



负责任 GEO 倡议实践指南

GEO红皮书（2026）

生成式引擎优化的边界、风险与治理

GEO Red Book 2026: Boundaries, Risks, and Governance

出品：深圳市象新数字科技有限公司

发布日期：2026年7月15日

官网：www.xxdigits.com

摘要 SUMMARY

生成式 AI 技术浪潮席卷而来，随着生成能力与用户规模的迅速跃迁，内容创作、市场营销、品牌传播等领域正面临前所未有的变革。GEO（Generative Engine Optimization生成式引擎优化）作为新兴业态应运而生，野蛮生长的同时也引发了诸多乱象。

作为40余家权威机构联合发布《负责任GEO治理倡议》的落地实践，2026年6月，每经AI智库联合业内专家联合编撰 **《GEO红皮书（2026）：生成式引擎优化的边界、风险与治理——负责任 GEO 倡议实践指南》**。红皮书为合规、负责任的 GEO 厘清边界与目标：让真实信息被正确理解、让用户在更少噪声中做出更好判断、让企业竞争聚焦能力与责任。

以“真实优先、安全治理”为核心理念，红皮书首次系统性梳理来自价值观与商业伦理、信息质量、AI系统攻击三个层面的九大类型风险。同时，红皮书提出合规治理框架与路径，提供从立项诊断到服务商评估的全周期实操指南，向行业各方发出负责任GEO治理倡议，呼吁共建可信的AI信息生态，试图为品牌企业、GEO服务商、AI平台及相关机构法务合规团队建立完整可落地的GEO治理体系，推动GEO回归正向价值。

目录导航 CONTENTS

前言

一 负责任GEO理念与价值观

- GEO 的三个核心目标
- GEO 实践的原则
- GEO 红线与风险识别

二 GEO 红线风险总览

- 风险分级和类型框架速览
- 九大风险类型
- 红线风险详情
- 处置建议与替代方案

三 GEO 风险应对与治理

- 价值观层治理
- 信息质量层治理
- AI系统攻击层治理
- 风险评分矩阵
- 证据保全与责任归属

四 GEO 合规治理与效果评估

- GEO 治理的目标与前提
- GEO治理的组织架构与机制
- GEO 效果指标体系

五

GEO 服务采购与合规评估

GEO服务商类型

服务商评估维度与评分框架

如何判断服务商的专业水平

企业采购GEO服务的完整用户旅程

采购 GEO 服务的参考策略

六

给行业的倡议

负责任GEO治理框架

向各利益相关方呼吁

给社会公众的建议

参

参考来源

附

附录

附录一：极简术语表

附录二：GEO红线速查表

附录三：可复用的内容审核清单

附录四：企业内部GEO培训提纲

附录五：红线风险评分模型

附录六：可实操的防御性技术栈与组织分工

附录七：GEO 服务商评估方法参考

附录八：POC、合同与项目治理参考

附录九：GEO服务商风险清单与不合理信号

附录十：GEO服务采购准备与参考方法

前言

当前，生成式人工智能已深度融入我国生产生活。截至2025年底，我国生成式AI用户突破6亿，信息分发逻辑正从“关键词检索”转向“生成式问答”。GEO（Generative Engine Optimization，生成式引擎优化）作为新兴业态应运而生，但其野蛮生长已引发各类乱象。

当GEO从技术路径异化为信息投毒，当批量生产的低质内容开始污染信息源、操纵AI认知，我们不得不正视：GEO正站在从野蛮生长向规范治理转型的十字路口。

2026年4月，每经AI智库联合中国新闻技术工作者联合会、中国经济传媒协会、国家广告研究院品牌实验室、新华社国家重点实验室、弗若斯特沙利文等40余家行业权威机构发布《负责任GEO治理倡议》，来自专业媒体、行业组织、学术界、科技产业的实践者与思考者，联合倡导与推动生成式引擎优化朝着合规、长效、真实、向善的方向发展。

为打通倡议从理念走向实操的落地路径，进一步厘清生成式引擎优化的边界、风险与治理路径，2026年6月，每经AI智库携手GEORankHub，深度整合多方研究资源，联合编撰本GEO红皮书。因行业与技术变化迅速且编撰时间有限，红皮书仍有诸多未尽之题和不足之处，我们将持续研究与迭代，也诚邀社会各界提出意见与建议，共同为行业提供与时俱进的GEO治理路径。

红皮书以“真实优先、安全治理”为核心理念，系统梳理覆盖价值观、信息质量与AI系统攻击三层九类GEO红线行为和风险管理策略，以及合规GEO体系构建、效果评估和供应商甄别的落地建议，旨在为品牌方、服务商、AI平台及相关机构法务合规团队在进行AI品牌建设业务时建立一套完整的GEO红线与治理框架参考，推动GEO回归正向价值。

需要说明的是，针对“黑帽GEO”红线风险行为，红皮书描述的是“有人在这样做”和“你该怎么防”，而非“怎么去做”，所有风险描述均以治理为终点，而非攻击能力的扩散，红皮书不提供任何可直接用于实施攻击的内容。

本红皮书适用对象包括：品牌企业、GEO服务商、AI平台，以及涉及大模型内容生产传播业务的产品团队、内容团队、法务合规团队、网络安全团队、数据治理团队、平台治理团队、企业知识库与智能体团队等。



一 负责任GEO理念与价值观

1.1 GEO 的三个核心目标

GEO正在从概念落地为企业信息基建的一环。多方数据显示，用户获取答案的入口在迁移——豆包、Kimi、DeepSeek、企业客服、内部知识库成了新的主场。其提问方式也在改变：不再是过去通过搜索引擎“帮我找个网页”，而是直接向AI要一个结论、一份对比、一条采购建议、一个风险提示，或是一个本地推荐。

