



ROLE MIGRATION / TASKS · RIGHTS ·
INTERFACES

AI时代市场部 岗位重构图谱

把九类岗位拆成可观察任务，决定哪些应保留、增强、协同、
执行、重组或退出，并重画判断权与结果责任。

先拆任务，再谈编制

避免把技术暴露误读为裁员比例

判断权不随模型能力下放

用影响、可逆性与证据决定授权

九类岗位各有迁移边界

品牌、内容、创意到 CRM 与管理

给出 14/30/90 天路径

任务盘点、岗位试点与组织迁移

出版说明

竹势 AI 营销智库出品

本报告面向企业董事会、经营管理层与市场组织负责人，讨论生成式 AI 与执行型智能体进入市场部后，岗位如何围绕经营结果、任务组合、判断权与协作接口重新设计。报告中“任务暴露”指技术能力可能影响任务的程度；“能力增强”指员工借助 AI 改善速度、覆盖或质量；“自动化采用”指组织把任务交给系统执行；“岗位变化”指岗位内部任务权重、责任与接口改变；“就业净影响”和“人员裁减”则属于组织层与劳动力市场层结果。六者不得相互替代。

报告使用公开可核验材料。企业披露、平台披露和供应商材料均按其披露性质使用，不把自报效率直接解释为独立因果效果。来源索引采用纯文本格式，不提供第三方可点击链接。

概念	定义	本报告使用边界
任务暴露	某类任务在技术能力上可能受到影响	预测性或结构性指标，不等于已自动化
能力增强	AI协助人完成任务	不能直接推断岗位减少
自动化采用	组织真实部署并让系统执行	受流程、数据、治理与成本约束
岗位变化	任务、权责、接口和评价方式改变	岗位名称可保留
就业净影响	新增、消失、转岗等合计变化	需要劳动力市场或企业人事数据
人员裁减	企业明确减少人员	不得从暴露比例推断

- 1. 执行摘要与十大核心判断
- 2. 第一章 总论：岗位重构的对象与边界
- 3. 第二章 岗位 DNA 四层模型
- 4. 第三章 任务拆解：建立可观察的工作单元
- 5. 第四章 ROLE-6 任务迁移矩阵
- 6. 第五章 判断权阶梯与风险治理
- 7. 第六章 九岗任务迁移图谱（上）
- 8. 第七章 九岗任务迁移图谱（下）

- 9. 第八章 新角色、共享能力与组织接口
- 10. 第九章 产能上升后的二阶瓶颈
- 11. 第十章 KPI、绩效与经营归因
- 12. 第十一章 人员转型与能力邻接
- 13. 第十二章 案例卷宗与专家视角
- 14. 第十三章 14/30/90 天实施路线与管理工具
- 15. 结语与来源索引

BOARD BRIEF 02

执行摘要与十大核心判断

竹势 AI 营销智库出品

BOARD DECISION / 董事会判断

岗位重构的最小单位是任务；真正要重画的是判断权、交付接口与结果责任，而不是按技术暴露比例删减职位。

任务先于编制

可逆性先于自动化率

判断权不随模型能力下放

经营结果高于产量

核心判断	董事会含义
岗位重构的最小单位是任务，不是人数	同一岗位内，信息搜集、初稿生成、版本适配等任务可能快速迁移，而品牌主张、预算约束、客户承诺与责任承担仍需人负责。先定编后拆任务，会把技术暴露误读为人员结论。
岗位名称可能稳定，岗位内容会剧烈重排	品牌经理、投放经理、CRM 经理等名称可以保留，但其工作重心将从亲自生产转向目标定义、系统编排、证据审核、例外处理和跨部门协调。
可评测性与可逆性比“模型能否完成”更重要	可自动执行的任务必须同时具备清晰输入、可验证输出、足够数据、低不可逆损失和明确异常升级。技术能力只是必要条件。
AI 扩产会放大判断稀缺性	内容、图像、受众组合和投放建议增加后，企业真正稀缺的是选什么、不选什么、何时停止、谁承担后果。审核吞吐若不重构，产能上升将转化为返工和风险。
算法执行不能替代预算责任	媒介和投放自动化降低操作成本，却可能强化平台目标、归因偏差和局部最优。目标函数、实验设计、增量测量、品牌安全和预算上限必须由企业掌握。
高频互动适合机器协助，客户承诺不可无人负责	社媒和 CRM 中的分类、摘要、建议和低风险回复可以自动化；退款、权益、价格、医疗金融等敏感承诺，以及危机和弱势客户场景必须设人工升级。
研究岗位不会因检索变快而失去价值，反而更依赖问题定义	信息处理速度提高后，样本偏差、证据等级、因果误判和合成信息污染成为更大风险。研究者的价值转向问题框定、方法选择、反证和证据审计。
新角色来自接口，而不是时髦头衔	AI 内容运营、工作流编排、模型风险、知识资产等角色，本质上处理跨任务接口。多数企业应先建立共享责任，再视规模专职化。
培训结业不等于岗位迁移	人员转型必须通过真实任务、可观察产出和责任升级完成。工具课可作为入门，但岗位资格应由项目作品、风险处理和经营结果证明。
组织迁移应分阶段并可退出	14 天做任务普查，30 天在单一岗位链路试点，90 天重画接口、KPI 与人员安排。每阶段都应有证据门、员工沟通和回退方案。

国际劳工组织 2025 年更新强调，生成式 AI 对职业的影响首先表现为任务暴露，且更可能带来岗位转型而非机械替代；世界经济论坛的 2025 年雇主调查则给出到 2030 年岗位创造与岗位消失并存的情景估计。两类证据都不支持把暴露比例直接改写为裁员事实。NIST AI RMF 与其生成式 AI 配套文件进一步表明，组织在部署时必须把治理、情境映射、测量和管理连成闭环。

第一章 总论：岗位重构的对象与边界

竹势 AI 营销智库出品

OLD

岗位名称

职责宽泛 · 接口隐形 · 结果难归因



NEW

任务责任单元

输入 · 动作 · 输出 · 判断 · 异常 · owner

1.1 反共识：不要用“自动化率”决定编制

市场部长按专业分工建立岗位：品牌负责主张，内容负责表达，创意负责概念，媒介负责触达，投放负责账户，社媒负责互动，CRM 负责关系，研究负责证据，管理负责资源与问责。生成式 AI 与智能体跨越这些边界，使一个系统可以同时检索、生成、适配、执行和汇总。但跨越任务边界不等于跨越责任边界。自动化率只描述动作被机器承担的比例，无法回答动作是否有价值、错误是否可逆、谁批准、谁对客户与预算负责。

Drucker 的目标与责任观提醒管理者，岗位不是活动清单，而是对结果的承诺。Mintzberg 对管理工作的观察则说明，真实工作由短促、碎片化、关系密集的活动构成，正式职位说明书常常不能反映异常处理与协调负担。AI 进入后，这种差异更大：显性生产任务减少，隐性审查、接口和例外任务增加。

误读	真实含义	董事会纠偏
“某岗位 60% 暴露，所以可裁 60%”	暴露是潜在技术影响，不是采用、替代或净就业结果	要求任务级采用证据与业务容量分析
“每人配一个助手即可”	个人效率可能提高，但流程接口和责任未改变	以端到端任务链试点
“产出翻倍就是成功”	可能同步增加低价值产出、审核和资产混乱	加入经营结果与质量成本
“平台自动优化，所以企业无需投手”	平台执行目标不等于企业增量和风险目标	保留预算、实验与终止权

1.2 六种结果必须区分

层次	观察问题	可用证据	不能推出
技术暴露	模型理论上能影响哪些任务	任务评估、基准测试	真实采用率
员工增强	员工是否更快或更好	实验、对照任务、质量评分	岗位减少
流程采用	系统是否进入正式流程	日志、审批、系统记录	净就业影响
岗位变化	任务权重和责任是否改变	新岗位卡、RACI、时间分配	行业就业趋势
组织人员变化	招聘、转岗、自然流失或裁减	人力资源数据与公告	技术单一归因
宏观净影响	职业增长与消失合计	劳动力统计与雇主调查	单个企业结论

BOARD BRIEF 04

第二章 岗位 DNA 四层模型

竹势 AI 营销智库出品

01

经营结果

岗位为何存在

02

任务组合

具体做什么

03

决策权

谁能选择与否决

04

协作接口

与谁交换什么

技能是支撑层，不能替代岗位设计。

岗位 DNA 四层模型依次是经营结果、任务组合、决策权和协作接口。技能是支撑层，不能反过来定义岗位。企业常见错误是先行“会使用哪些 AI 工具”，再把工具能力包装成新岗位；结果是岗位目标不清、权限不明、与销售和法务反复冲突。

层级	关键问题	输出物	常见失败
经营结果	岗位为哪项收入、品牌、客户或风险结果负责	结果树、北极星指标	只写活动数量
任务组合	为结果持续执行哪些任务	任务词典、频率与输入输出	沿用旧 JD
决策权	哪些判断可建议、共决、限权执行或禁止机器决定	判断权阶梯	默认下放
协作接口	向谁取数、交付、升级和接受审查	接口契约、SLA、RACI	无人处理异常

