行中证据	只凭审批截图
撤销	定义自动与人工撤销条件	只有管理员手工关闭

第4章 任务—风险双轴授权图

竹势 AI 营销智库出品



双轴的第一轴是任务属性：频率、可逆性、标准化程度和反馈速度；第二轴是风险暴露：客户影响、品牌不可逆性、金额、数据敏感度、监管强度和可回滚性。频率高不等于可自动化，金额低也不等于低风险。公开发布一条错误医疗内容金额可能为零，但品牌与监管后果很高。

表9 双轴授权五档

档位	适用条件	控制模式	示例
L1 自动执行	高频、可逆、低外部影响、无敏感数据	白名单 + 自动监控 + 事后抽样	内部标签整理、素材命名
L2 抽样复核	标准化、可回滚、客户影响有限	机器执行，按比例抽样与漂移监控	低风险 SEO 元数据、内部摘要
L3 事前批准	对外可见或涉及客户，但可回滚	专业复核后单人批准	常规社媒发布、CRM 邮件草稿
L4 双钥匙	高金额、敏感数据、强监管或难回滚	业务批准 + 风险/专业批准，任一可停止	预算大幅调整、医疗金融触达
L5 禁止自动化	违法、不可接受伤害、缺乏可靠控制	仅可辅助建议，不可自动执行	伪造评价、歧视性定价、绕过同意

表10 风险评分维度与阈值

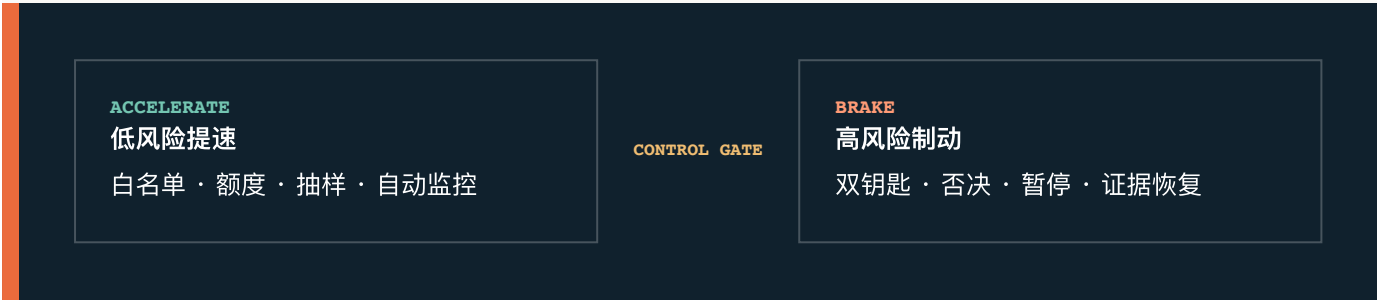
维度	1分	3分	5分
客户影响	内部使用	小范围客户可见	大规模或弱势群体
品牌不可逆性	完全可撤回	会被截图传播	涉及信任、危机或虚假陈述
金额	无直接金额	在部门授权额度内	显著影响 P&L 或客户权益
数据敏感度	公开数据	一般个人信息	敏感个人信息/跨境
监管强度	一般行业	广告/消费者规则显著	金融、医疗、未成年人等
可回滚性	即时自动回滚	需人工修复	不可逆或损害已扩散

表11 授权决策规则

总分/单项	授权建议	必要控制
≤8 且无单项5	L1-L2	监控、抽样、月度复核
9-15	L3	事前专业复核、明确 A
16-22 或任一单项5	L4	双钥匙、红队、暂停与回滚
>22 或法律禁止	L5	禁止自动化，仅限建议

第5章 低风险提速与高风险制动

竹势 AI 营销智库出品



过度审批的典型后果是低风险任务排队、高风险任务被机械盖章。企业应把控制强度放在“不可逆、规模化、敏感数据和强监管”上，把低风险高频任务转为自动执行、抽样复核和异常触发。双钥匙不是所有事项双签，而是两种不同能力的独立判断：经营负责人确认价值与资源，专业/风险负责人确认红线与残余风险。

表12 低风险任务去审批设计

控制	替代的人工动作	运行指标
白名单	逐件确认允许的工具和账户	白名单命中率、越权尝试
额度	逐次确认小额预算或频次	额度使用率、超限次数
抽样复核	100%人工检查	样本缺陷率、漏检率
自动回滚	人工发现后修复	平均恢复时间
漂移告警	定期凭感觉检查	规则/模型/数据漂移事件

表13 高风险双钥匙与停止机制

环节	业务钥匙	风险/专业钥匙	停止触发	恢复条件
客户数据调用	确认目的与客户价值	确认合法性、最小必要与同意	范围扩大、敏感数据出现	重新评估并缩减范围
预算调整	确认收益假设与额度	确认异常、欺诈、财务控制	单日偏差或异常消耗	核对数据与人工复位
高监管内容	确认业务必要性	确认专业准确与法规	证据不足、禁用表述	补证并重新批准
危机回应	确认立场与资源	确认法律、事实与披露风险	事实冲突、舆情升级	指挥官批准新版本

表14 human in the loop 失效模式

失效模式	机制	检测指标	改进
监督疲劳	高频审批导致注意力下降	平均审阅时长、连续通过率	抽样+异常优先
机械盖章	批准者缺信息或责任感	无理由批准占比	强制证据与差异高亮
能力错配	审阅者不懂专业/数据	复核后缺陷率	资格认证与分层复核
自动化偏差	过度信任机器建议	人类推翻率异常低	盲测、反向挑战
责任倒挂	执行者承担无权纠偏责任	事故后升级失败	A 与资源权匹配

第6章 责任链五问与事件归责

竹势 AI 营销智库出品

1	2	3	4	5
谁设目标	谁给权限	谁看证据	谁能停止	谁负责纠偏

事故归责应沿事件链分析，而不是寻找一个“最后点击的人”。Reason 的瑞士奶酪模型说明事故通常来自多层防线孔洞对齐；Perrow 的正常事故理论提醒高耦合、复杂系统中意外组合不可完全消除。因此责任应分为目标责任、数据与规则责任、工具授权责任、输出批准责任、监控纠偏责任。最终 A 对结果与修复负责，但不意味着所有过错都归于同一人。

表15 责任链五问

问题	责任对象	应留证据	常见空洞
谁定义目标	目的、受众、成功标准	目标卡、风险偏好	目标模糊却追责执行者
谁提供数据与规则	数据合法性、质量、品牌/业务规则	数据目录、规则版本	错误数据无人负责
谁授权工具	账户、权限、模型和连接器	授权记录、期限、最小权限	供应商默认权限
谁批准输出/行动	内容、触达、预算、发布	证据包、批准理由	批准者无信息
谁监控并纠偏	运行、异常、投诉、恢复	告警、停机、复盘	事故后才发现