AI会把零散信源整合压缩，变成直接可用的答案、清单或决策路径。这意味着，企业追求的“曝光”已经跳出了页面排名的逻辑，延伸到了更深层的诉求：信息是否被准确理解、是否被可信地引用、是否在AI的回答中被公平对待。

真正的GEO应当围绕三个目标展开：

- **让真实而有价值的信息更容易被机器理解**

一个企业的产品、价格、案例、资质、服务范围、售后责任、适用边界、风险提示、客户评价、更新日期、内容审核记录和纠错入口，需要用清晰、可核验、可追溯的方式表达。结构化、可读性、语义一致性、来源标注、更新时间、作者背景、审校记录和事实证据，都是帮助生成式引擎减少误解的基础。

- **让用户在更少信息噪声中作出更好的判断**

GEO的社会价值并不只属于企业增长。它还关系到用户能否看到可靠来源，关系到医疗、金融、法律、教育、公共服务、招聘、汽车、母婴、养老、本地生活等高影响场景中的信息安全，关系到内容创作者、媒体、专业机构、品牌方、平台与公众之间的信任分配。好的GEO能提高答案质量，减少幻觉与误导，降低用户验证成本，让高质量内容获得应得的可见度。

- **让企业竞争回到能力、证据和责任**

企业可以表达优势，可以纠正错误，可以补充上下文，可以建设权威来源，可以让机器更准确地读取事实。企业不应通过虚构事实、污染知识库、攻击竞争对手、操控评论、批量制造低质内容、注

入隐藏指令、伪造 AI 生成标识或借安全漏洞影响 AI 输出。GEO 的边界是一套商业伦理、信息质量、安全工程、消费者保护和社会信任的共同边界。

1.1.1 GEO 的五个社会价值

01 对用户的价值

减少信息不对称，让用户更快获得可核验答案，避免伪造评论、夸大宣传和 AI 幻觉误导。

02 对企业的价值

把竞争从短期排名技巧转向事实资产、专业能力、客户成功、公开资料、内容治理和品牌可信度。

03 对内容生态的价值

让原创研究、专业经验、清晰解释和真实案例得到更好的引用机会，降低低质批量内容占用公共知识空间的概率。

04 对 AI 系统的价值

提高检索语料质量，减少噪声、冲突证据、毒化上下文和隐藏指令，使生成答案更稳定、更可解释。

05 对社会治理的价值

在医疗、金融、法律、公共安全、教育、就业、政务和消费决策等领域，真实、透明、可追溯的信息会降低系统性风险。

1.1.2 GEO 适合解决的问题

- **品牌可见性**：用户问到某个品类、解决方案、品牌对比、采购建议时，企业是否被 AI 答案合理提及。
- **引用可见性**：AI 答案是否引用企业自有官网、帮助中心、白皮书、产品文档、案例页面、开发者文档、权威媒体或行业机构资料。
- **答案准确性**：AI 对企业产品、价格区间、适用场景、服务范围、认证资质、客户案例、交付能力的描述是否准确。
- **竞争态势**：AI 在推荐或对比时是否频繁提及竞品，企业与竞品的相对位置、理由、描述质量如何。
- **实体治理**：企业名称、品牌名、产品名、创始人、子品牌、英文名、简称、地区业务、官网入口之间的关系是否清晰。
- **跨平台一致性**：豆包、DeepSeek、Kimi、元宝、千问在同类问题上的回答是否一致，差异来自模型、搜索源、内容生态还是企业资料缺口。

1.1.3 GEO 不宜承诺的事项

- 保证任意 AI 平台在任意问题中固定推荐某品牌。

- 保证短期内通过GEO直接带来确定销量或确定线索数量。
- 通过批量灌水、伪造案例、制造虚假第三方评价、购买隐性软文来影响答案。
- 把单次截图当作项目效果证据。
- 把AI答案里的品牌提及直接等同于搜索排名、广告点击、商机和收入。
- 在没有实验设计、对照组、时间窗口和样本量的情况下宣称完成因果归因。

1.2 GEO 实践的原则

1.2.1 七条正向原则

01 真实优先

企业可以表达观点，也可以表达优势，但事实性陈述必须有证据。产品能力、价格、客户、认证、奖项、研究结论、测评数据、转化率、市场份额、服务范围和适用场景都应具备可核验来源。