2.1 岗位卡的重写规则

先写结果，再写任务；先写例外，再写日常。

把“负责内容运营”改写为可观察的输入、动作、输出、判断、频率和异常。

每项任务必须绑定一个结果所有者和一个异常所有者。

任何机器执行权都必须附带额度、范围、可逆性、日志和停机条件。

岗位卡字段	示例：投放组合经理
经营结果	在预算和品牌安全约束下提高增量获客价值
核心任务	目标设计、实验组合、预算分配、创意学习、异常升级
机器权限	生成建议、形成草稿、在限额内暂停异常素材
人工终审	预算跨账户迁移、归因口径、品牌安全与重大加码
关键接口	财务、数据、法务、平台、创意、销售
评价方式	增量效果、风险事件、实验学习速度、预算偏差

第三章 任务拆解：建立可观察的工作单元

竹势 AI 营销智库出品

触发

何时开始

输入

数据与约束

动作

人机步骤

输出

完成定义

判断

选择与否决

异常

升级与恢复

O*NET 与 ESCO 等职业体系的价值在于把职业拆为任务、知识、技能与工作活动；企业落地时仍需进一步加入频率、对象、系统、判断和异常。过粗会看不见迁移机会，过细则造成编排成本和责任碎片化。建议采用“一个可独立验收、可分配责任、可记录异常的工作单元”作为最小粒度。

任务字段	说明	示例
触发条件	何时开始	新产品立项或数据异常
输入	必需材料与数据	产品卖点、受众、历史效果
动作	人或系统执行步骤	检索、归纳、生成、审核、发布
输出	可验收产物	策略简报、素材、账户草稿
判断点	需要取舍的节点	是否采用主张、是否加预算
异常	超出规则的情形	敏感舆情、数据缺失、成本突增
频率与容量	日/周/月以及峰值	每日 20 条、季度战役
证据	完成与质量记录	日志、版本、审批、效果数据

3.1 任务日志与价值流

14 天任务普查不要求员工记录每一分钟，而是覆盖典型周、峰值周与异常事件。记录“开始原因、等待对象、返工次数、风险判断、输出用途”比记录工时更有价值。Theory of Constraints 提醒企业，局部提速会把瓶颈推向下一环节；因此任务普查必须同时记录等待和队列。

观察维度	诊断信号	可能动作
高频低判断	重复适配、整理、格式化	增强或自动执行
低频高风险	危机回应、重大预算	保留人工并强化演练
高等待	依赖法务、数据或领导审批	重画接口与授权
高返工	Brief 不清或标准冲突	前移约束与样例
高隐性劳动	协调、催办、解释	建立共享看板与接口契约

第四章 ROLE-6 任务迁移矩阵

竹势 AI 营销智库出品

R 保留 高责任判断
O 增强 人主导、AI扩展
L 协同 跨角色交接
E 执行 低风险、可评测
R 重组 任务重新打包
E 退出 无经营价值

ROLE-6 将任务迁移分为保留 Retain、增强 Orchestrate、协同 Link、自动执行 Execute、重组 Rebundle、退出 Exit。六类不是成熟度阶梯，而是不同的设计结果。高价值、高判断密度任务可以保留；高价值且可被 AI 扩展的任务适合增强；跨角色任务需要协同；低风险、可评测任务可以执行；因新技术而需要重新组合的任务进入重组；对经营无贡献或由历史流程遗留的任务应退出。

类别	适用条件	例子	控制要求
R 保留	高判断、高风险、低可逆	品牌主张终审、危机立场	双人复核与责任人
O 增强	高价值、人主导、AI 扩展	洞察综合、创意探索	证据与版本记录
L 协同	跨系统、跨角色、需交接	营销到销售线索移交	接口 SLA 与异常所有者
E 自动执行	规则清晰、可评测、可回滚	格式适配、标签、低风险报表	额度、日志、回滚
R 重组	旧岗位任务被重新打包	投手转实验与组合经理	新岗位卡与 KPI
E 退出	低价值、重复或无使用者	无人阅读的周报	停止条件与节省再分配

4.1 六维迁移评分卡

维度	1 分	3 分	5 分	决策影响
经营价值	无明确使用者	支持局部指标	直接影响收入/品牌/客户	高价值不等于自动化
判断密度	规则明确	需少量取舍	需语境、伦理与权衡	高分偏保留/增强
可评测性	无法客观验收	可人工抽检	可自动或对照验证	高分利于执行
数据条件	缺失或无权使用	部分可用	完整、及时、授权明确	低分先治理数据
风险与可逆性	高损失难恢复	可局部修复	低风险易回滚	低风险才下放
交互复杂度	多方敏感协商	有限交接	单一系统规则任务	复杂任务偏协同

评分不应机械求和。任何涉及法定责任、重大客户承诺、不可逆预算、敏感人群或公开危机的任务，即使可评测性很高，也应设置否决门。Simon 的有限理性说明决策者不可能掌握全部信息；AI 可以扩大搜索与方案空间，却不能消除目标冲突和责任承担。

第五章 判断权阶梯与风险治理

竹势 AI 营销智库出品

AI 建议

人独立判断

AI 草拟

人审后使用

共同决策

证据与反证并列

限权执行

额度 · 白名单 · 可停机

禁止机器决定

不可逆 · 高责任 · 弱势群体

判断权阶梯把“人机协同”转换为可审计授权。层级包括 AI 建议、AI 草拟、人类共同决策、限权机器执行、禁止机器决定。Parasuraman 等关于自动化层级的研究提示，自动化可以分布在信息获取、分析、决策选择和行动执行的不同阶段。企业不应把生成能力直接等同于行动权。

层级	机器可以做什么	人必须做什么	典型场景
AI 建议	汇总证据、提出选项	决定是否进入讨论	品牌定位、危机预警
AI 草拟	形成可编辑初稿	核验事实、语气与责任	内容、回复、合同外沟通
共同决策	提供模拟与反证	共同选择并记录理由	媒介组合、实验设计
限权执行	在额度和白名单内行动	设定约束、抽检、处理异常	暂停异常广告、发送低风险提醒
禁止机器决定	不得作出或发出决定	人类专属判断与签批	重大承诺、敏感权益、裁员与歧视性决策

风险信号	治理动作
不可逆或难补救	提高人工层级，要求双人签批
反馈延迟	限制自动执行规模，先做小样本
目标可被代理指标替代	加入对照、增量或质量指标
数据含敏感个人信息	最小化访问、用途限制与留痕
供应商自报效果	标注披露性质，要求独立验证
公开传播与危机场景	预设停机、升级与统一发言人

BOARD BRIEF 08

第六章 九岗任务迁移图谱（上）

竹势 AI 营销智库出品

品牌

主张与危机立场保留

内容

证据供应链增强

创意

候选扩展、概念终审

媒介

组合与增量重组

投放

限额执行、实验负责

九岗图谱不使用统一模板替代专业差异。品牌岗位面对长期意义与一致性；内容岗位面对证据、叙事和资产供应链；创意岗位面对概念、审美与文化语境；媒介岗位面对触点组合与平台结构；投放岗位面对实验和边际预算。

6.1 品牌岗位

品牌岗位的重构重点不是减少一个固定比例的人，而是把低判断、可回滚动作迁出，把核心主张、价值观、危机立场、品牌架构等责任显性化。岗位负责人需要对 AI 产出的选择、发布和后果承担可追踪责任。

消退任务	增强任务	新增任务/角色	不可下放判断	关键接口
手工汇总竞品口号、重复检查格式	多方案定位探索、品牌资产一致性审查	品牌语义治理、AI资产授权管理员	核心主张、价值观、危机立场、品牌架构	产品、法务、创意、销售、管理层

任务编号	代表任务	ROLE-6	迁移条件
品牌-1	手工汇总竞品口号	Execute	竞品来源可追溯、分类规则稳定，输出仅作内部扫描且允许人工抽查
品牌-2	多方案定位探索	Orchestrate	必须同时输入品牌战略、目标受众与真实产品证据，候选方案不得自动发布
品牌-3	品牌语义治理	Rebundle	当品牌资产跨渠道复用频繁、版本冲突可量化时，设统一语义规则与资产责任人
品牌-4	核心主张	Retain	涉及品牌定位、价值观或公共立场时，由品牌最高责任人基于长期取舍终审
品牌-5	产品	Link	产品事实、上市节奏或承诺变化须由产品负责人确认后写入品牌资产库

6.2 内容岗位

内容岗位的重构重点不是减少一个固定比例的人，而是把低判断、可回滚动作迁出，把事实、观点边界、署名责任、敏感表达等责任显性化。岗位负责人需要对 AI 产出的选择、发布和后果承担可追踪责任。

消退任务	增强任务	新增任务/角色	不可下放判断	关键接口
从零起草、平台机械改写、标签整理	选题组合、证据增强、叙事编辑、资产复用	内容供应链经理、事实核验编辑	事实、观点边界、署名责任、敏感表达	品牌、研究、SEO、社媒、法务