表16 五类事故的归责路径

事故	首要控制缺口	事件链责任	最终 A 的义务
模型错误	验证与适用边界不足	模型/规则所有者、复核者、批准者	暂停、纠正、评估影响
错误发布	发布权限和复核失效	内容 R、品牌 eXamine、发布 A、系统授权者	撤回、通知、流程修复
预算失控	额度、监控、回滚不足	投放负责人、财务控制、系统所有者	止损、核账、恢复
隐私泄露	目的、最小权限、供应商控制不足	数据所有者、系统授权者、处理方	遏制、通知、补救
歧视输出	数据、测试、影响评估不足	产品/营销所有者、模型与数据团队、风险审查者	停止使用、纠正受影响群体

表17 事故复盘的责任与非责任

应追问	不应接受
谁拥有可防止或减轻该失效的控制权	“AI 自己做的”
其是否获得足够信息、时间和资源	“已经有人在环”
控制是否按设计运行，设计本身是否合理	“日志显示流程跑完了”
供应商、内部配置和使用情境分别贡献了什么	“合同写了供应商负责”
如何调整授权、阈值、证据和培训	只处罚最后操作人

第7章 跨职能与外部伙伴共同责任

竹势 AI 营销智库出品



跨 CMO、销售、产品、数据、IT、法务与代理商的工作流，应按决策对象分轨，而不是建立一个庞大的联合审批委员会。业务 P&L 负责人通常对经营目标与资源最终 A；CMO 对品牌和营销专业判断拥有决定权；CIO/CDO 对系统授权和数据平台控制负责；法务/隐私/内控拥有特定红线否决权与风险接受程序；代理商对合同交付、事实说明和其控制范围内的缺陷负责，但内部采用与上线仍必须有企业 A。

表18 跨职能决策权轨道

决策对象	主决策者	强制咨询/复核	否决/停止	最终 A
营销目标与预算	业务单元负责人/CMO	CFO、销售、产品	超风险偏好时风险职能	P&L 负责人
品牌表达	CMO/品牌负责人	业务、法务、专业专家	违法或重大声誉风险	CMO
客户数据使用	数据所有者+业务提出者	隐私、信息安全	隐私/安全负责人	业务数据处理负责人
系统权限	CIO/CDO/系统所有者	业务、内控、安全	安全负责人	系统所有者
代理商交付上线	内部项目负责人	品牌、法务、采购	相关风险职能	内部项目负责人

表19 供应商共同责任证据边界

责任域	企业必须掌握	供应商应提供	不能外包
模型/平台变更	版本清单、适用场景、回归测试	重大变更通知、已知限制	是否继续使用的决定
数据处理	目的、数据分类、访问范围	处理位置、子处理者、安全措施	合法性与最小必要判断
运行事件	内部日志、客户影响、恢复记录	平台事件、根因、补救	客户通知与经营纠偏
质量与偏差	本企业情境测试	通用评测与限制	针对客户/行业的适用性
退出与迁移	导出、替代、停用方案	可迁移格式、删除证明	业务连续性

表20 代理商权责契约最低条款

条款	必须写清
决策对象	代理商可建议、可制作、可发布或仅可提交草稿
账号权限	具体账户、操作范围、时间和额度
证据包	素材来源、版权、事实依据、模型/工具使用说明
审批与 SLA	谁批准、超时如何处理、不得默认批准
事故协同	暂停、通知、保全证据、客户沟通
版本责任	谁维护模板、品牌规则、模型和工作流版本
退出	权限回收、数据删除、资产移交

第8章 八个完整案例：成功、受限与失败

竹势 AI 营销智库出品

CASE DOCKET / 八个责任事件档案

主体与时间 · 场景 · 行动 · 机制 · 结果 · 边界 · 管理含义

案例1：腾讯：把 Responsible AI 原则嵌入集团治理

表21 案例卷宗：腾讯：把 Responsible AI 原则嵌入集团治理

项目	内容
主体与时间	腾讯；2024—2026
业务情境	大型平台企业同时面对内容、社交、广告、云与大模型业务，风险横跨隐私、安全、公平和技术滥用。
具体动作	公开 Responsible AI 原则，并在 ESG 披露中将 AI 风险治理、隐私保护、数据安全、算法公平和防止技术滥用纳入治理框架。
作用机制	以集团原则定义不可突破边界，再由业务与技术团队将原则转成具体控制，形成集团风险偏好与场景执行之间的接口。
结果证据	官方原则与 2025 ESG 报告披露治理方向；未披露单一营销工作流的审批时延或事故率。
边界与限制	原则披露不等于每个工作流已形成唯一 A、额度与停止权；外部无法验证运行有效性。
权责启示	集团级原则必须下沉为权责契约卡，否则会停留在价值声明。

案例2：中国平安：高监管集团的 AI 伦理与风险治理

表22 案例卷宗：中国平安：高监管集团的 AI 伦理与风险治理

项目	内容
主体与时间	中国平安；2022—2026
业务情境	金融、医疗与养老场景同时涉及敏感个人信息、专业判断和高监管要求。
具体动作	发布 AI 伦理治理政策，强调以人为本、人类自主、安全可控、公平正义、开放透明；持续在可持续发展报告中披露数字化与 AI 治理。
作用机制	把 AI 风险放入既有集团治理、审计与风险管理体系，而非设立孤立的“AI 团队负责”。
结果证据	官方政策和年度报告可核验治理架构；未披露具体营销自动化的差错率与授权阈值。
边界与限制	高监管集团经验不能直接复制到小企业；控制成本与组织层级不同。
权责启示	风险接受权应进入既有风险治理轨道，专业判断与经营批准不可混为一体。

案例3：中国网信部门：生成合成内容标识责任链

表23 案例卷宗：中国网信部门：生成合成内容标识责任链

项目	内容
主体与时间	国家互联网信息办公室等部门；2025—2026
业务情境	生成式内容在制作、导出、传播和平台分发中跨多个主体，单靠最终发布者难以完成追溯。
具体动作	《人工智能生成合成内容标识办法》区分显式与隐式标识，明确服务提供者与传播服务提供者的责任；2025 年 9 月 1 日实施，后续执法通报显示未标识、隐式标识不规范等会被处置。
作用机制	把责任沿内容生命周期分配到制作、文件元数据、传播和用户声明环节，形成可追踪链条。
结果证据	监管与执法通报；不是企业经营绩效案例。
边界与限制	标识能改善知情与追溯，但不能证明内容真实、合法或无偏差。
权责启示	营销流程必须把标识作为发布前证据项，不能把“平台会自动处理”当控制。

案例4：Air Canada 聊天机器人：企业不能把错误承诺推给自动系统

表24 案例卷宗：Air Canada 聊天机器人：企业不能把错误承诺推给自动系统

项目	内容
主体与时间	Air Canada / 加拿大不列颠哥伦比亚省民事解决法庭；2022—2024
业务情境	客户根据航空公司网站聊天机器人的错误信息申请丧亲票价退款，企业主张机器人是独立信息源。
具体动作	争议进入裁决，裁决认定企业应对其网站信息负责。
作用机制	客户触达系统代表企业行动；部署者拥有内容、系统和纠偏控制，不能把机器拟人化后切断责任链。
结果证据	公开裁决广泛引用；本指南将其作为责任倒挂反例。
边界与限制	个案法域与具体事实有限，不能直接等同所有国家法律结论。
权责启示	客户可见 AI 输出必须有明确企业 A、可纠正渠道和知识版本责任。