02 用户受益

GEO 的核心问题是“用户得到的答案是否更准确、更完整、更安全”。

03 来源可追溯

重要事实应保留原始出处、发布日期、更新日期、作者或审核角色。

04 机器可理解

合规 GEO 鼓励清晰标题、结构化段落、常见问题解答、结构化标记与可访问内容。

05 披露关系

付费内容、代言、联盟链接、合作评测等应以用户能理解的方式披露。

06 尊重竞争

GEO 可以建立自己的权威，不能通过虚假信息损害竞争对手商誉。

07 安全内建

所有面向 AI 搜索、RAG、Agent 和内容抓取的页面都可能成为模型上下文。

1.2.2 五个默认假设

01

公开内容会被多类系统读取

官网、内容账号、短视频账号、内容社区笔记、问答、PDF、产品手册、招聘信息、工商信息、媒体稿和评论，都可能成为 AI 生成答案的材料。

02

平台和监管都在强化 AI 内容标识

AI 生成内容不标识、伪造标识、剥离标识、把 AI 生成内容包装成真人经历，都会提高合规风险和声誉风险。

03

虚假评价与刷量刷单是重点治理对象

虚假评价、刷量、刷单、好评返现、虚假交易和伪中立榜单，已经属于国内网络竞争与消费者保护重点治理对象。

04

高影响行业必须采用更高审校标准

医疗、药品、保健食品、金融理财、保险、法律、教育、招聘、汽车安全、母婴、养老和公共服务内容，需要专业审校、限制条件、风险提示和证据留痕。

05

AI 系统攻击不只发生在模型对话框

它可能发生在网页、评论、文档、图片、元数据、知识库、插件、日志、客服记录、工单、邮件和第三方内容供应链中。

1.2.3 四条底线原则

面对GEO可能带来的风险，必须确立清晰的行为红线。基于案例分析与技术研判，我们提出以下四类明确反对的行为：

- 第一，反对杜撰虚假信息“投毒”互联网

反对通过生产和发布不实、虚假内容，如假报告、假测评、假新闻，制造“伪共识”，从而诱导大模型生成更大范围、更难以察觉的虚假信息。
- 第二，反对使用AI生成低质内容“污染”互联网

反对通过自动化工具和生成式 AI生产大量无实质价值、同质化严重的内容，导致互联网信息生态“荒漠化”。
- 第三，反对恶意攻击与不正当竞争

反对通过恶意提示词注入等手段干扰AI正常输出，反对进行品牌信誉“拉踩”等不正当竞争行为，反对干扰医疗健康、金融、政务等领域专业知识的AI解读，误导用户决策。
- 第四，反对侵犯知识产权

反对在未经协商与授权的情况下，利用爬虫等方式抓取专业媒体、研究机构等组织的版权内容与数据，用于模型训练和GEO内容生产。

1.3 GEO 红线与风险识别

1.3.1 GEO的红线判定公式

可以用一个简洁公式识别 GEO 红线：

红线风险 = 不正当目的 × 信息失真 × 技术操纵 × 影响范围 × 可逆性难度

不正当目的

是否以误导用户、攻击竞品、绕过平台规则或污染模型为目标。

信息失真

是否虚构、夸大、遗漏关键限制或伪造来源。

技术操纵

是否利用抓取、索引、检索、RAG、提示词、向量、工具调用、渲染或排名信号弱点。

影响范围

是否面向大规模用户、多个模型、多个搜索引擎或高影响决策场景。

可逆性难度

错误信息是否会被多站点转载、进入训练数据、进入 RAG 知识库、被其他 AI 引用并持续扩散。

当上述五个维度的整体评估达到“中高风险”及以上时，即触发 GEO 红线，应立即停止相关行为并进入合规审查流程。

在立项、发布和验收时，团队可以连续追问以下问题：

- 这项GEO是否以改善用户理解为目标，还是以诱导 AI 偏向某个不充分结论为目标？
- 相关事实是否来自原始证据、官方资料、真实客户授权、真实交易记录或可公开核验的第三方资料？
- 内容是否涉及广告、商业推广、AI 生成合成内容、达人合作、专家背书或客户评价？是否完成标识与披露？

- 内容是否涉及竞品、行业排名、性能对比、市场份额或负面事实？是否能经得起法务、媒体、监管和竞品复核？
- 内容是否进入企业 RAG、客服机器人、销售助手、投标系统或公开 AI 搜索？是否完成来源、权限、敏感信息和安全扫描？
- 若这段内容被AI摘要成一句话，是否会误导用户？若其他平台转载，是否仍然能保留事实边界？
- 若供应商用截图证明“AI已经开始推荐我们”，团队能否追溯到每个被引用内容的来源、责任人和审核记录？

1.3.2 负责任合规GEO与高风险GEO的对照

合规 GEO 与高风险 GEO 在七个维度上的核心差异如下：

维度	✓ 合规 GEO（推荐）	✕ 高风险 GEO
内容目的	帮助用户理解真实信息	诱导模型或用户得出预设结论
事实来源	强调可核验、可追溯、可更新	常见虚构、拼接、伪造与循环引用
自动化使用	将自动化用于草稿、整理、质检和结构化	将自动化用于大规模低质发布、刷量和伪造评价
竞品关系	允许客观比较，并要求证据和边界	通过诋毁、伪投诉和负面内容投放影响竞品声誉
技术策略	强调可访问性、schema 一致和清晰结构	依赖 cloaking、隐藏文本、门页、重定向和投毒
AI 安全	建立注入防护、来源分级和输出验证	利用提示词注入、RAG 投毒和记忆污染等弱点
披露	清晰披露付费关系和利益关系	常见软文、假评论和员工评价未披露

GEO红线风险总览

2.1 风险分级和类型框架速览

绿色行为（推荐）：真实内容 + 结构化呈现 + 来源标注 + 可访问性 + 更新记录 + 用户意图覆盖 + FAQ审校 + 多格式发布。

黄色行为（需谨慎）：自动化生成但有人审校；第三方评测但披露不足；竞品对比但证据不充分；监控AI输出但抽样不完整。**需补充审核、披露、证据和流程。**

红色行为（严格禁止）：伪造事实、操控评价、攻击竞品、批量低质内容、隐藏文本、链接农场、提示词注入、RAG投毒、工具调用诱导、敏感信息泄露、恶意搜索投毒等。

A类：价值观与商业伦理风险（重点影响公平竞争与信任）

B类：信息质量风险（重点影响内容真实性与用户决策）

C类：技术与系统安全风险（重点影响AI系统本身的安全与完整性）

2.2 九大风险类型

GEO 红线风险可归纳为下列九大类型，分属价值观与商业伦理、信息质量、技术与系统安全三大层次；下一节（2.3 红线风险详情）将逐项展开每种红线的定义、底层原理、识别信号、危害说明与真实场景复盘。