任务编号	代表任务	ROLE-6	迁移条件
内容-1	从零起草	Execute	资料有明确来源、模板固定且发布前仍经过事实与品牌审核
内容-2	选题组合	Orchestrate	选题必须关联经营目标、受众动作与内容预算，系统只生成有证据的候选组合
内容-3	内容供应链经理	Rebundle	当多渠道版本、审核队列和资产复用形成稳定流水线时，设内容供应链责任
内容-4	事实	Retain	关键事实、署名观点、比较性主张和监管敏感表达必须由编辑责任人确认
内容-5	品牌	Link	品牌口径变更、研究证据更新或法务意见必须通过版本化接口同步

6.3 创意岗位

创意岗位的重构重点不是减少一个固定比例的人，而是把低判断、可回滚动作迁出，把创意命题、审美取舍、文化语境、最终提案等责任显性化。岗位负责人需要对 AI 产出的选择、发布和后果承担可追踪责任。

消退任务	增强任务	新增任务/角色	不可下放判断	关键接口
尺寸延展、初级草图、素材检索	概念发散、快速原型、变体测试	创意系统设计师、生成资产制片	创意命题、审美取舍、文化语境、最终提案	品牌、内容、制作、媒介、采购

任务编号	代表任务	ROLE-6	迁移条件
创意-1	尺寸延展	Execute	仅限既定母版、授权素材和固定规格，生成结果可回滚且不改变核心概念
创意-2	概念发散	Orchestrate	创意 brief、受众张力与品牌禁区齐备，系统用于扩展方案空间而非替代评审
创意-3	创意系统设计师	Rebundle	当生成资产量、版权核验和多版本制作成为持续负荷时，集中为创意系统职能
创意-4	创意命题	Retain	创意命题、文化语境、审美取舍及品牌主视觉由创意负责人终审
创意-5	品牌	Link	品牌、制作与采购需共同确认授权范围、制作可行性和供应商责任

6.4 媒介岗位

媒介岗位的重构重点不是减少一个固定比例的人，而是把低判断、可回滚动作迁出，把目标受众、渠道角色、品牌安全和长期触达逻辑等责任显性化。岗位负责人需要对 AI 产出的选择、发布和后果承担可追踪责任。

消退任务	增强任务	新增任务/角色	不可下放判断	关键接口
报表拼接、排期整理、基础受众重叠分析	跨渠道情景模拟、触点组合、库存风险判断	媒介组合架构师、平台审计负责人	目标受众、渠道角色、品牌安全和长期触达逻辑	投放、财务、数据、代理商、平台

任务编号	代表任务	ROLE-6	迁移条件
媒介-1	报表拼接	Execute	平台数据口径已统一、重复转化已标记，自动报表不直接触发预算变化
媒介-2	跨渠道情景模拟	Orchestrate	渠道角色、预算边界和增量假设已定义，模拟结果必须由实验或外部证据校准
媒介-3	媒介组合架构师	Rebundle	当跨平台组合、合同退出和品牌安全成为常态决策时，设组合架构责任
媒介-4	目标受众	Retain	目标受众、渠道功能、频次上限和长期覆盖逻辑由媒介负责人决定
媒介-5	投放	Link	投放、财务、数据和代理商须按同一归因口径、结算周期与异常 SLA 交换信息

6.5 投放岗位

投放岗位的重构重点不是减少一个固定比例的人，而是把低判断、可回滚动作迁出，把预算上限、归因口径、加码停机、异常处置等责任显性化。岗位负责人需要对 AI 产出的选择、发布和后果承担可追踪责任。

消退任务	增强任务	新增任务/角色	不可下放判断	关键接口
手工调价、常规账户巡检、低风险规则执行	实验设计、创意学习、边际预算配置	增长实验经理、算法监督员	预算上限、归因口径、加码停机、异常处置	媒介、创意、销售、财务、数据

任务编号	代表任务	ROLE-6	迁移条件
投放-1	手工调价	Execute	仅在白名单账户和小额额度内执行，设置日损失上限、回滚与异常停机
投放-2	实验设计	Orchestrate	实验需预先定义假设、对照、样本窗口和成功阈值，避免平台事后归因
投放-3	增长实验经理	Rebundle	当预算组合、实验队列和算法监督成为持续职责时，重组为增长实验管理
投放-4	预算上限	Retain	预算上限、跨账户迁移、重大加码停机及敏感受众策略由业务责任人签批
投放-5	媒介	Link	媒介、创意、销售和财务必须回传库存、素材、线索质量与边际利润信号

第七章 九岗任务迁移图谱（下）

竹势 AI 营销智库出品

社媒

公开承诺与危机升级

CRM

同意、频控与权益边界

研究

问题与因果解释保留

管理

资源、绩效与最终问责

7.1 社媒岗位

社媒岗位优先迁移的是可按语义标签稳定分流、可在发布前预览的监听与草拟任务；迁移上限由公开可见性、情绪强度和传播速度共同决定。日常评论可由系统分类并给出回复候选，但涉及赔偿承诺、群体冲突、公共议题或快速扩散事件时，必须切换到统一事实源和人工发言权。社媒负责人要与客服约定工单移交时限，与品牌和公关约定危机口径，与法务约定禁答事项，并对自动回复身份、暂停条件和升级记录负责。

消退任务	增强任务	新增任务/角色	不可下放判断	关键接口
评论分类、日常摘要、标准问答草拟	社区语义理解、议题节奏、关系维护	社区风险指挥、社媒知识运营	公开承诺、冲突回应、危机升级、弱势群体沟通	客服、品牌、法务、公关、产品

任务类型	机器作用	人工责任	升级条件
低风险监听与分流	聚类评论、识别重复问题并生成待审回复	抽查语义标签，维护官方事实库和禁答清单	负面情绪骤升、传播速度异常或出现人身/公共安全风险
关系与争议处置	汇总对话历史、利益相关者和可选回应	判断公开立场、承诺尺度与回应时机	媒体介入、跨平台扩散、重大客户或群体冲突
社区治理建设	监测规则失效、沉淀案例并建议升级树	维护发言权、暂停机制、危机演练与交接 SLA	自动回复越权、同类误判反复或监管要求变化
不可下放决定	提供证据包和回应草案，不直接发布	对赔偿、公共议题、危机声明和弱势群体沟通负责	品牌/公关最高责任人或法务联合签批

7.2 CRM岗位

CRM 岗位的迁移顺序应由客户授权、旅程阶段和动作可逆性决定。名单清洗、偏好归并和低风险触达草稿可以先行自动化；价格例外、退款权益、敏感属性推断和高价值客户挽留不得由模型单独决定。CRM 负责人需要把同意范围、频控、退出渠道和人工接管写入旅程规则，并与销售明确线索转交阈值、与客服明确承诺回写、与数据团队明确字段用途和保留期限。

消退任务	增强任务	新增任务/角色	不可下放判断	关键接口
名单清洗、分群标签、常规触达草稿	旅程编排、下一最佳行动、流失风险解释	客户旅程编排师、同意与偏好管理员	权益承诺、价格例外、敏感触达、人工接管	销售、客服、数据、法务、产品

任务类型	机器作用	人工责任	升级条件
数据整理与低风险触达	去重客户记录、归并偏好并草拟授权范围内的信息	核对同意状态、频控和字段用途，抽查分群偏差	退订上升、投诉集中、同意记录缺失或数据来源不明
旅程与挽留决策	预测流失、建议下一最佳行动并模拟触达节奏	结合客户价值、关系历史和服务问题决定动作	涉及价格歧视风险、脆弱客户、重大退款或高价值客户
同意与偏好治理	发现冲突字段、过期授权和触达过载	维护用途清单、删除请求、人工接管和跨渠道频控	敏感属性被推断、授权用途扩张或跨境/第三方共享
不可下放决定	形成客户背景与选项，不自动承诺权益	批准价格例外、赔偿、敏感触达和终止关系	销售/客服业务负责人按额度签批，超额升级法务或管理层

7.3 研究岗位

研究岗位适合迁移的是来源发现、转录、编码建议和证据整理，不适合迁移的是研究问题、样本适用性和因果结论。自动处理只有在来源可追溯、重复材料可去重、关键数字能回到时才进入正式证据链；合成样本只能用于探索，不得替代真实消费者证据。研究负责人要为每项结论标明证据等级、时间边界和可推翻条件，并与产品、品牌和数据团队约定决策问题、数据口径与复核责任。

消退任务	增强任务	新增任务/角色	不可下放判断	关键接口
检索、转录、初步编码、资料格式化	研究设计辅助、证据地图、反证搜索、情景推演	证据审计师、合成样本治理负责人	问题定义、样本与方法选择、因果解释、结论边界	战略、产品、品牌、数据、外部机构