案例5：Rite Aid：监督存在仍可能系统性失灵

表25 案例卷宗：Rite Aid：监督存在仍可能系统性失灵

项目	内容
主体与时间	Rite Aid / 美国联邦贸易委员会；2012—2023
业务情境	零售商在数百家门店部署人脸识别，错误匹配对消费者造成羞辱、排斥和其他伤害风险。
具体动作	FTC 指控其缺乏合理程序与保障；和解禁止其五年内使用相关技术，并要求风险评估、测试、监控、培训和消费者投诉机制。
作用机制	问题不是“有没有操作员”，而是测试情境、偏差、培训、告警与纠偏系统整体失效。
结果证据	FTC 官方案件与命令；属于执法和解，不是法院对全部事实的终局判决。
边界与限制	人脸识别并非典型营销内容，但直接说明客户识别、个性化和线索评分的责任风险。
权责启示	human in the loop 必须被评估有效性，不能作为免责标签。

案例6：Snap My AI：上线前影响评估不能后补

表26 案例卷宗：Snap My AI：上线前影响评估不能后补

项目	内容
主体与时间	Snap / 英国信息专员办公室；2023—2024
业务情境	Snap 将 My AI 从付费用户扩展至全部用户，涉及包括未成年人在内的大规模个人数据处理。
具体动作	ICO 因数据保护风险评估不足展开调查并发出初步执法通知；后续在 Snap 改进风险评估后结束调查，同时强调组织不得忽视数据保护风险。
作用机制	影响评估必须发生在线上与扩容决策之前，并随用户范围变化重新评估。
结果证据	ICO 官方调查结论：最终未作罚款，说明纠偏可以改变监管结果。
边界与限制	不代表该产品所有风险均被消除。
权责启示	授权期限与用户范围必须绑定；从小范围扩到全量应触发再批准。

案例7：Google 奥运广告撤下：创意效率不能替代品牌判断

表27 案例卷宗：Google 奥运广告撤下：创意效率不能替代品牌判断

项目	内容
主体与时间	Google；2024
业务情境	奥运期间播出的 Gemini 广告展示父亲借助 AI 为女儿写信，遭到公众批评，认为弱化了儿童真实表达。
具体动作	Google 随后撤下广告，并表示广告意在展示 AI 可作为创意起点。
作用机制	技术能力正确并不等于品牌语境正确；缺少对文化、情感与受众反应的专业 eXamine，会把可执行内容变成品牌反噬。
结果证据	公司公开回应与广泛媒体报道；未披露销售或品牌资产量化影响。
边界与限制	舆论反应具有情境性，不能推导所有 AI 辅助创意都应禁止。
权责启示	高曝光品牌内容必须保留人类品牌判断与反向挑战，而非只做事实/合规审核。

案例8：Coca-Cola AI 节日广告：规模化生产与真实性张力

表28 案例卷宗：Coca-Cola AI 节日广告：规模化生产与真实性张力

项目	内容
主体与时间	The Coca-Cola Company；2024
业务情境	品牌用生成式 AI 制作节日广告，引发对视觉质感、情感真实性和创意劳动的争议。
具体动作	品牌继续探索 AI 创意，同时面对公众对“机器感”和经典资产处理的批评。
作用机制	AI 显著降低变体与制作门槛，却会放大品牌资产不可逆性；产能控制与品牌审批必须分离。
结果证据	公开广告与公司/媒体披露；未披露完整制作流程、成本和商业结果。
边界与限制	外部无法确认各版本具体生成比例与内部审核机制。
权责启示	越能规模化生成，越需要明确品牌 A、素材权利证据、受众测试和停止阈值。

BOARD BRIEF 10

第9章 十二个营销 workflows RACI-X 示例

竹势 AI 营销智库出品

内容	广告	社媒	线索	CRM	客服
洞察	定价	活动	舆情	代理商	数据

表29 十二个工作流 RACI-X 总表

工作流	R	最终A	eXecute	eXamine	eScalate	授权档
洞察扫描	市场情报经理	CMO/策略负责人	AI 抓取与聚类	分析师复核来源	重大负面/证据冲突	L2
选题推荐	内容策略经理	内容负责人	AI 生成候选	品牌/业务复核	敏感议题	L2
内容生成	内容编辑	内容负责人	AI 起草与适配	事实+品牌复核	证据不足	L3
品牌审核	品牌经理	CMO/品牌总监	AI 规则预检	品牌专家	核心资产偏离	L3-L4
自动发布	渠道运营	渠道负责人	AI 定时发布	发布前检查/抽样	账号异常/高风险内容	L2-L3
社媒互动	社群运营	客户体验负责人	AI 回复白名单问题	抽样+投诉复核	情绪升级/承诺/退款	L2-L4
广告预算	投放经理	CMO或P&L负责人	AI 在额度内调整	财务/投放复核	超额度/异常消耗	L3-L4
线索评分	增长运营	销售负责人	AI 评分与路由	数据/销售抽样	偏差或关键客户	L3
CRM触达	CRM运营	客户经营负责人	AI 个性化与发送	隐私/品牌复核	敏感数据/退订异常	L3-L4
客户数据调用	数据分析师	数据处理业务负责人	AI 查询授权数据	隐私/安全复核	目的变化/敏感信息	L4
舆情响应	公关经理	CMO或危机指挥官	AI 监测和草拟	法务+事实复核	重大危机	L4
代理商交付	代理商项目经理	内部营销项目负责人	代理商工具执行	品牌/法务/采购复核	权利不清/数据外传	L3-L4

表30 洞察扫描权责契约示例

项目	配置
决策对象	洞察扫描的输出、外部动作及相关数据使用
R	市场情报经理
A	CMO/策略负责人
机器执行	AI 抓取与聚类
专业复核	分析师复核来源
升级条件	重大负面/证据冲突
授权方式	L2
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表31 选题推荐权责契约示例

项目	配置
决策对象	选题推荐的输出、外部动作及相关数据使用
R	内容策略经理
A	内容负责人
机器执行	AI 生成候选
专业复核	品牌/业务复核
升级条件	敏感议题
授权方式	L2
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表32 内容生成权责契约示例

项目	配置
决策对象	内容生成的输出、外部动作及相关数据使用
R	内容编辑
A	内容负责人
机器执行	AI 起草与适配
专业复核	事实+品牌复核
升级条件	证据不足
授权方式	L3
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表33 品牌审核权责契约示例

项目	配置
决策对象	品牌审核的输出、外部动作及相关数据使用
R	品牌经理
A	CMO/品牌总监
机器执行	AI 规则预检
专业复核	品牌专家
升级条件	核心资产偏离
授权方式	L3–L4
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表34 自动发布权责契约示例