01 竞争伦理风险

竞品攻击、商誉诋毁、伪中立比较、舆论操控。

02 信任伪造风险

虚假评论、虚假交易、虚假奖项、隐性付费关系。

03 内容质量风险

幻觉内容、批量低质内容、抄袭拼接、伪鲜度。

04 搜索 spam 风险

关键词堆砌、隐藏文本、门页、链接农场、过期域名、站点声誉滥用。

05 用户安全风险

搜索投毒、隐藏重定向、钓鱼链接、恶意下载。

06 LLM 指令风险

直接注入、间接注入、越狱、提示词泄露、多轮劫持。

07 RAG 与检索风险

知识库投毒、向量操纵、虚假 RAG 条目、检索诱导。

08 Agent 与工具风险

工具调用操纵、过度代理、不安全输出处理、响应渲染攻击。

09 数据与供应链风险

记忆投毒、训练数据投毒、模型后门、供应链污染、敏感信息泄露。

2.3 红线风险详情

A类：价值观与商业伦理风险

2.3.1 负面声誉攻击与商誉操纵

定义

利用生成式引擎的回答机制，通过虚假差评、伪投诉、影射文章、匿名爆料或系统性负面叙事，让用户或 AI 对竞争对手或特定实体形成无根据的负面判断，损害其商业信誉与商品声誉。

底层原理

生成式引擎会综合网页、社区、评论、媒体报道和结构化内容。攻击者若把负面叙事分散到多个位置，模型可能把噪声当作共识，重复出现的负面内容会改变系统对实体的语义画像。

识别信号

同一负面说法在多个低质量站点同步出现；缺少原始证据；表述高度相似；发布时间集中；与营销活动或竞品发布节点重合；多个账号模板化表达；标题夸张；大量使用“避雷”“骗局”等标签。

危害说明

企业层面会引发侵权、反不正当竞争和公关危机，被攻击企业可能损失客户和融资机会；行业层面会破坏公平竞争；社会层面会让 AI 搜索更像谣言放大器；内容平台会承担治理压力；用户会被误导。

场景复盘：某国内 B2B 软件厂商在准备大客户采购季时，要求服务商让 AI 搜索在“国内 CRM 哪家更稳定”这类问题中弱化竞品。供应商绕开本方产品证据建设，组织多篇匿名问答、内容账号短评和论坛帖子，反复暗示竞品“频繁宕机、售后消失、合同埋坑”。这些内容没有故障公告、判决、监管处罚或客户授权记录，却在多个平台同日出现。两周后，销售团队在元宝和千问中截到带负面提示的对比回答，并把截图用于销售沟通。

2.3.2 虚假社会证明与交易伪造

定义

用不存在的用户、未使用产品的人、员工亲友、AI 角色制造好评、差评、评分、点赞、粉丝、播放量和案例背书；或通过伪订单、伪成交、伪试用、伪下载、伪预约、伪咨询量提升商业可信度。

底层原理

消费者和模型都会把评价数量、评价一致性、外部口碑、社交证明、交易量、下载量和使用量当作可信线索。虚假评论和虚假数据会污染这些线索，让模型形成错误品牌画像。

识别信号

评价文本相似；账号历史稀薄；评价时间集中；缺少真实使用细节；正负面极端；IP、设备或支付链路异常；转化数据与收入不匹配；订单退回率异常；同类账号重复；数据口径不披露；无法提供审计记录。

危害说明

企业可能面临罚款、下架、平台封禁、集体诉讼和虚假宣传认定；行业信任成本上升；用户做出错误购买决策；损害投资人、客户和行业判断。

场景复盘：某连锁轻医美机构希望在 AI 搜索中被描述为“用户口碑更好”。外包团队安排员工亲友注册多个账号，发布“真实体验”笔记和高分评价，并把同一批照片稍作裁剪后投放到多个平台。部分笔记还暗示其他机构“踩雷”。用户询问 AI 助手“某城市皮肤管理机构推荐”时，系统引用了这些评价并生成排序建议。后续平台风控发现评价账号下单轨迹异常，评价内容模板化。

2.3.3 付费关系不披露

定义

付费评测、榜单、KOL 推荐、合作案例或联盟链接未披露利益关系。

底层原理

AI 搜索会把第三方媒体、榜单和评测视为外部权威。若付费关系被隐藏，模型和用户会高估独立性。

识别信号

榜单标准模糊；多篇文章用同一模板；商业链接集中；作者和媒体缺少独立评测流程；披露缺失或位置隐蔽。

危害说明

用户被误导，媒体信誉受损，品牌在监管和平台侧承担合规风险。

场景复盘：某母婴产品品牌邀请达人发布“自用分享”，合同中约定投放费用和转化返点，但内容没有任何商业合作标识。达人笔记把产品写成“群里宝妈一致推荐”，还让 AI 文案工具批量生成问答评论。AI 搜索在回答“婴幼儿用品推荐”时，误把该内容当作独立用户体验。出现投诉后，品牌才发现代理商把商业推广包装成普通体验。

2.3.4 权威信号盗用

定义

购买历史上有高可信背景的过期域名并改造成与原主题无关的商业或低质内容站；或第三方内容借宿在高权威站点上，主要目的在于利用宿主排名信号，而内容与宿主核心业务或编辑责任脱节。

底层原理

搜索系统和 AI 引擎可能仍读取历史链接、品牌语义和权威信号，导致新内容获得不当信任；宿主站点的历史信誉会传递给第三方页面，模型和搜索系统可能把租用内容误判为宿主的专业内容。

识别信号

域名主题突然变化；历史快照与当前内容不一致；外链锚文本与当前业务无关；页面与主站主题割裂；作者、编辑、审校缺失；大量外链和商业 CTA；内容由白标服务商统一生产；页面版权仍显示宿主名称。