任务类型	机器作用	人工责任	升级条件
资料处理与证据编目	检索、转录、初步编码并标注出处	验证原始来源、去除重复转载、维护证据等级	关键数字无法回、来源相互引用或时间口径不一致
研究设计与解释	提出假设、问卷草案、反证路径和分析方案	选择样本、方法、因果识别与外推边界	样本偏差无法校正、相关性被误作因果或结论影响重大投资
合成数据与自动研究治理	记录生成条件、比较真实样本并提示不确定性	限定用途、设计验证集、审计模型偏差与可重复性	合成样本替代真实调查、模型版本变化导致结论漂移
不可下放决定	整理支持与反对证据，不宣布最终结论	定义问题、接受或否定假设并承担研究结论	研究负责人签署；高风险决策需业务、法务或专业机构复核

7.4 管理岗位

管理岗位可以把派单、进度汇总和例会材料交给系统，但目标冲突、资源取舍、绩效后果和组织承诺不能外包。迁移边界由决策影响范围、员工权益和预算不可逆性确定：系统可提示容量缺口和异常，管理者必须解释取舍、配置资源并承担结果。市场负责人还要与 CHRO 约定岗位变化与申诉机制，与 CIO 和法务约定权限及审计，与财务和业务部门约定投入上限、收益口径和停机条件。

消退任务	增强任务	新增任务/角色	不可下放判断	关键接口
派单催办、周报拼接、例会纪要	目标分解、系统编排、容量管理、例外处置	人机运营负责人、AI 治理委员会秘书	资源配置、绩效、授权、组织承诺与最终问责	CEO、CHRO、CIO、法务、财务、业务部门

任务类型	机器作用	人工责任	升级条件
运营记录与容量预警	汇总任务、依赖、队列和进度偏差	校验承诺、调整优先级并处理信息缺口	关键路径延误、审核队列超载或跨部门SLA 连续失守
目标与资源编排	模拟不同目标、人员和预算组合的后果	处理目标冲突、配置容量并解释取舍	涉及人员调整、重大预算、客户承诺或不可逆资源锁定
人机治理与例外机制	记录权限使用、异常和规则命中情况	设授权额度、申诉渠道、停机方案和责任链	员工权益受影响、监控超出必要范围或重大模型/数据事故
不可下放决定	提供方案、预测和风险提示，不形成组织承诺	批准绩效后果、岗位变化、资源配置与最终问责	CEO/CMO/CHRO 等相应权责主体依治理章程签批

品牌岗位的典型决策场景

品牌岗位最容易被误判为“主要做文案和视觉规范”。事实上，品牌经理的核心工作是维护企业对市场的长期承诺：选择一个主张意味着放弃其他主张，允许一次促销表达意味着判断它是否侵蚀价格体系，批准一个生成式视觉意味着判断其是否延续品牌的识别、文化与伦理边界。AI 可以同时提供数十种定位表达、竞品语义地图和资产一致性检查，但方案空间扩大后，取舍负担更重。品牌负责人应建立“主张—证据—受众—渠道—禁区”的决策记录，使每次选择可解释、可复盘。对于新品、小范围渠道和短周期内容，可以授权系统进行草拟和预检；涉及品牌架构、核心价值、公共议题和危机立场时，应由明确的品牌责任人终审。品牌团队还需与产品部门约定事实有效期，与销售约定不可夸大的承诺，与法务约定高风险表达的审批等级。其绩效不应以生成资产数量衡量，而应观察品牌一致性、主张被业务采用的程度、重大偏差和学习周期。

内容岗位的典型决策场景

内容岗位会出现最显著的“事务减少、责任增加”。检索、转录、提纲、初稿、摘要、标题变体、渠道适配和元数据整理都可被显著增强；但内容团队必须承担证据来源、事实时效、观点边界、作者责任和资产生命周期。企业应把内容生产改造成受控供应链：选题进入前先说明经营目标和受众动作，初稿必须绑定证据，编辑环节区分事实修改、观点修改和品牌修改，发布后将表现、投诉和销售反馈回写。AI 生成的低成本会诱发“多发总比少发好”，因此必须设置内容预算、候选上限和退出规则。长篇研究、专家署名、监管敏感行业、比较性广告和客户案例需要更高层级审核。内容岗位的未来并非单纯“会写得更快”，而是成为证据编辑、叙事架构师和资产运营者，能够决定哪些内容值得生产、哪些内容应合并、哪些内容必须停止。

创意岗位的典型决策场景

创意岗位的可迁移任务集中在素材搜索、草图、构图变体、尺寸延展、简单修图和快速原型。不可替代的部分是把商业问题转化为创意命题，识别文化语境，判断概念是否真正有差异，决定哪种表达值得承担品牌风险。生成式工具会把“做一个版本”变成“做一百个版本”，但审美与策略并不会按同样速度扩张。创意负责人应限制进入评审的候选数量，并要求每个候选说明其要改变的认知、情绪或行为。版权、肖像、训练数据和风格模仿风险应在制作前处理，而非上线前最后一轮补救。对于常规电商图和可量化广告素材，可以采用机器变体与快速测试；对于品牌主视觉、社会议题和长期资产，应保持人类主导的概念与终审。创意绩效应同时看商业效果、原创性、品牌适配、制作可行性和投放后学习，不能只以点击率或制作速度评价。

媒介岗位的典型决策场景

媒介岗位的价值将从采购和排期操作转向触点组合、平台权力管理和增量覆盖。算法可以预测受众、自动分配库存和优化竞价，但平台通常围绕自身可观测目标运行，企业则需要同时考虑品牌长期记忆、跨渠道重叠、价格、内容环境和客户体验。媒介负责人应把平台建议视为输入，而不是结论；建立渠道角色定义，明确哪些渠道负责认知、考虑、转化或关系维护；对集中度、品牌安全、频次和数据依赖设上限。跨平台归因容易重复计算同一转化，因此应使用实验、地理分组、时间序列或其他适当方法验证增量。媒介岗位还需与采购管理合同退出，与数据团队管理口径，与创意团队建立素材学习回路。真正的升级不是少做报表，而是能够解释为什么投入这个触点、为什么停止另一个触点，以及组合在不确定条件下如何保持韧性。

投放岗位的典型决策场景

投放岗位的手工操作会快速减少，但预算责任不会消失。系统可以提出出价、预算、定向和素材调整，甚至在限额内执行；人必须设计目标函数、识别代理指标、决定实验结构、控制不可逆损失并承担加码或停机责任。一个高转化率的自动策略可能只是在收割已有需求，或通过优惠吸引低价值客户；一个较低短期回报的策略可能在建立新客和品牌记忆。投放经理应从账户维护者升级为增长实验经理：维护假设池，分配探索与利用预算，设置对照，区分平台归因与企业增量，记录失败原因。机器执行权应采用额度和白名单，预算跨账户迁移、重大出价变化、敏感受众和品牌安全事件必须升级。绩效应看边际增量、预算偏差、实验有效率、异常损失和学习复用，而不是单一 ROAS。

社媒岗位的典型决策场景

社媒场景具有高频、公开、情绪化和不可完全撤回的特点。AI 适合做舆情聚类、评论分类、常见问题草拟、对话摘要和低风险回复建议；但公开承诺、冲突回应、重大投诉、社会议题、未成年人和弱势群体沟通必须由人处理。社媒负责人应建立语义边界与升级树，不仅按关键词拦截，还要按意图、情绪、传播速度和潜在伤害分级。自动回复需要明确身份，避免让客户误以为已获得正式承诺。危机期间应降低或暂停自动化，统一事实源和发言权。社媒团队还需把社区反馈结构化回传产品、客服和品牌，而不是只交付互动量。其未来岗位更接近社区关系经理和风险指挥：日常由系统提高响应覆盖，关键时刻由人维护信任、同理心和组织立场。

CRM 岗位的典型决策场景

CRM 与会员运营可以使用 AI 完成标签、分群、下一最佳行动建议、旅程草拟、流失预警和触达内容适配，但数据可用并不等于可以任意使用。客户同意、频率偏好、敏感属性、权益承诺和公平性必须进入任务规则。系统建议的高转化触达可能造成过度打扰、价格歧视或对脆弱客户施压。CRM 负责人应建立“用途—数据—动作—额度—退出”的闭环，对不同渠道和客户阶段设频控，对退款、价格例外、金融医疗信息和高价值客户设置人工接管。与销售的接口要明确何时从营销培育转为人工跟进，与客服明确问题和承诺如何回写，与法务明确同意和删除请求。岗位绩效应看客户终身价值、留存、投诉、退订、触达公平和旅程学习，而不是发送量。

研究岗位的典型决策场景

研究岗位的输入处理会被深度增强：检索、翻译、转录、编码、摘要、竞品对比和初步假设生成都能加速。风险也随之增加，因为模型会把不完整材料写成连贯叙事，把相关关系包装成因果，把重复转载误判为多源证据。研究负责人应建立证据等级、来源独立性、时间边界、样本适用性和反证要求。合成受访者或模拟人群可用于探索问题和测试问卷，但不能在缺少验证时替代真实消费者。任何关键结论应说明什么证据可能推翻它。研究岗位未来更像问题架构师和证据审计师：定义决策需要什么信息，选择方法，辨别样本与测量偏差，解释不确定性，并把结论转成可行动而不过度承诺的建议。绩效可看决策采用、预测校准、证据缺陷率和研究资产复用。