项目	配置
决策对象	自动发布的输出、外部动作及相关数据使用
R	渠道运营
A	渠道负责人
机器执行	AI 定时发布
专业复核	发布前检查/抽样
升级条件	账号异常/高风险内容
授权方式	L2-L3
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表35 社媒互动权责契约示例

项目	配置
决策对象	社媒互动的输出、外部动作及相关数据使用
R	社群运营
A	客户体验负责人
机器执行	AI 回复白名单问题
专业复核	抽样+投诉复核
升级条件	情绪升级/承诺/退款
授权方式	L2-L4
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表36 广告预算权责契约示例

项目	配置
决策对象	广告预算的输出、外部动作及相关数据使用
R	投放经理
A	CMO或P&L负责人
机器执行	AI 在额度内调整
专业复核	财务/投放复核
升级条件	超额度/异常消耗
授权方式	L3–L4
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表37 线索评分权责契约示例

项目	配置
决策对象	线索评分的输出、外部动作及相关数据使用
R	增长运营
A	销售负责人
机器执行	AI 评分与路由
专业复核	数据/销售抽样
升级条件	偏差或关键客户
授权方式	L3
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表38 CRM触达权责契约示例

项目	配置
决策对象	CRM触达的输出、外部动作及相关数据使用
R	CRM运营
A	客户经营负责人
机器执行	AI 个性化与发送
专业复核	隐私/品牌复核
升级条件	敏感数据/退订异常
授权方式	L3–L4
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表39 客户数据调用权责契约示例

项目	配置
决策对象	客户数据调用的输出、外部动作及相关数据使用
R	数据分析师
A	数据处理业务负责人
机器执行	AI 查询授权数据
专业复核	隐私/安全复核
升级条件	目的变化/敏感信息
授权方式	L4
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表40 舆情响应权责契约示例

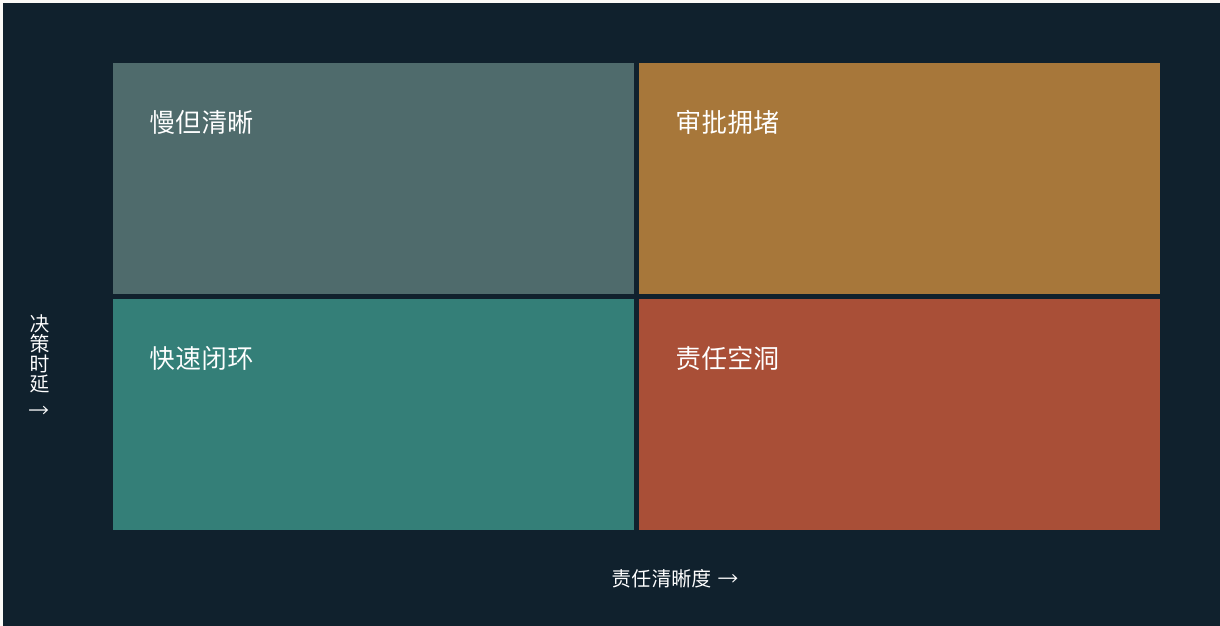
项目	配置
决策对象	舆情响应的输出、外部动作及相关数据使用
R	公关经理
A	CMO或危机指挥官
机器执行	AI 监测和草拟
专业复核	法务+事实复核
升级条件	重大危机
授权方式	L4
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

表41 代理商交付权责契约示例

项目	配置
决策对象	代理商交付的输出、外部动作及相关数据使用
R	代理商项目经理
A	内部营销项目负责人
机器执行	代理商工具执行
专业复核	品牌/法务/采购复核
升级条件	权利不清/数据外传
授权方式	L3–L4
证据包	目标、输入来源、规则版本、输出差异、批准理由、日志
SLA	低风险按小时；高风险按事件指挥机制
停止权	A、相关风险职能和系统所有者按对象拥有

第10章 决策延迟与责任空洞热区图

竹势 AI 营销智库出品



权责设计不能只看事故，也要看速度。Galbraith 的信息处理观点指出，不确定性越高，组织越需要增加信息处理能力，而不是简单增加层级；Simon 的有限理性提醒批准者不可能处理无限信息。决策延迟热区通常来自多重审批、重复复核、隐形否决与信息包不完整。责任空洞则来自多人 A、供应商灰区、系统权限无人负责和异常无人接管。

表42 热区识别矩阵

热区	可观察信号	根因	修复
多重审批	同一对象连续三次以上批准	对象未拆分、风险偏好不清	按对象设单一批准者
重复复核	不同部门检查同一内容	标准未共享	统一证据包与差异化复核
无人问责	事故后先找会议	最终 A 无资源或未指定	绑定资源权与纠偏义务
隐形否决	C 长期拖延不回复	咨询权变成事实否决	咨询 SLA，超时升级而非默认批准
代理商灰区	内部与外部互称对方负责	合同与运行矩阵脱节	共同责任表+内部 A
平台灰区	以平台默认设置作为控制	缺乏内部系统所有者	权限清单与版本责任

表43 决策时延指标

指标	定义	管理用途
端到端决策时延	从证据齐备到行动生效	识别整体速度
等待占比	等待审批时间/总周期	定位层级堵点
一次通过率	无需返工的批准比例	判断输入与标准质量
无理由批准率	无实质理由的批准占比	发现机械盖章
例外升级率	进入异常流程的比例	判断阈值与稳定性
越权尝试率	被系统阻断的非授权动作	判断权限配置与培训