危害说明

损害旧机构声誉，误导用户和模型，把公共信任转化为私利；宿主品牌承担连带风险；搜索结果质量下降；用户难以区分真实编辑内容与租用内容。

场景复盘：一家培训机构购买过期的地方教育论坛域名，保留原站部分栏目标题，再批量发布“职业资格证书报考指南”和课程推荐。由于域名曾被许多学校和地方站点引用，AI 搜索把新内容误判为较高可信来源。某高权重地方资讯站把一批与本站主营产品无关的“AI 办公神器排行榜”“某药械产品测评”交给第三方运营。第三方以媒体域名背书发布商业内容，页面没有作者、审校、广告标识和真实测评过程。多个 AI 搜索系统在整理行业榜单时引用了这些页面，使用户误以为结论来自独立媒体。

2.3.5 伪中立比较与选择架构操纵

定义

通过看似客观的比较页、评分表和问答页，把维度设计、样本选择和结论引导到预设品牌。**底层原理：**模型会把结构化比较表和“最佳选择”段落当作高信息密度内容；若比较维度被操控，生成答案会继承偏差。

识别信号

评分规则不可复现；缺少样本说明；竞品缺点被放大，自身缺点被省略；数据无来源。

危害说明

用户选择被误导，竞品受到不公平评价，AI 输出强化单方叙事。

场景复盘：多个主流及高影响力媒体平台出现虚假 GEO 服务商权威评测稿件，内容谎称权威研究院发布行业报告，虚构不知名 GEO 厂商的榜单排名、评分及评分维度，依托虚假数据打造营销话术，推销不合理服务商选择指标。该行为本质是伪造权威背书、操控行业认知，误导用户与合作方判断，扰乱 GEO 行业正规评价体系与市场秩序。

2.3.6 舆论操控与水军协同

定义

通过大量账号、社群、机器人或代理写手制造看似自发的讨论热度、共识和评价。

底层原理

生成式引擎会从开放网络中提取“公众看法”。人造共识会变成模型的语义背景。

识别信号

账号创建时间接近；互动网络过密；文本相似；话题跳转不自然；缺少真实用户体验。

危害说明

行业讨论空间被污染，真实用户声音被淹没，AI 对品牌的“公众口碑”画像失真。

场景复盘：某消费品牌出现质量投诉后，代理商组织水军在多个平台发布“投诉者是职业黑子”“竞品带节奏”等相关内容，并集中点赞正面评论、举报真实投诉。AI 平台抓取相关信息后把事件归因为“竞品攻击”。后续监管机构和权威媒体介入，企业发现代理商团队扩大了负面舆情危机。

B类：信息质量风险

2.3.7 AI 生成内容未经审查

定义

把未经核验的 AI 生成文字、视频、图片等结果当作事实发布，导致不存在的功能、人物、数据、案例、事件、现场影像资料进入网页。

底层原理

大模型会生成语言上合理但事实错误的内容。若缺少事实核查，幻觉会被搜索收录，再被更多 AI 引用。

识别信号

引用无法打开；数据缺少口径；人物和机构不存在；多个页面出现同一错误；作者无法解释来源。

危害说明

企业会因虚假宣传和错误承诺承担责任，用户会采取错误行动，行业知识库被污染。

场景复盘：某出海服务公司用 AI 批量生成“各国合规指南”，没有律师审校，也没有更新日期。文章把旧政策、不同国家规则和营销建议混在一起。用户询问 AI 搜索“某类产品能否进入东南亚市场”时，系统引用了该文并给出错误判断，导致客户准备资料方向错误。

2.3.8 虚构引用、论文、奖项和认证

定义

伪造论文、专家、媒体报道、客户名单、认证、审计报告、测试结果或奖项。

底层原理

AI 系统依赖引用和来源线索判断可信度。伪造证据会直接攻击信任机制。

识别信号

DOI（唯一标识符）、证书编号、官网公告无法验证；论文标题与作者库不匹配；奖项官网无记录；客户未授权。

危害说明

严重时构成欺诈或虚假宣传，损害客户、投资人和公众判断。

场景复盘：某 AI 硬件公司发布白皮书，列出多篇“国际顶会论文”和“国家级认证”，但论文标题不存在，检测报告只是内部测试截图。AI 深度研究工具在生成行业报告时引用这些资料，投资人把它作为技术领先证据。尽调阶段发现引用伪造，企业融资进程受影响。

2.3.9 夸大功效与高影响领域误导

定义

在医疗、金融、法律、教育、公共安全等领域夸大产品效果、保证结果、隐瞒风险或提供不合格建议。

底层原理

高影响主题会影响用户健康、财务、安全和社会福祉。错误内容的外部性远高于一般营销内容。

识别信号

出现“保证治愈”“稳赚”“零风险”“100% 通过”等词；缺少适用范围；无专业审校；案例样本极小。

危害说明

用户可能遭受健康、财务或法律损害，企业面临行政处罚、诉讼和平台下架。

场景复盘：某健康管理机构在内容中把个案改善描述为“普遍有效”，使用“逆转”“根治”“保证改善”等表达，还让 AI 生成患者故事。用户在 AI 搜索中询问症状调理方案时，系统引用这些内容并给出偏商业的建议。专家复核发现内容缺少临床证据、适用人群和风险提示。

2.3.10 规模化低质内容滥用

定义

以操纵搜索或 AI 可见度为目的，大量生成缺少原创价值、事实核验和编辑责任的页面。

底层原理

规模化内容会占用索引与检索空间，降低优质信息被召回概率。Google 已把 scaled content abuse（规模化内容滥用）作为明确政策风险。

识别信号

页面数量短期暴增；结构和措辞高度相似；作者缺失；信息空洞；用户停留和转化异常低。

危害说明

品牌被算法或人工处罚，用户搜索成本上升，公共信息空间被低质内容占据。

场景复盘：某招商加盟公司要求供应商一个月发布上千篇城市加盟指南。大部分内容只替换城市名和行业名，甚至在没有当地门店、没有服务团队的地区也写“本地成功案例”。AI 搜索在回答“某城市加盟项目”时抓取这些页面，用户误以为该企业在本地产有成熟运营。