管理岗位的典型决策场景

市场管理者过去大量时间用于分派、催办、汇总和会议；这些事务可以被系统显著减少。新的管理负担是设定目标、配置人机容量、设计权限、处理例外、解决跨部门瓶颈和承担最终结果。管理者不能把机器生成的计划直接视为承诺，也不能在资源不变的情况下要求团队审核无限增加的产出。应建立固定运营节奏：每周看任务队列和异常，每月看经营结果、质量成本和能力迁移，每季度审查岗位卡与供应商。员工监控必须遵守必要性和比例原则，避免把日志变成无边界绩效监控。管理岗位的核心绩效是组合结果、系统韧性、决策速度、人才成长和重大风险，而不是“使用 AI 的次数”。当系统失败时，管理者仍是资源、授权和结果责任的最终承担者。

第八章 新角色、共享能力与组织接口

竹势 AI 营销智库出品

销售

线索定义 · 时限 · 退回

客服

承诺 · 升级 · 人工接管

产品

事实版本 · 反馈闭环

数据 / IT

口径 · 权限 · 停机

法务 / 采购

禁区 · SLA · 退出

新角色通常出现在旧岗位之间的空白地带：谁维护品牌上下文，谁管理工作流状态，谁验证自动化，谁处理权限与模型风险，谁把学习写回知识资产。社会技术系统思想强调，技术与社会组织必须共同优化；只增加工具而不调整权责，会形成“系统有能力、组织无责任”的断层。

新型责任	适合专职的条件	小团队做法	主要产出
AI 内容运营	内容量大、渠道多、品牌风险高	由内容负责人兼任并设共享规范	资产流水线、审核规则
工作流编排	跨系统任务多且稳定复用	市场运营兼任	流程图、异常与日志
知识资产治理	资料量大、更新频繁	研究或品牌共享负责	知识目录、有效期与引用规则
模型与自动化风险	强监管、客户影响大	与法务/IT 联合委员会	风险登记、测试与停机
增长实验管理	预算大、渠道多、增量难测	投放负责人升级	实验组合与证据库

8.1 市场与八类组织接口

接口	交换对象	必须明确的契约	异常所有者
销售	MQL、SQL、商机反馈	定义、时限、退回原因	收入运营负责人
客服	问题、情绪、承诺与升级	机器人边界、人工接管	客服负责人
产品	卖点、路线、反馈	版本与事实有效期	产品市场负责人
数据	指标、口径、权限	数据血缘、刷新频率	数据负责人
IT	系统、账号、日志	接入、变更、停机	系统所有者
法务	内容、隐私、合同	审核等级、SLA、禁区	法务责任人
采购	供应商、成本、退出	数据归属、锁定、审计	采购负责人
代理商	Brief、资产、效果	交付格式、证据和复盘	代理商管理者

Conway 定律表明，系统结构会反映组织沟通结构。企业若希望形成端到端 AI 营销链路，不能只采购一个“全能平台”，还必须重画跨部门沟通和决策路径。Hammer 的流程再造思想在此仍有价值：围绕客户结果重组流程，而不是在旧部门边界上加速局部步骤。

BOARD BRIEF 11

第九章 产能上升后的二阶瓶颈

竹势 AI 营销智库出品



AI 最容易制造的管理错觉是“供给越多越好”。当草稿、素材、受众组合和优化建议成倍增加，审核、选择、存储、版本、上线和学习回收会成为新瓶颈。若没有容量上限，团队会把更多时间花在筛选和解释上，甚至降低决策质量。

二阶问题	形成机制	领先指标	治理动作
审核拥堵	草稿量超过审批容量	待审队列、平均等待	限定进入审核的候选数
返工上升	约束未前移、标准冲突	返工轮次、否决原因	结构化 Brief 与自动预检
资产混乱	多版本无唯一来源	重复率、过期资产	资产 ID、状态和有效期
决策瓶颈	领导成为唯一终审	升级集中度、超时	风险分级与授权梯度
学习失真	只保留成功、归因不清	实验覆盖与对照率	证据库和失败复盘
风险外溢	机器跨系统执行	越权、投诉、回滚	最小权限与停机机制

Goodhart 定律提醒管理者，当指标成为目标后，它会失去作为指标的价值。内容数量、节省工时和采用率一旦直接绑定绩效，员工与系统都会倾向于生产更多可计数活动，而不是更高的客户价值。

BOARD BRIEF 12

第十章 KPI、绩效与经营归因

竹势 AI 营销智库出品

PERSON

个人判断

事实 · 异常 · 实验质量

TEAM

团队闭环

转化 · 品牌 · 周期

SYSTEM

系统可靠

回滚 · 权限 · 单位成本

BUSINESS

经营结果

增量收入 · 留存 · 现金效率

岗位重构后的绩效应分为个人、团队和系统三层。个人层衡量判断、交付和风险处理；团队层衡量端到端结果与协同；系统层衡量稳定性、质量成本和治理。任何一层都不能单独替代经营结果。

层级	指标类别	示例	禁止替代
个人	判断与责任	事实错误率、异常处置、实验质量	内容数量
团队	流程与结果	线索转化、品牌质量、上市周期	采用率
系统	可靠性与风险	成功率、回滚率、权限事件、单位成本	节省工时
经营	收入与客户	增量收入、留存、品牌指标、现金效率	平台归因报表

岗位	首要经营指标	质量/风险指标	学习指标
品牌	品牌资产健康与战略一致性	重大偏差和危机事件	主张验证周期
内容	有效触达与转化贡献	事实/合规/品牌错误	内容假设命中率
创意	创意对业务目标的增量贡献	版权与品牌风险	概念到学习周期
媒介	增量覆盖与组合效率	品牌安全与集中风险	渠道实验覆盖
投放	边际增量回报	预算偏差与异常损失	实验有效率
社媒	关系质量与问题闭环	投诉、公开失误	议题学习速度
CRM	客户价值与留存	不当触达、同意违规	旅程优化周期
研究	决策采用与预测校准	样本与证据缺陷	反证和复用率
管理	组合结果与组织能力	重大失控与人才流失	迁移成熟度

第十一章 人员转型与能力邻接

竹势 AI 营销智库出品

可在岗迁移

核心判断相同

项目训练

新增跨系统能力

专业认证

法律 · 隐私 · 安全 · 统计

不宜强迁移

尊重选择并配置替代路径

能力邻接图从员工已有能力出发，判断未来任务与现有能力之间的距离。邻接不是性格标签，而是可通过作品和任务证据验证的迁移关系。Argyris 的双环学习强调，不仅要改进动作，还要检验支配动作的目标和假设；刻意练习则要求在明确标准、即时反馈和重复挑战中提升。

迁移类型	判定	发展方式	例子
可在岗迁移	核心判断相同，工具与流程变化	导师制+真实任务	资深编辑转内容供应链经理
需项目训练	新增跨系统与数据能力	8-12 周项目认证	投手转增长实验经理
需专业认证	涉及法律、隐私、安全或统计	外部课程+考试+监督实践	模型风险或隐私治理
不宜强迁移	价值观、兴趣或基础能力不匹配	尊重选择、内部转岗或外部招聘	不把优秀文案强制转数据工程

员工转型证据	最低要求
任务作品	至少完成一个真实业务闭环
判断记录	说明选择、否决与风险理由
异常处理	成功处理至少一种边界情形
协作接口	完成跨部门交付并获得验收
经营结果	指标改善或明确的学习结论
复盘能力	能把经验沉淀为规则或资产

员工沟通必须明确：任务迁移不等于预设裁员；试点数据不直接用于个体惩罚；哪些能力将被投资；何时可能涉及岗位合并、外部招聘或自然流失。CHRO 应参与设计，而不是在技术上线后负责安抚。

第十二章 案例卷宗与专家视角

竹势 AI 营销智库出品

八个案例证明八种机制，不汇总成一个“AI ROI”

企业自报、随机现场实验与受限 / 反噬案例按不同证据等级使用；每个案例都单列动作、机制、结果与外推边界。

案例1 | 阿里巴巴/天猫，2025—2026：商家 AI 工具进入营销与经营链路

阿里巴巴官方材料披露，天猫在商家侧升级生成内容、营销出价和经营辅助能力；2025 年官方材料称自动出价模型平均提升活动 ROI 约 12%，2025 年“双 11”相关披露又称 AI 出价引擎带来 12% 的营销 ROI 提升。