第11章 指标、审计与董事会治理

竹势 AI 营销智库出品



董事会不应逐项审批营销 AI，但应批准风险偏好、重大自动化边界、最终问责轨道和季度审议机制。COSO 内控框架强调控制环境、风险评估、控制活动、信息沟通与监督；Deming 的系统观则提醒管理层不要只处罚个体，而应修复产生错误的系统。季度审议应同时看经营结果与控制健康度，避免“零事故”通过压制使用实现。

表44 权责设计平衡计分卡

维度	核心指标	警戒信号
速度	决策时延、等待占比	低风险任务持续排队
质量	返工率、缺陷率、复核有效率	通过率极高但事故增加
控制	越权、停机、回滚、日志完整性	越权未阻断或无回滚
学习	例外关闭周期、复盘落地率	同类事故重复
经营	收入、转化、成本、客户体验	提效但业务结果无变化
公平与信任	投诉、差异影响、标识合规	弱势群体负面影响

表45 董事会季度审议问题

问题	所需证据
哪些 AI 动作已经获得客户可见或资金执行权？	场景清单、授权档位、额度
每个高风险对象的唯一最终 A 是谁？	权责契约卡
本季度发生了哪些停机、越权与重大例外？	事件与复盘
哪些审批在制造延迟而未降低风险？	时延热区图
供应商变更是否触发再测试和再授权？	版本与变更记录
哪些风险被正式接受，何时到期？	风险接受台账

第12章 情境裁剪：规模、地域、监管与成熟度

竹势 AI 营销智库出品

SIZE 组织规模	REGION 地域规则	RISK 监管强度	MATURITY 能力成熟度
--------------	----------------	--------------	-------------------

表46 组织情境裁剪

情境	推荐设计	不可省略
小团队	岗位可兼任，但对象必须拆分；CEO/CMO 可为多个对象 A	机器不得成为 A；高风险双钥匙可由外部专业顾问补位
跨国公司	集团红线+区域执行权；数据跨境与本地监管单独轨道	区域最终 A、集团停止权、变更通知
高监管行业	专业复核、风险接受、双钥匙和留痕更强	适用性验证、客户权益与申诉
低成熟度	先限于建议与内部执行，建立场景清单	账号、数据、日志和停止权
高成熟度	扩大受限执行，强化自动监控和持续测试	版本化授权与再认证

表47 AI 成熟度与授权演进

阶段	机器权力	人类控制重点	升级条件
探索	建议权	事实核验与场景边界	稳定质量、无敏感数据
试点	低风险执行权	逐件/抽样复核、日志	缺陷率与恢复达标
运行	额度内执行权	异常监控、抽样、风险接受	跨周期稳定、审计通过
规模化	跨系统受限执行	版本治理、供应商与组合风险	自动控制有效且可回滚
自治受控	高频闭环执行	董事会风险偏好与持续保证	仅限高度标准化可逆任务

第13章 30/60/90 天权责改造台

竹势 AI 营销智库出品

DAY 01–30 盘点对象与空洞	DAY 31–60 重配令牌与阈值	DAY 61–90 运行证据与再认证
----------------------	----------------------	-----------------------

路线图以真实决策流为起点。Argyris 的双环学习说明组织不能只修正一次错误，还要检视产生决策的目标、规则和假设；OODA 强调观察、判断、决策、行动的循环速度。90 天的目标不是完成制度文件，而是让至少三类真实场景在压力测试下运行并留下证据。

表48 90天权责改造总台

阶段	重点	交付物	退出门槛
0—30天	画真实决策流、盘点权限与事故	场景清单、现状 RACI、热区图、风险分级	确认三个试点和唯一 A
31—60天	三类场景试运行、红队与停机演练	RACI-X、权力令牌、证据包、演练记录	能暂停、回滚、归责
61—90天	调整授权、指标、系统权限与制度	制度、权限配置、季度看板、供应商附录	经营与控制指标同时达标

表49 三类试点组合

试点类型	示例	目的
低风险高频	内部摘要、素材标签、SEO 元数据	验证去审批与抽样复核
客户可见中风险	常规内容发布、社媒标准回复	验证专业复核、升级和回滚
高风险关键动作	预算调整、敏感数据触达、危机响应	验证双钥匙、停止与事件归责

表50 红队压力测试清单

测试	预期控制
模型输出虚假事实	eXamine 拦截并记录来源缺口
账号权限被扩大	系统自动拒绝并通知系统所有者
预算数据延迟或异常	停止自动调整，转人工核验
客户请求删除/退订	立即停止触达并同步系统
供应商模型版本突变	冻结高风险场景并回归测试
批准者连续快速通过	触发监督疲劳告警与抽查

第14章 管理工具与角色清单

竹势 AI 营销智库出品

权责契约卡

对象 · 角色 · 令牌 · 额度 · 期限 · 证据 · 升级 · 停止 · 恢复 · 复盘

表51 权责契约卡模板

字段	填写要求
决策对象	精确到输出、动作、数据与金额
RACI-X	每个角色与动作可观察
权力令牌	建议/执行/批准/停止/风险接受
证据包	来源、测试、版本、差异、批准理由
SLA	正常、异常、危机三档
升级阈值	金额、客户、数据、品牌、监管
停止与恢复	谁可停、恢复证据与批准者
日志	身份、时间、输入、输出、工具、版本
复盘	频率、负责人、规则变更
版本责任	谁维护并何时再认证

表52 董事会/CEO清单

行动	评审问题	指标
批准风险偏好	哪些动作绝不自动化？	高风险场景覆盖率
指定最终 A	A 是否有资源与纠偏权？	无人 A 数量=0
季度审议	速度和风险是否平衡？	时延、事故、经营结果
要求退出能力	供应商或模型失效能否停用？	恢复与迁移时间

表53 CMO清单

行动	评审问题	指标
拆分品牌与经营对象	谁对品牌不可逆性负责？	品牌返工与事故
建立内容证据包	事实、权利、标识是否齐全？	一次通过率
优化低风险审批	哪些可以抽样而非逐件？	等待占比
危机停止权	谁能在分钟级暂停？	平均遏制时间

表54 CIO/CDO清单

行动	评审问题	指标
最小权限	AI 真实可调用什么？	越权尝试/阻断率
版本治理	模型/连接器变化是否再认证？	未评估变更数
日志与回滚	能否还原事件链？	日志完整率、恢复时间
供应商控制	是否获得变更与事件证据？	供应商 SLA

表55 法务/内控清单

行动	评审问题	指标
定义否决边界	哪些法律/风险红线必须停？	红线覆盖率
审查风险接受	谁有权接受何种残余风险？	到期未复核数
验证复核有效性	审核者是否有能力与信息？	复核后缺陷率
事件归责	是否沿五问而非找替罪羊？	重复事故率

表56 CHRO与执行团队清单

行动	评审问题	指标
岗位卡纳入权力令牌	员工知道能做、不能做和如何申诉吗？	授权理解度
复核者资格	谁具备专业复核能力？	认证覆盖率
执行者申诉	不合理目标或规则如何升级？	申诉关闭周期
绩效不奖机械通过	是否奖励发现问题和安全停机？	有效停止/改进数