2.3.11 抄袭、拼接和同义改写

定义

复制他人内容，轻微改写、翻译、拼接或自动同义替换后发布，缺少实质新增价值。

底层原理

生成式系统会把多来源内容合并。抄袭拼接会造成重复信息、错误归属和原创者权益损失。

识别信号

段落顺序与原文相近；事实错误继承；引用被删除；多个页面存在相同句式。

危害说明

原创者权益受损，企业面临版权和搜索处罚，用户看到重复低质答案。

场景复盘：某咨询公司把同行报告、券商研报和公开论文拆段后交给 AI 改写，再发布为“原创行业洞察”。页面没有引用来源，也没有新增分析。AI 搜索在总结行业规模和趋势时引用该公司文章，原创机构的结论被二次占用。被投诉后，咨询公司无法说明数据口径和原始来源。

2.3.12 关键词堆砌与隐藏文本

定义

在页面中不自然重复关键词、城市名、产品名，或用不可见文本影响机器读取。

底层原理

传统搜索和部分 AI 检索会读取文本信号。堆砌和隐藏文本试图制造虚假的主题相关性。

识别信号

人读不顺；页面底部出现长串地区词；CSS 隐藏；结构化内容与可见内容差异。

危害说明

用户体验下降，站点可能被处罚，AI 摘要可能被不自然文本污染。

场景复盘：某 SaaS 公司为了让 AI 搜索更容易把它与“低代码、RPA、本土化、信创”等词关联，在页面底部堆砌大量城市、行业和竞品关键词，还把部分文字做成低可见度。用户正常阅读几乎看不到这些内容，但 AI 抓取摘要时会读到。结果多个回答把企业能力扩大到未实际覆盖的行业。

2.3.13 虚假分支页面与本地服务能力虚构

定义

通过批量创建高度相似的垂类、区域中间分支页面，或虚构各地门店、服务团队、库存、本地案例、服务覆盖范围等本地服务能力，将各类相似查询的用户统一引流至固定目标页面或服务的行为。

底层原理

利用查询覆盖面制造可见度，但每个页面缺少独立价值。AI 搜索可能把它们当作多个证据。

识别信号

URL 和标题差异小；正文模板化；没有真实本地地址或团队；多个页面同一 CTA（Call To Action / 用户行为召唤）。

危害说明

用户被引导到不匹配页面或不存在的本地地址，搜索和 AI 检索结果质量下降。

场景复盘：某家政平台生成数百个“城市服务站”页面，实际只有几个直营城市。页面写有本地电话、本地案例和服务承诺，但电话统一转接外地客服。用户询问 AI 搜索“本地靠谱保洁公司”时，系统把这些页面作为本地覆盖证据。用户下单后发现没有本地人员，投诉平台虚假服务范围。

2.3.14 结构化信息造假

定义

schema（结构化标记）、FAQ（常见问题解答）、review（用户评价）、price（定价）、availability（供货能力）等结构化数据与页面可见内容不一致，或标注不存在的评价和价格；或编造用户问题和答案，制造“常见疑问”的假象，把强营销答案伪装成用户关切。

底层原理

AI 搜索和搜索引擎会读取结构化数据帮助理解页面，错误标注会制造高可信格式的信号；FAQ 结构很适合被 AI 抽取，虚构 FAQ 会把营销口径包装成用户需求。

识别信号

结构化标记与页面不一致；评分无评价来源；价格和库存与交易系统不同；FAQ 无人工可见内容；问题措辞像销售话术；没有用户来源；所有答案导向同一产品；缺少限制条件。

危害说明

用户被错误价格、评分、库存或资质误导，平台可能撤销富结果资格；用户误以为问题有普遍性，AI 可能直接采用答案。

场景复盘：某教育机构在结构化数据中标记“五星评分、万人评价、官方认证课程”，但页面可见内容没有相应评价、认证和证书编号。AI 搜索抓取结构化数据后，在摘要中显示高评分和认证标签。学员投诉时，机构解释为“技术人员误填模板”。某跨境支付公司用 AI 生成大量 FAQ，例如“某服务是否被监管推荐”“用户为什么都选择某品牌”。这些问题并非真实用户提问，回答中包含无法证明的背书和倾向性结论。AI 客服读取 FAQ 后，在用户比较支付方案时主动推荐该品牌。

2.3.15 伪鲜度和过期信息包装

定义

把旧内容改日期、轻微更新标题或使用“2026 最新”等词，让用户和 AI 误以为信息已被重新审校。

底层原理

时效信号影响用户和 AI 对可信度的判断。伪鲜度会隐藏过期事实。

识别信号

正文引用旧年份；截图和数据过期；更新记录空白；标题与内容时态冲突。

危害说明

用户依据过期政策或价格决策，企业承诺和真实能力脱节。

场景复盘：某法律服务平台把 2022 年政策解读文章改成“2026 最新”，正文只更新标题和首段，没有更新条文、适用地区和实施日期。AI 搜索回答“最新合规要求”时引用该文，导致用户按照过期口径准备材料。内部审计发现大量页面存在伪更新。