案例要素	内容
主体与时间	阿里巴巴/天猫，2025—2026
业务情境	阿里巴巴官方材料披露，天猫在商家侧升级生成内容、营销出价和经营辅助能力；2025 年官方材料称自动出价模型平均提升活动 ROI 约 12%，2025 年“双 11”相关披露又称 AI 出价引擎带来 12% 的营销 ROI 提升。
具体动作	天猫把商家内容生成、经营机会识别、自动出价和客服辅助嵌入商家经营产品，并以自动出价模型在活动预算范围内执行竞价调整。
作用机制	系统把分散在素材制作、市场信号读取和竞价操作中的高频决策压缩到同一数据回路；实时转化信号缩短调价延迟，规模化素材供给扩大可测试组合，从而可能改善预算匹配效率。
结果证据	结果为企业官方披露，能够证明平台采用和自报效果，但不能独立证明所有商家获得同等增量，也不能排除同期流量、促销和模型取舍影响。
边界与启示	平台目标函数、商家利润与品牌长期价值可能不完全一致；企业仍需保留预算、品牌和实验责任。管理含义：投手岗位应转向约束设计、增量实验与平台审计。

案例2 | 腾讯，2025—2026：AI 广告推荐与营销服务增长

腾讯 2025 年年度业绩披露，营销服务收入同比增长 19% 至人民币 1450 亿元，并称 AI 能力改善广告定向；2026 年第一季度营销服务收入同比增长 20%，官方称升级 AI 驱动广告推荐并扩展微信生态闭环营销能力。

案例要素	内容
主体与时间	腾讯，2025—2026
业务情境	腾讯 2025 年年度业绩披露，营销服务收入同比增长 19% 至人民币 1450 亿元，并称 AI 能力改善广告定向；2026 年第一季度营销服务收入同比增长 20%，官方称升级 AI 驱动广告推荐并扩展微信生态闭环营销能力。
具体动作	腾讯升级广告推荐、定向与转化链路，在微信生态内连接内容触点、用户行为和商业转化信号，并扩展闭环营销产品。
作用机制	更连续的一方行为信号提高受众与广告匹配的可观测性，推荐系统据此改善排序和库存定价；闭环链路减少跨平台信号丢失，使模型更快获得反馈。
结果证据	这是上市公司经营披露，能支持“AI 与广告平台能力强化同时发生”，但不能把全部收入增长单独归因于 AI，广告负载、用户参与和行业需求亦有贡献。
边界与启示	平台收入增长不等于广告主增量利润；企业要建立自己的增量测量与品牌安全指标。管理含义：媒介与投放岗位不能只接受平台报表。

案例3 | 中国平安，2025：AI 服务代表承担高频客户服务

平安 2025 年中期业绩披露，2025 年上半年 AI 服务代表提供约 8.82 亿次服务，占客户服务总量约 80%。

案例要素	内容
主体与时间	中国平安，2025
业务情境	平安 2025 年中期业绩披露，2025 年上半年 AI 服务代表提供约 8.82 亿次服务，占客户服务总量约 80%。
具体动作	平安将大量标准查询、流程指引和常规服务交由 AI 服务代表处理，同时保留人工坐席承接复杂咨询、投诉、情绪沟通和合规例外。
作用机制	高频问题被标准化后，系统可并行响应并保持知识口径一致，释放人工容量处理需要同理心、授权和跨系统判断的长尾场景；分层路由决定效率能否转化为服务质量。
结果证据	数据来自公司披露，证明大规模采用与服务占比；未提供独立客户满意度对照，不能据此推断人员裁减或所有场景质量提高。
边界与启示	金融服务包含敏感权益与合规义务，客户承诺、投诉、弱势群体和重大例外应人工接管。管理含义：CRM 与社媒岗位应明确语义、额度和升级边界。

案例4 | 京东，2019—2026：智能客服与商家经营 AI

京东官方曾披露在 618 期间智能客服处理 90% 的咨询并响应超过 3.8 亿次；2026 年又宣布在 618 中让 AI 参与经营诊断、营销投放、直播转化、客服和售后，面向超过一百万商家。

案例要素	内容
主体与时间	京东，2019—2026
业务情境	京东官方曾披露在 618 期间智能客服处理 90% 的咨询并响应超过 3.8 亿次；2026 年又宣布在 618 中让 AI 参与经营诊断、营销投放、直播转化、客服和售后，面向超过一百万商家。
具体动作	京东在客服链路部署意图预测、情绪识别和智能应答，并进一步向商家提供经营诊断、营销投放、直播转化、客服与售后等 AI 工具。
作用机制	意图预测把咨询提前路由到合适知识或人工队列，减少等待和重复转接；商家工具则把前端需求信号回流到营销与服务动作，形成跨触点的经营反馈回路。
结果证据	早期服务比例和响应量来自公司官方博客，属于企业自报；跨年度材料口径不同，不能合并为趋势，也不能推断净就业影响。
边界与启示	意图预测错误、情绪误判和复杂售后仍需人工；商家端采用效果存在行业差异。管理含义：客服岗位应升级为复杂关系管理与异常处理。

案例5 | Harvard Business School / P&G, 2024—2025: The Cybernetic Teammate 现场实验

研究团队与宝洁开展随机现场实验，791 名专业人员参与真实产品创新任务，用于比较个人、团队及 AI 辅助条件下的表现。

案例要素	内容
主体与时间	Harvard Business School / P&G, 2024—2025
业务情境	研究团队与宝洁开展随机现场实验，791 名专业人员参与真实产品创新任务，用于比较个人、团队及 AI 辅助条件下的表现。
具体动作	研究团队将宝洁专业人员随机分配到个人或双人团队、使用或不使用生成式 AI 的条件，让参与者完成真实产品创新任务，并比较不同协作结构。
作用机制	AI 为个人即时补充跨职能知识、生成候选方案并支持迭代，使原本依赖另一专业角色的信息整合部分内嵌到个人 workflow；其价值来自降低专业边界摩擦，而非单纯加快写作。
结果证据	证据来自研究论文与公开研究说明，样本和任务明确；结果适用于特定创新任务和组织环境，不能直接外推到所有营销岗位或长期绩效。
边界与启示	短期任务表现不包含长期品牌责任、组织学习、客户关系与人员发展。管理含义：创意与研究岗位应设计人机团队，而非简单替人。

案例6 | Klarna, 2024—2025：客服自动化与人工能力再平衡

Klarna 2024 年官方披露 AI 助手承担大量客服对话并声称相当于数百名全职人员的工作量；随后公司管理层公开讨论重新增加人工服务，以改善客户体验和质量。

案例要素	内容
主体与时间	Klarna, 2024—2025
业务情境	Klarna 2024 年官方披露 AI 助手承担大量客服对话并声称相当于数百名全职人员的工作量；随后公司管理层公开讨论重新增加人工服务，以改善客户体验和质量。
具体动作	Klarna 先以 AI 助手承接大比例客服会话、压缩人工招聘与服务成本，随后公开调整方向，重新强调人工客服和更高质量的客户接触。
作用机制	自动化在标准问题上扩大并发和降低单位处理成本，但当复杂问题、关系修复和情绪沟通占比上升时，缺少人工判断会把节省转化为体验损失；重新配置人工容量是对质量约束的反馈。
结果证据	最初效果主要为公司与供应商自报；后续人员再平衡说明自动化采用并不等于稳定替代，也不能仅凭早期效率推断长期组织结构。
边界与启示	过度强调处理量可能损害关系质量；高复杂和高价值客户需要人工。管理含义：这是受限/反噬案例，必须设置质量门、客户分层和可随时启用的回退容量。

案例7 | 可口可乐，2024：生成式 AI 节日广告争议

可口可乐在节日传播中使用生成式 AI 制作广告，引发受众对质感、真实性和品牌情感的争议。

案例要素	内容
主体与时间	可口可乐，2024
业务情境	可口可乐在节日传播中使用生成式 AI 制作广告，引发受众对质感、真实性和品牌情感的争议。
具体动作	可口可乐在标志性节日广告中使用生成式 AI 参与画面与制作，并以新技术版本延展长期品牌资产。
作用机制	生成技术降低镜头和版本制作成本，却也让细节一致性、人物真实感和情感工艺更容易被受众放大审视；当作品调用强记忆符号时，微小失真会通过品牌期待差距触发反感。
结果证据	争议能证明品牌反应风险，但公开资料不足以独立量化其对销量或品牌资产的净影响。
边界与启示	效率不能证明品牌增量；高知名度品牌需要更严格的审美与文化终审。管理含义：创意岗位的稀缺价值转向命题、文化判断和最终取舍。

案例8 | Boston Consulting Group / HBS，2023：锯齿状能力边界现场实验

案例要素	内容
主体与时间	Boston Consulting Group / Harvard Business School 等，2023
业务情境	研究团队在波士顿咨询公司开展预注册现场实验，758 名顾问完成一组位于或超出生成式 AI 能力边界的咨询任务，用于观察不同任务类型下的人机协作表现。
具体动作	参与顾问被随机分配到不使用 AI、使用 GPT-4、使用 GPT-4 并接受协作方法说明等条件；组织把真实知识工作拆为创意、分析和判断任务，并记录质量、速度及错误。
作用机制	在模型能力覆盖的任务上，AI 提供结构、草稿和快速迭代，扩大个人可完成的工作范围；在能力边界外，流畅但错误的输出会诱发过度依赖，导致专业人员减少独立验证并放大判断失误。
结果证据	同行评审前研究报告显示，在能力边界内，使用 AI 的顾问完成更多任务、速度更快且质量更高；在边界外任务上，使用 AI 的参与者正确率下降。结果来自特定模型、咨询任务和短期实验，不代表所有岗位或长期绩效。
边界与启示	组织必须识别“锯齿状能力边界”，不能按工具可用性整体迁移岗位。管理含义：研究、内容和管理岗位要由高不确定任务保留独立验证、反证和人工否决权。