表57 外部伙伴清单

行动	评审问题	指标
明确交付边界	可直接发布还是只能交草稿？	越权事件
提供证据	来源、版权、模型与版本是否可核验？	证据完整率
事故协同	多久通知、谁保全证据？	通知与恢复时长
退出移交	权限、数据、资产如何回收？	回收完成率

专题深化：权责系统的七个运行难题

竹势 AI 营销智库出品

运行难题

真正的治理质量，取决于制度在压力、异常和版本变化时如何运作。

一、批准者无信息：证据包必须服务决策，而非堆积材料

批准失败常常不是批准者不负责，而是其看到的信息不支持判断。内容审批页面只展示最终文案，却隐藏来源、机器改写幅度、受众范围、历史版本和敏感数据使用；预算审批只展示建议值，却不展示数据延迟、归因不确定性和回滚成本。有效证据包应采用“决策摘要—关键差异—风险与替代方案—原始证据索引”的分层结构。批准者需要知道自己正在批准什么对象、在什么假设下生效、有哪些尚未消除的风险，以及失败后能否撤回。证据越多不等于信息越充分；没有围绕决策对象组织的材料只会制造机械盖章。

二、审核者无能力：复核权必须与资格、时间和工具绑定

eXamine 不是把任意员工放进审批流。事实复核需要来源识别能力，品牌复核需要理解品牌资产和受众语境，隐私复核需要理解处理目的、数据分类与法律基础，投放复核需要识别归因偏差和异常消耗。企业应为不同复核对象建立资格标准、最大负荷与辅助工具。连续审批数量过高、平均审阅时长过短、批准理由高度重复，都是监督疲劳或机械盖章信号。对于高风险场景，复核者还应能够直接触发停止，而不是只能提出建议后等待业务负责人处理。

三、执行者无申诉权：安全升级不能被绩效压力压制

一线执行者最早看到规则冲突、客户反应和系统异常，却常因没有正式申诉通道而继续执行。若团队绩效只奖励发布数量、触达量或预算消耗，员工会倾向于忽视不确定性。权责契约应明确“善意停止”保护：执行者基于证据暂停高风险动作，不因短期目标未完成受到惩罚；升级后必须在 SLA 内获得裁决；裁决和规则变更写回系统。申诉权不是削弱最终 A，而是增加前线传感能力。最终 A 仍需快速裁决、调配资源并承担延迟或继续运行的后果。

四、模型与 workflow 版本变化：授权不能永久继承

同一名称的模型、连接器或 workflow 在版本更新后可能表现不同。供应商调整安全策略、上下文窗口、工具调用逻辑或默认数据保留设置，都会改变既有风险判断。因此授权必须绑定版本或可识别配置，并设置重大变更触发器：模型主版本变化、数据源变化、客户范围扩大、从草稿生成升级为自动发布、从内部数据升级为个人信息、从单一渠道扩展到跨境渠道。触发后不是所有流程停摆，而是高风险场景冻结，低风险场景进入加密抽样，直至回归测试完成。

五、停止权的悖论：能停机但无人敢停，等于没有停止权

停止权必须在技术、制度和激励三个层面同时存在。技术上要有一键暂停、账户撤权、预算上限、队列冻结和版本回滚；制度上要写明谁可在什么证据下停止、通知谁、何时复核；激励上要保护善意停止，避免把停机视为“拖业务后腿”。高可靠组织研究强调对异常保持敏感、尊重一线专业判断。营销组织可以借鉴这一点：重大舆情、异常消耗、客户投诉突增、数据来源不明和标识失效，都应有预先约定的停止触发。恢复运行需要新的证据，而不是领导口头要求“先上线再说”。

六、共同责任的错觉：每一方负责一部分，不代表整体有人负责

代理商负责制作、平台负责模型、企业 IT 负责接入、市场部负责发布，这种分工看似完整，实际可能没有任何一方对端到端客户影响负责。共同责任必须通过接口契约连接：每一方说明其控制范围、提供证据、通知变更和处理事故；企业内部仍指定一个对上线采用与经营结果负责的 A。该 A 不需要亲自掌握所有技术细节，但必须有权暂停供应商、调整预算、召集专业人员、通知客户和改变流程。没有这些能力的“项目联系人”不能承担最终问责。

七、零事故不是唯一目标：治理要允许受控学习

若管理层只追求零事故，最简单的做法是禁止 AI 接触真实客户和真实系统，企业也就失去学习机会。更合理的目标是把风险暴露限制在可接受范围内，并提高发现、遏制、恢复和学习能力。低风险场景可用小受众、低额度、短授权和可回滚设计获得真实反馈；高风险场景先在沙箱、历史数据和红队环境中验证。指标既要记录事故，也要记录有效停止、及时升级、快速恢复和规则改进。成熟治理的标志不是从不出错，而是错误不会在无人知情、无人负责和无法回滚的状态下扩散。

表60 证据包分层模板

层级	内容	决策价值
决策摘要	对象、建议动作、受众、金额、时点	让批准者明确正在决定什么
关键差异	与已批准版本、规则或基线的变化	把注意力集中于新增风险
风险与替代	残余风险、替代方案、回滚方案	支持真正取舍而非二元点击
验证证据	来源、测试、样本、异常与限制	判断专业可靠性
运行条件	额度、期限、监控、停止阈值	把批准转成受限授权

表61 授权再认证触发器

触发器	默认动作	再认证证据
模型主版本变化	冻结L4场景，L2–L3加密抽样	回归测试、已知限制、差异说明
数据范围扩大	暂停新增数据调用	目的、最小必要、合法性与权限
受众规模扩大	降低额度并重新批准	小范围结果、投诉与偏差
草稿升级为自动执行	重新风险分级	回滚、监控、停止演练
供应商或子处理者变化	暂停敏感数据流	合同、位置、安全和删除机制
监管或平台规则变化	更新规则并验证	合规解释、规则版本、培训

表62 善意停止与申诉机制

要素	制度要求	反面信号
触发资格	R、eXamine、系统所有者和风险职能均可按对象停止	只有最高领导能停
证据门槛	合理怀疑即可临时暂停，高风险无需先证明损害	必须等事故发生
响应SLA	明确分钟/小时级裁决	升级后长期无人回应
绩效保护	善意停止不计为执行失败	因暂停扣绩效
闭环	裁决、恢复与规则变更留痕	口头恢复、无复盘

专题深化：六个高频营销场景的阈值化设计

竹势 AI 营销智库出品

阈值化场景

按影响面、额度、客户承诺和数据风险配置不同的控制强度。

一、内容自动发布：按影响面而非内容长度分级

内容发布风险不应按“长文、短文”区分，而应看受众规模、品牌资产、事实密度、监管主题和撤回难度。日常活动提醒可以在已批准模板、固定受众和可撤回渠道内采用 L2；涉及价格、功效、客户权益、竞争比较或高管表态的内容至少进入 L3；金融、医疗、未成年人、危机回应和重大公司事项通常进入 L4。系统应自动比较当前稿与已批准模板的差异，突出新增事实、数字、承诺和敏感表述，复核者只需聚焦变化。发布后监控不应只看阅读量，还应看投诉、纠错、删除、平台警告和异常扩散。