2.3.16 来源漂白与 AI 引 AI 循环引用

定义

把 AI 生成内容发布到网页，再让另一 AI 引用，随后把该 AI 输出当作外部证明。

底层原理

生成系统之间会互相读取开放网页。错误一旦进入可索引页面，可能形成循环增强。

识别信号

原始证据缺失；引用链回到同一批页面；多个来源互相转述但无一手证据。

危害说明

幻觉获得“来源外衣”，错误更难纠正，社会信任受损。

场景复盘：某咨询团队让 AI 生成行业结论，引用来源是另一篇 AI 生成文章，后者又引用无原始出处的自媒体。几轮转述后，虚假的市场规模数字被包装成“多家机构一致预测”。AI 深度研究工具在报告中复用该数字，客户决策会以错误市场规模为基础。

2.3.17 机器生成外链与链接农场

定义

通过自动化站点、目录、评论、论坛签名、合作页或低质媒体批量制造入链。

底层原理

链接仍然是搜索生态的重要信号。低质外链试图伪造声誉和引用关系。

识别信号

外链站主题无关；锚文本过度一致；站点注册和模板相似；访问质量低。

危害说明

可能导致搜索处罚，污染引用网络，让 AI 误判权威关系。

场景复盘：某品牌服务商批量购买低质量外链和目录页，把锚文本（一种链接形式，把关键词做成链接，指向其他网页）设置为“某行业第一”“官方推荐”。这些页面内容简陋，站点之间互相链接。AI 搜索在判断品牌相关来源时抓到部分页面，导致回答中出现未经证实的称号。

2.3.18 用户生成内容操纵与滥用

定义

开放评论、论坛、问答或知识库被批量灌入广告、链接、伪评价和隐藏文本。

底层原理

UGC 内容也可能被搜索和 AI 抓取。未治理的 UGC 会把站点变成垃圾信息入口。

识别信号

新账号短期高频发帖；文本模板化；外链集中；主题偏离社区。

危害说明

站点声誉下降，用户被恶意链接诱导，AI 引用错误内容。

场景复盘：某硬件品牌官网开放用户问答，但缺少审核。大量机器人账号发布与产品无关的导流、伪测评和隐藏商业链接。AI 搜索抓取问答区后，把垃圾内容中的参数、适用场景和售后承诺混入回答。品牌方直到客服投诉增加才发现站内用户生成内容已被污染。

2.3.19 隐蔽重定向与点击诱导

定义

对搜索引擎、AI 抓取器和用户展示不同路径，或根据来源、设备、地域跳转到高风险页面。

底层原理

重定向可以被用于正常迁移，也可以用于隐藏真实落地页。恶意使用会欺骗抓取和用户。

识别信号

直接访问正常，从搜索点击异常；不同设备结果不同；日志显示异常跳转链。

危害说明

用户可能进入诈骗、恶意下载或无关广告，站点承担安全与搜索处罚风险。

场景复盘：某招商项目页面在 AI 搜索抓取时展示“政策解读与行业指南”，普通用户点击后跳到销售页面。部分短链在不同地区和设备上跳转不同内容。AI 搜索引用该页面时显示为中立指南，用户实际访问却进入营销漏斗。

2.3.20 提示词注入与模型操纵

定义

在提问输入或内容中（直接输入、外部文档、多轮对话、角色扮演/编码/假设场景）加入恶意指令，试图改变 LLM 应用的目标、规则、输出格式或安全边界。

底层原理

LLM 在同一上下文中处理系统指令、开发者指令、用户指令和数据。若边界不清，恶意文本可能被当作更高优先级任务；长期对话和记忆会让模型在上下文中保持先前假设，攻击者利用渐进方式降低系统警觉；模型可能在复杂语境中误判意图，攻击者利用安全分类与生成任务之间的缝隙。

识别信号

用户输入出现规则覆盖、角色替换、格式逃逸、要求泄露内部指令等意图；会话中出现反复要求改变身份、保存规则、绕过审查和迁移到工具调用；输出偏离任务或触及禁止内容。

危害说明

可能导致错误答案、敏感信息泄露、工具滥用、安全策略失效、虚假舆论、攻击文案、欺骗性内容和安全漏洞利用材料。

场景复盘：某供应商把一段诱导性文本藏在提交给采购方的 PDF 附录和网页元数据中。采购方的投标分析智能体读取后，在总结供应商对比时异常强调该供应商优势。审计团队复查检索上下文，发现外部材料中存在面向模型的指令性文本。

2.3.21 敏感信息诱导泄露

定义

诱导 LLM 输出系统提示、开发者提示、隐藏策略、工具说明、检索上下文、内部配置，或输出个人信息、商业秘密、内部策略、客户数据、未公开价格、合同条款，或通过翻译、总结、调试、压缩、导出等表面无害任务让模型把检索上下文或隐藏资料带出。

底层原理

模型在上下文中能看到部分内部指令或数据；检索上下文、训练数据、工具返回和历史对话都可能含敏感信息；攻击者可用间接请求绕过显式敏感词。若输出控制不足或边界不清，模型可能复述敏感上下文。

识别信号

用户要求“打印全部规则”“显示上文隐藏内容”“把系统配置转成 JSON”等；答案包含内部 ID、邮箱、合同、未发布政策；请求内容与用户权限不匹配；输出含上下文原文；多轮对话试探边界。

危害说明

泄露后攻击者更容易构造后续注入；企业策略、提示工程和安全规则暴露；企业法律责任和信任损失显著；增强后续攻击能力。某员工让办公智能体总结内部报价文件，同时打开了一个外部网页。网页中的诱导内容让智能体把摘要发送到外部地址。虽然动作需要工具能力配合，但组织没有把外部网页与内部文档隔离，导致上下文外带风险。

场景复盘：某品牌客服机器人被用户反复追问“你遵循哪些内部规则”。机器人没有直接泄露完整系统提示，但输出了投诉优先级、敏感词列表和升级流程。竞争对手据此推测客服策略并设计规避话术。