12.9 十位具名专家视角

专家	核心视角	岗位重构启示
Peter Drucker	目标与责任	岗位应由结果责任定义，AI 技能只是支撑。
Henry Mintzberg	真实管理工作	碎片化协调和例外处理常被职位说明书遗漏。
Herbert Simon	有限理性	AI 扩大搜索能力，但目标冲突与责任仍需组织解决。
Raja Parasuraman	自动化层级	信息获取、分析、选择与执行可采用不同自动化程度。
W. Edwards Deming	系统观	绩效问题常来自系统而非个体，局部提速会暴露系统缺陷。
Eliyahu Goldratt	约束理论	产能提升会把瓶颈推到审核、数据或决策。
Chris Argyris	双环学习	需检验目标和规则，而不仅优化执行动作。
Michael Hammer	流程再造	围绕客户结果重组端到端流程。
Daniel Kahneman	判断偏差	快速生成会放大锚定、确认和过度自信风险。
Karl Popper	可证伪性	研究与增长实验要设计可能推翻假设的证据。

第十三章 14/30/90 天实施路线与管理工具

竹势 AI 营销智库出品

DAY 01-14

任务盘点

工作日志 · 风险基线 · owner

DAY 15-30

岗位试点

ROLE-6 · 判断权 · 接口契约

DAY 31-90

组织迁移

KPI · 人员路径 · 阶段门 · 退出

13.1 14 天：任务盘点与风险基线

阶段	关键动作	交付物	退出门
第 1-3 天	确定一个业务链路 with 岗位范围	范围说明、结果树	无结果所有者则暂停
第 4-8 天	任务日志、接口和异常观察	任务词典、等待与返工图	样本不覆盖峰值则补充
第 9-12 天	ROLE-6 评分与判断权分级	迁移评分卡、风险登记	高风险无控制则不试点
第 13-14 天	选择试点任务与员工沟通	试点章程、沟通清单	员工与主管未理解则延后

13.2 30 天：岗位链路试点

工作流	要求
选择	一个端到端链路，任务数 5–12 个
基线	速度、质量、返工、风险和经营指标
授权	机器权限、额度、日志、回滚和人工升级
运行	至少覆盖正常、峰值与异常场景
比较	与基线或对照流程比较，不只看自报节省
复盘	决定保留、扩展、重组或退出

13.3 90 天：组织迁移

模块	90 天结果
岗位	发布新岗位卡和判断权
流程	接口契约、SLA、异常所有者
人员	能力邻接、项目认证、招聘或共享能力
治理	风险分级、审计、停机与供应商退出
绩效	个人、团队、系统三层 KPI
资产	知识、内容、实验和版本治理
管理节奏	周运营、月复盘、季度岗位审查

13.4 岗位重构工作包

工具	用途	最低字段
任务清单	描述真实工作	触发、输入、输出、判断、异常、频率
迁移评分卡	决定 ROLE-6	价值、判断、评测、数据、风险、交互
新岗位卡	明确未来责任	结果、任务、权力、接口、KPI
接口契约	跨部门交付	对象、格式、SLA、退回、升级
RACI	明确责任	负责、批准、协商、知会
试点看板	管理证据	基线、运行、质量、风险、结果
人员计划	安排转型	邻接、训练、认证、选择
沟通清单	保护信任	目的、边界、数据用途、申诉、退出

董事会阶段门	必须回答的问题
价值门	是否改善真实经营结果或关键能力？
质量门	错误、返工和审核成本是否可接受？
风险门	权限、客户、品牌、隐私和预算是否受控？
人员门	岗位责任与员工安排是否清晰、公平、可执行？
规模门	扩大规模后瓶颈、成本和治理是否仍成立？
退出门	何时停止、回滚或更换供应商？

专题： 供应商、代理商与外部能力边界

市场部岗位重构常被误解为企业内部问题，实际上大量任务由广告平台、内容平台、代理商、数据供应商和模型供应商共同完成。外部伙伴引入 AI 后，企业必须重新定义交付证据与责任：代理商提交的不应只是成品，还应包括所用事实、版本、审批、素材权利和实验设计；平台提供的自动优化不应免除企业预算责任；模型供应商的安全承诺不能替代企业自己的访问控制和日志。采购合同需明确数据用途、输出权属、人工参与、分包、事故通报、审计、模型变更和退出迁移。对关键流程，应保留供应商替换和人工降级方案，避免岗位知识被黑箱平台吸走。内部岗位也要从“收作业”转向“管理接口与证据”，能够评价外部系统是否真正改善经营结果。

专题： 容量规划与岗位定编的正确顺序

定编应发生在任务迁移试点之后，而非之前。企业首先估算需求容量：内容和活动量、客户互动峰值、投放账户复杂度、研究项目数、审核等级和异常概率；再估算机器容量、人工审核容量和跨部门等待；最后才讨论岗位合并、自然流失、转岗或招聘。自动化降低单项任务工时，并不必然降低总人力需求，因为服务覆盖、测试数量和治理任务可能增加。管理层应至少观察一个完整业务周期和一个峰值周期，区分一次性建设工作与长期运营工作。涉及人员调整时，要把技术变化、业务收缩、组织合并和绩效问题分开说明，不能把所有变化笼统归因于 AI。公平的顺序是先证明流程可行，再明确未来岗位，再提供迁移与选择，最后根据稳定容量作人员决策。

专题：员工参与、申诉与心理安全

岗位重构会改变员工的身份、发展预期和对管理的信任。若企业以秘密试点或“效率竞赛”推进，员工可能隐藏真实流程、避免报告错误或用大量低价值产出证明存在感。项目应在开始时说明收集哪些任务数据、用途是什么、是否用于个体评价、谁可查看以及保存多久；员工应能指出系统错误、拒绝不安全操作并申请人工复核。主管需要奖励发现风险和停止错误自动化的行为，而不是只奖励使用率。对于被合并或显著变化的岗位，应提供具体的能力邻接、学习时间、导师和内部机会。心理安全并不意味着取消绩效责任，而是让员工可以在不受不当惩罚的情况下暴露系统缺陷。只有真实问题进入管理视野，组织才能形成可靠的人机工作系统。

专题：强监管与高声誉风险场景的加严规则

金融、医疗、教育、公共服务、未成年人、上市公司信息披露以及涉及重大公共利益的营销活动，需要比一般消费品更严格的判断权配置。企业应把法规义务、适用人群、禁止主张、证据标准和人工签批写入任务定义；对模型输出进行事实来源校验、适当性检查和歧视风险测试；对客户可见的自动互动明确机器身份和人工渠道。任何可能影响权益、价格、资格、诊断、投资或重大公司信息的动作，不应由营销系统独立决定。高风险场景还需要更长的日志保存、定期抽样、独立审查和事故演练。效率收益若不能在合规、客户理解和可纠正性前提下实现，就不应被视为成功。董事会应接受高风险流程自动化速度较慢，这是治理质量而非落后。

专题：岗位重构的董事会审查问题

董事会不需要逐项审批工具配置，但必须审查岗位重构背后的经营假设和责任安排。第一，重构是否服务于明确的收入、客户、品牌、效率或风险目标，还是只因技术热度启动。第二，试点是否覆盖真实业务、峰值和异常，而非经过挑选的演示任务。第三，企业是否把平台自报、模型基准、员工节省工时和经营增量区分开来。第四，哪些判断权被下放，额度、白名单、日志、回滚和停机是否完整。第五，岗位变化后是否出现无人负责的接口，尤其是营销与销售、客服、产品、数据、IT、法务和代理商之间。第六，员工是否获得真实迁移机会，任务数据是否被合比例地使用。第七，规模扩大后审核、数据、成本和异常容量是否同步增加。第八，若供应商涨价、模型变化、监管要求升级或质量下降，企业是否能退出并维持关键运营。第九，管理层是否对最终结果承担责任，而不是把失败归因于系统。第十，岗位重构是否形成可复用的组织能力，包括任务词典、判断记录、接口契约、知识资产和人才梯队。上述问题应在每季度经营审查中持续回答，因为岗位设计不是一次性组织图调整，而是伴随技术、客户和业务策略变化的动态治理。成熟企业不会追求一次确定所有未来岗位，而是建立一套能够持续观察、试验、纠偏和退出的机制。