二、广告预算调整：额度、节奏与数据质量共同决定执行权

预算自动化最容易被误写为“AI 可在预算范围内优化”。真正可运行的授权至少包含账户、活动、日上限、单次调整幅度、累计调整幅度、最短观察窗口、归因数据新鲜度和异常消耗阈值。低风险成熟账户可允许 AI 在日预算的 5%—10% 内调节，并要求每次调整后保留观察窗口；新品、重大促销和高波动账户应降低额度或采用事前批准。若转化数据延迟、像素失效、渠道口径变化或异常点击上升，系统必须停止自动调节。这里的核心不是给出统一百分比，而是把额度与数据可靠性、回滚能力和 P&L 承受力绑定。

三、社媒互动：把“标准问题”与“企业承诺”分开

社媒自动回复可以覆盖营业时间、公开产品信息、活动规则入口和已批准的常见问题，但涉及退款、赔偿、合同、健康建议、价格承诺、个人情况和情绪升级时必须转入人工。企业应维护“可直接回答、需要检索后回答、必须升级、禁止回答”四类意图表。机器不得通过连续对话逐步拼接出超出单条授权范围的承诺。对客户表达不满、威胁投诉、提及媒体或监管机构的对话，应同时触发客户服务和舆情轨道，而非继续由机器人安抚。

四、线索评分与客户触达：营销效率不能覆盖公平与隐私

线索评分的最终输出往往影响销售资源分配、价格优惠和客户体验，因此不能只由模型团队负责。业务负责人应定义评分用途和可接受差异，数据团队负责数据质量与特征来源，隐私职能审查目的和最小必要，销售负责人对资源分配结果 A。模型不得把敏感属性或其明显代理变量用于不合理差别待遇。企业应按客户群体比较触达率、拒绝率、转化率和投诉，而不是只看总体 AUC 或转化提升。客户退订、反对自动化处理或要求人工沟通时，应有明确停止和申诉通道。

五、舆情响应：危机指挥权必须预先指定

舆情场景中最大的责任空洞是“大家都在看，没有人能定”。日常监测可由 AI 自动执行并由公关团队抽样复核；当出现高传播速度、主流媒体介入、监管点名、重大事实争议或人身安全风险时，应进入危机模式。危机指挥官成为该事件沟通对象的唯一 A，法务拥有违法披露停止权，业务负责人提供事实与补救资源，AI 只能汇总证据、生成版本和模拟反应。任何自动发布、自动删除或批量回应都应冻结。恢复常规模式必须记录事实结论、对外口径、未决风险和后续监控周期。

六、代理商交付：把创意所有权、工具使用与发布责任拆开

代理商可能使用自有模型、平台或外包制作链路，企业不能只要求“保证合规”。权责契约需分别明确素材来源与版权、客户数据是否可进入外部工具、生成内容标识、品牌规则版本、可否直接操作企业账号、谁保存中间稿与生成记录、供应商变更是否通知。代理商对其制作过程和声明真实性负责，企业内部项目负责人对是否采用和上线 A。若企业要求代理商使用指定工具或在极短时限内上线，也应承担相应的目标与资源责任，不能把所有事故风险单向转嫁。

表63 六类场景阈值示例

场景	自动执行边界	必须升级条件	最终A
内容发布	已批准模板、低风险主题、小范围、可撤回	新增事实/承诺、强监管、危机、高曝光	内容/品牌负责人
预算调整	白名单账户、额度内、数据新鲜、可回滚	异常消耗、数据延迟、超幅度、重大活动	P&L或投放负责人
社媒互动	公开FAQ、无个性化承诺	退款、投诉、健康/法律、情绪升级	客户体验负责人
线索评分	已批准特征、仅作优先级建议	敏感属性、重大权益、群体差异异常	销售/客户经营负责人
舆情响应	监测、聚类、内部草拟	监管、媒体、快速扩散、事实冲突	危机指挥官
代理商交付	草稿制作、限定工具和数据	直接发布、版权不清、数据外传、供应商变化	内部项目负责人

表64 批准者最小信息集

对象	必须看到	不得隐藏
内容	受众、事实来源、版本差异、标识、撤回方案	模型生成比例、关键新增承诺
预算	当前消耗、归因时延、调整幅度、最坏损失	数据缺口、异常流量
数据调用	目的、字段、客户范围、保留期、处理方	敏感字段和跨境路径
客户触达	人群规则、频次、退订、投诉与人工接管	排除名单和弱势群体影响
供应商变更	版本、子处理者、已知限制、迁移影响	默认设置变化

专题深化：治理中的四条边界

竹势 AI 营销智库出品

第一，专业否决权不等于经营接管。法务、隐私、安全、医学或财务控制人员应能在其专业红线内阻断行动，但不应因此成为活动收入、客户增长或品牌表现的最终 A。经营负责人必须在红线内重新选择目标、方案和资源。这样既保护专业独立性，也防止业务把失败归咎于控制职能。

第二，风险接受权不能口头下放。接受残余风险意味着组织明知控制不能完全消除风险，仍决定继续，因此必须限定风险类别、损失上限、持续时间和补偿措施。到期后若没有重新评估，授权应自动失效。对可能违法、侵害基本权利或超出组织风险偏好的事项，不存在可由普通业务负责人接受的空间。

第三，日志保留不能无限扩张。为了问责而记录输入、输出、身份、版本和操作是必要的，但日志本身可能包含个人信息、商业秘密和敏感内容。企业应定义最小记录字段、访问权限、保留期限和删除程序，并把审计可得性与数据最小化同时纳入设计。记录一切既不经济，也可能制造新的泄露面。

第四，自动化禁区应保持可解释和可更新。禁止自动化清单不能只写抽象伦理原则，而应列出具体动作，例如未经授权抓取敏感客户数据、生成虚假评价、自动作出对客户有重大不利影响且无申诉渠道的决定、绕过平台标识、在危机状态下自动对外回应。禁区应随法规、业务模式和控制能力变化更新，但任何放宽都必须经过正式风险接受与压力测试。

表65 四条治理边界的责任配置

边界	拥有决定权	必要证据	升级条件
专业否决边界	专业风险岗位	法律/专业标准、事实和影响	业务与专业意见冲突
风险接受边界	经授权风险所有者	残余风险、上限、期限、补偿控制	超风险偏好或到期
日志边界	系统所有者+隐私/内控	审计目的、字段、访问与保留期	记录敏感度扩大
自动化禁区	董事会/CEO批准，职能维护	法律、权利影响、不可逆性	拟放宽或出现新场景

实施时还应坚持最小充分原则：每增加一个审批节点，都要说明它降低了哪一类具体风险；每放开一项自动执行权，都要说明其额度、期限、证据和撤销条件。无法回答这两个问题的控制，要么是形式主义，要么是未经治理的授权。