专题：年度岗位组合审查机制

岗位重构不应只在组织调整或预算削减时发生。企业可建立年度岗位组合审查，把九类岗位放在同一张经营地图中，检查任务需求是否变化、哪些自动化已稳定、哪些新风险出现、哪些共享能力开始成为瓶颈。审查应使用三类证据：流程证据包括任务量、等待、返工、异常和机器执行日志；经营证据包括客户、收入、品牌、成本和风险结果；人员证据包括能力邻接、转岗表现、关键人才负荷和团队信任。管理层应识别“名义岗位未变但责任已转移”的情况，及时更新岗位卡、授权和薪酬；也要识别“技术可以做但组织不应做”的任务，明确保留人工的理由。岗位组合审查还应考虑业务季节性和战略变化，例如新

品期需要增强研究与创意，增长期需要加强媒介、投放和 CRM，危机期需要扩大品牌、社媒和法务协同。审查结论不应只形成组织图，而应形成未来十二个月的任务迁移清单、人员发展投入、系统建设优先级、供应商调整和风险演练计划。这样，岗位重构才会成为经营能力建设，而非一次性的技术项目。

专题： 重构完成后的最低运行标准

一个岗位重构方案只有进入日常运行，才算完成。最低运行标准包括：每项关键任务有唯一结果所有者；机器权限与人工责任可被员工理解；异常能够在约定时间内升级；内容、预算、客户和研究证据都有版本与来源；系统变更不会绕过岗位负责人；员工可以报告错误和请求人工复核；绩效同时观察经营、质量、风险与学习；供应商退出不会导致核心知识和流程丢失。市场运营负责人应每月抽查任务链，每季度审查判断权和接口，每年重新评估岗位组合。任何无法说明责任人、数据来源、审批状态和回滚方式的自动化，都不应被视为正式生产能力。企业宁可少自动化几个任务，也不要形成无人负责、无法追溯、无法停止的执行链。

最终审查还应确认，任何新增效率都没有通过隐性加班、审核外包或风险转移获得。若内容产能上升但投诉、退订、返工、预算波动或关键人才负荷同步恶化，说明岗位设计仍未完成。组织应把这些负效应纳入质量成本，并在扩张前修复。岗位重构的成功标准，是同等或更低风险下获得更好的经营结果与更强的组织学习能力。

结语

竹势 AI 营销智库出品

先重构工作，再讨论编制

只有当任务迁移、责任接口、经营结果和人员路径都有证据时，岗位合并或新增才是经营决策，而不是技术想象。

AI 时代市场部岗位重构的核心，不是为旧岗位寻找一个技术附属品，也不是按预测比例削减人员，而是把经营结果、任务组合、判断权和组织接口重新连成一套可运行的责任系统。低价值事务应退出，标准任务可自动执行，专业人员应被增强，跨部门工作需要协同，高风险判断必须保留，人和机器之间的新接口则需要重组。企业真正要建设的不是“更多产出”，而是更高质量的选择、更短的学习周期、更明确的责任和更可控的执行。

来源索引（纯文本，可搜索）

[S1] 国际劳工组织（ILO），Pawel Gmyrek 等，《Generative AI and Jobs: A Refined Global Index of Occupational Exposure》，2025-05-20，任务级职业暴露指数。

[S2] 国际劳工组织（ILO），《Generative AI and Jobs: A 2025 Update》，2025-05-20，政策简报。

[S3] 世界经济论坛（WEF），《Future of Jobs Report 2025》，2025-01-07，全球雇主调查与 2030 情景估计。

[S4] 经济合作与发展组织（OECD），《The Effects of Generative AI on Productivity, Innovation and Entrepreneurship》，2025-06，实验研究综述。

[S5] 经济合作与发展组织（OECD），《Generative AI and the SME Workforce》，2025-11，多国中小企业调查。

[S6] 经济合作与发展组织（OECD），《OECD Employment Outlook 2023: Artificial Intelligence and the Labour Market》，2023-07-11，七国工作场所调查与文献综述。

[S7] 美国国家标准与技术研究院（NIST），《Artificial Intelligence Risk Management Framework 1.0》，2023-01，Govern-Map-Measure-Manage。

[S8] 美国国家标准与技术研究院（NIST），《Generative Artificial Intelligence Profile, NIST AI 600-1》，2024-07-26。

[S9] ISO/IEC，《ISO/IEC 42001:2023 Artificial Intelligence Management System》，2023，管理体系标准。

- [S10] 美国劳工部 O*NET, 《O*NET Content Model》, 持续更新, 职业任务与工作活动体系。
- [S11] 欧盟委员会 ESCO, 《European Skills, Competences, Qualifications and Occupations》, 持续更新, 职业与技能分类。
- [S12] David Autor, 《Why Are There Still So Many Jobs?》, Journal of Economic Perspectives, 2015, 任务与技术替代框架。
- [S13] Erik Brynjolfsson, Danielle Li, Lindsey Raymond, 《Generative AI at Work》, 2023/2025, 客服现场研究。
- [S14] Fabrizio Dell’Acqua、Edward McFowland III、Ethan Mollick 等, 《Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of AI on Knowledge Worker Productivity and Quality》, 2023, Boston Consulting Group 758 名顾问预注册现场实验。
- [S15] Harvard Business School / Procter & Gamble, Fabrizio Dell’Acqua 等, 《The Cybernetic Teammate》, 2025, 791 名专业人员现场实验。
- [S16] Raja Parasuraman, Thomas Sheridan, Christopher Wickens, 《A Model for Types and Levels of Human Interaction with Automation》, IEEE, 2000。
- [S17] Peter F. Drucker, 《The Practice of Management》, 1954, 目标与责任。
- [S18] Henry Mintzberg, 《The Nature of Managerial Work》, 1973, 管理工作的真实结构。
- [S19] Herbert A. Simon, 《Administrative Behavior》, 1947, 有限理性。
- [S20] W. Edwards Deming, 《Out of the Crisis》, 1986, 系统与质量管理。
- [S21] Eliyahu M. Goldratt, 《The Goal》, 1984, 约束理论。
- [S22] Chris Argyris, Donald Schön, 《Organizational Learning II》, 1996, 双环学习。
- [S23] Michael Hammer, James Champy, 《Reengineering the Corporation》, 1993, 流程再造。
- [S24] Daniel Kahneman, 《Thinking, Fast and Slow》, 2011, 判断与偏差。
- [S25] Karl Popper, 《The Logic of Scientific Discovery》, 1959, 可证伪性。
- [S26] Melvin Conway, 《How Do Committees Invent?》, 1968, 组织沟通与系统结构。
- [S27] 阿里巴巴集团, 《Tmall Doubles Down on Brand Growth with Expanded Merchant Support》, 2025-03-28, 商家 AI 工具与自动出价自报。
- [S28] 阿里巴巴集团, 《AI Powers Large-scale Applications at the World’s Largest Shopping Festival》, 2025-10-16, 商家 AI 团队与营销 ROI 自报。

[S29] 阿里巴巴集团,《Tmall Deepens Brand Investment and Rolls Out Agentic AI Tools》, 2026-03-27, 品牌与商家经营披露。

[S30] 腾讯控股,《2025 Annual and Fourth Quarter Results》, 2026-03-18, 营销服务收入与 AI 广告能力。

[S31] 腾讯控股,《2026 First Quarter Results》, 2026-05-13, AI 推荐与微信闭环营销。

[S32] 腾讯控股,《2025 Annual Report》, 2026-04-09, 业务与风险披露。

[S33] 中国平安,《2025 Interim Results》, 2025-08-26, AI 服务代表 8.82 亿次及 80% 服务量口径。

[S34] 中国平安,《2025 Annual Report》, 2026-03-26, 集团业务与 AI 应用披露。

[S35] 京东集团官方博客,《JD’s Smart Customer Service Can Predict Consumers’ Intentions》, 2019, 618 客服处理量自报。

[S36] 京东集团官方博客,《JD.com Announces First Fully AI-Integrated 618 Grand Promotion》, 2026-05-19, 商家经营链路。

[S37] 京东集团官方博客,《JD.com Announces Enhanced Spring Dawn Initiative with New AI Service Package》, 2024-03-12, 商家服务自报。

[S38] Klarna,《AI Assistant Handles Two-thirds of Customer Service Chats in its First Month》, 2024-02, 企业自报。

[S39] Klarna 管理层公开访谈与公司材料, 2025, 人工客服能力再平衡。

[S40] The Coca-Cola Company, 节日广告与生成式 AI 相关官方传播材料, 2024。

[S41] 欧盟,《Regulation (EU) 2024/1689 Artificial Intelligence Act》, 2024-07-12, 风险与人类监督要求。

[S42] 中国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》, 2023-07-13, 生成式 AI 服务治理。

[S43] 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》, 2021-08-20。

[S44] 国家市场监督管理总局,《互联网广告管理办法》, 2023-02-25。

[S45] International Organization for Standardization,《ISO 30414 Human Capital Reporting》, 2018, 人力资本指标。

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能力组合；“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。

[shichangbu.ai](#)[专题中心](#) · [返回顶部](#) ↑