结论

竹势 AI 营销智库出品

FINAL CONTROL PRINCIPLE

机器获得的是可撤销执行权，人类承担的是不可转移的判断与纠偏责任。

AI 市场部的权责设计不能停留在“人负责、AI 辅助”这类宽泛表述。真正可运行的系统必须把机器执行、人类专业判断、经营批准、风险接受、系统授权、异常升级与最终问责拆开，再通过额度、期限、证据、停止权和复盘重新连接。RACI-X 的价值不在增加三个字母，而在把责任从静态角色变成可观察动作。

董事会应坚持一条底线：任何具体决策对象只有一个最终 Accountable 的人类岗位。这个岗位必须拥有资源配置权、风险接受权或纠偏能力；否则“最终负责”只是组织语言。与此同时，唯一 A 不意味着所有权力集中。复杂企业必须用专业复核、否决/停止、风险接受和系统授权形成制衡，并用清晰的对象边界避免多人 A。

最成熟的治理并非审批最多，而是能让低风险高频任务快速运行，让高风险任务在正确节点减速，并能在异常发生时立即暂停、回滚、归责和学习。权责系统的最终验收标准，是经营结果、决策速度和风险控制同时改善。

附录A 经典思想的管理转译

竹势 AI 营销智库出品

表58 八个经典思想及 AI 情境修正

思想	原始主张	本题连接	AI情境边界	管理转译
Galbraith 信息处理	不确定性增加要求更强信息处理能力	审批层级不是唯一答案	AI增加信息量也制造噪声	用证据包、例外路由替代层层汇报
Simon 有限理性	决策者受信息、时间和认知限制	批准者不能阅读全部输出	自动摘要可能隐藏关键差异	呈现差异、风险与替代方案
Jensen/Meckling 委托代理	目标与信息不对称产生代理成本	供应商、代理商、内部团队目标不同	模型本身无代理责任主体资格	合同+运行控制+内部A
COSO 内控	控制环境、风险评估、控制活动、沟通和监督	AI 权责应进入既有内控	传统控制需适配概率输出	嵌入制度、权限、监控与审议
Reason 瑞士奶酪	事故来自多层防线缺口对齐	不只追最后操作人	AI增加隐藏耦合	沿事件链归责并修复多层控制
Perrow 正常事故	复杂紧耦合系统会出现意外组合	跨系统智能体不可假设零事故	不等于接受可避免风险	限制耦合、设置熔断和退出
Deming 系统观	多数质量问题来自系统	不以处罚个体替代改进	仍需处理故意违规	指标与复盘聚焦系统原因
Argyris 双环学习	修正行动同时检视目标与规则	事故后要调整风险偏好和授权	不是无限讨论	把例外转成规则与权力更新

附录B 具名专家视角

竹势 AI 营销智库出品

表59 十位专家视角与本报告判断

专家	研究/经验	具体主张	适用边界	本报告判断
Herbert Simon	组织决策与有限理性	组织通过程序与信息结构应对认知限制	不直接讨论生成式AI	批准界面必须压缩信息但保留关键差异
Jay Galbraith	组织设计与信息处理	不确定性越高越需提升横向信息处理能力	经典组织情境	避免用更多层级治理 AI
James Reason	人因与事故模型	事故通常是多层防线共同失效	源自安全关键行业	适用于 AI 事件链归责
Charles Perrow	复杂系统与正常事故	复杂、紧耦合系统难以完全避免意外	理论可能过于悲观	支持限制自动化耦合与设置停机
W. Edwards Deming	质量与系统管理	管理者应改进系统而非只责怪个人	需防止弱化个人故意违规责任	复盘应优先修复制度、数据和权限
Chris Argyris	组织学习	双环学习挑战目标与规则假设	实施需要心理安全	事故复盘应更新授权模型
Elham Tabassi	NIST AI 风险管理	AI 风险管理应贯穿治理、映射、测量、管理	NIST框架自愿使用	为 RACI-X 控制回路提供结构依据
Margrethe Vestager	欧盟数字与竞争治理经验	AI治理需按风险与角色分配义务	欧盟法域	跨国企业需按部署角色和风险分级
Alvaro Bedoya	FTC消费者保护与隐私执法	有偏差的自动监控会造成真实伤害，需测试与问责	美国执法情境	客户识别与评分不能只靠人工复核标签
Geoffrey Hinton	深度学习与AI风险公共讨论	先进AI能力与风险需要更强监督	对具体企业控制未给出完整方案	企业层面应转译为限权、监控和停止能力

附录C 公开纯文本来源索引

竹势 AI 营销智库出品

OECD, 《OECD AI Principles》, 2019, 2024 更新; 政府间 AI 原则, 强调透明、稳健与问责。

NIST, Elham Tabassi 等, 《Artificial Intelligence Risk Management Framework 1.0》, 2023; 自愿性跨行业风险管理框架。

NIST, Chloe Autio、Reva Schwartz 等, 《Generative Artificial Intelligence Profile》, 2024; 生成式 AI 风险行动清单。

ISO/IEC, 《ISO/IEC 42001:2023 Artificial intelligence management systems》, 2023; AI 管理体系要求。

European Union, 《Regulation (EU) 2024/1689 (Artificial Intelligence Act)》, 2024; 分阶段适用的 AI 法规。

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; 个人信息处理、自动化决策与处理者义务。

国家互联网信息办公室等, 《互联网信息服务算法推荐管理规定》, 2021 发布、2022 施行; 算法主体责任、审核评估与备案。

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; 生成式 AI 服务提供者义务。

国家互联网信息办公室等, 《人工智能生成合成内容标识办法》, 2025; 显式、隐式标识与主体责任, 2025-09-01施行。

Federal Trade Commission, 《FTC v. Rite Aid Corporation》, 2023—2024; 人脸识别不合理保障与五年禁用和解命令。

UK Information Commissioner’s Office, 《Snap My AI investigation conclusion》, 2024; 上线前数据保护影响评估。

Tencent, 《Tencent Responsible AI Principles》, 2024—2026; 集团 AI 安全治理原则。

Tencent, 《Environmental, Social and Governance Report 2025》, 2026; Responsible AI 与风险治理披露。

Ping An Group, 《Policy Statement on AI Ethics Governance》, 2022; AI 伦理原则与治理方向。

Ping An Group, 《2025 Sustainability Report》, 2026; 金融集团 AI 与可持续治理披露。

James Reason, 《Human Error》, 1990; 瑞士奶酪模型与组织事故。

Charles Perrow, 《Normal Accidents》, 1984; 复杂紧耦合系统事故理论。

Chris Argyris, 《Double Loop Learning in Organizations》, 1977; 组织双环学习。

COSO, 《Internal Control—Integrated Framework》, 2013; 内部控制五要素。

Jay R. Galbraith, 《Organization Design: An Information Processing View》, 1974; 组织信息处理。

PUBLISHER

关于竹势 AI 营销智库

竹势 AI 营销智库出品

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

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