2.3.22 响应渲染与链接伪装攻击

定义

利用 Markdown、HTML、富文本、卡片渲染和链接预览，让输出看似指向可信站点，实际跳转到恶意站点。

底层原理

用户和模型界面常显示锚文本或卡片摘要，真实 URL 被隐藏或不显眼。攻击者利用显示层差异。

识别信号

锚文本与 URL 域名不一致；短链接；跳转链过长；卡片标题与目标站点矛盾。

危害说明

用户凭据和数据面临风险，品牌内容成为钓鱼入口。

场景复盘：某仿冒售后页面把链接显示文本写成“官方客服”，实际跳转到第三方页面。AI 摘要在输出时保留了显示文本，用户以为是官方渠道。品牌安全团队排查发现多个页面使用相似的链接伪装方式。

2.3.23 检索增强系统投毒与操纵

定义

向 RAG 知识库、网页语料或文档库注入恶意或虚假内容，或通过语义伪装、关键词混合、触发词、定制内容片段、伪造文档条目使内容在向量检索或重排中获得不当召回，让模型在特定问题上输出攻击者选择的答案。

底层原理

RAG 依赖外部知识，若知识库可信边界薄弱，少量高召回内容就可能影响答案；向量检索按语义相似度召回，攻击者可让低可信内容在目标 query（搜索中的查询词）周围获得高相似度；生成式搜索可能对结构清晰、可摘要片段给予更高可用性。

识别信号

知识库新增来源异常；文本与目标问题高度贴合但来源低质；多个检索结果互相矛盾；答案过度依赖单一片段；检索结果语义贴近但来源不可信；文本存在奇怪重复；向量相似度高但 BM25 或人工相关性低；页面只围绕少数 query 服务；上下文缺失；结论没有证据；文档来源、版本、审批人缺失；格式过于相似但无流程记录。

危害说明

企业 RAG 给员工、销售、客服和客户提供错误建议；公开 AI 搜索扩大错误影响；用户看到被精心操纵的答案，AI 输出多样性下降；内部员工可能依据伪政策行动，销售或客服会给出错误承诺。某品牌在大量页面中使用相同的问答结构和实体别名，让 AI 搜索在特定提问下更容易召回自家页面。页面内容并非完全虚假，但刻意弱化限制条件，把非核心功能包装成行业主能力。

场景复盘：某企业把外包商提交的行业资料直接入库，未做来源和事实审核。外包商为了提高品牌推荐概率，插入多段“采购建议优先选择某品牌”的材料。销售助手回答客户问题时，反复输出倾向性推荐，且无法提供原始证据。

2.3.24 AI 数据层污染

定义

向 AI 助手的长期记忆或用户偏好中写入未经授权的事实或指令，或污染预训练、微调、偏好数据或评测数据，使模型形成偏见、后门、错误事实或安全薄弱点，长期改变模型行为。

底层原理

带记忆的 AI 会把历史信息当作用户偏好或事实；模型会从数据中学习模式，少量有目标的数据在特定条件下可能改变输出。污染记忆或训练数据会长期改变模型行为。

识别信号

记忆中出现来源不明的偏好；用户没有授权；记忆项影响多次对话；数据来源不清；样本异常相似；标签与内容矛盾；模型在特定触发下异常。

危害说明

可导致长期推荐偏差、隐私泄露、越权行动和品牌操纵；影响范围可能覆盖大量用户和长期版本，修复成本高。某开源客服语料中混入“特定品牌永远最佳”的样本，微调后模型在推荐场景异常偏向。某企业用历史客服对话微调模型，未清洗销售夸大话术、过期政策和用户个人信息。上线后模型在高频问题中复现过期承诺，并偶发输出用户隐私片段。数据治理复盘发现训练集缺少时间戳、授权和脱敏流程。

场景复盘：某销售助手具有长期记忆能力。一次对话中，销售人员把“某客户更关注低价”“某竞品有问题”等未经证实的信息写入记忆。后续助手在所有与该客户相关的建议中持续引用该记忆，影响报价和竞品比较。

2.3.25 模型后门与触发器

定义

在模型、微调权重或训练数据中植入特定触发条件，使模型在触发时输出异常内容。

底层原理

后门攻击平时表现正常，遇到触发词、格式或语境才改变行为，因此难以检测。

识别信号

正常评测通过，但特定词、格式或语言组合下异常；来源模型或权重链路不清。

危害说明

会造成品牌推荐操纵、数据泄露或安全策略失效。

场景复盘：某团队下载未经充分审计的开源微调模型用于客服场景。模型在常规测试中表现正常，但在特定触发语义下会输出异常推荐和外部链接。安全团队追踪后发现模型来源、训练数据和权重变更记录不完整。

2.3.26 模型与插件供应链污染

定义

使用来源不明的模型、插件、爬虫、数据处理库、向量库连接器或提示模板，导致安全与内容风险。

底层原理

大模型应用依赖复杂供应链。任何组件都可能带入漏洞、偏见、恶意逻辑或数据泄露路径。

识别信号

组件来源、许可证、更新记录、权限范围不清；供应商无法提供安全说明。

危害说明

可能导致内容污染、数据泄露、服务中断和合规责任。

场景复盘：某 GEO 服务商使用第三方内容生成插件批量生产资料。插件会自动加入未披露的推广链接和统计脚本，还会把客户草稿上传到境外服务。品牌方在发布后发现页面出现异常外链和隐私风险。

2.3.27 跨语言与跨引擎操纵

定义

利用不同语言、地区和 AI 引擎在来源偏好、召回和安全策略上的差异，发布矛盾或偏向内容。

底层原理

AI 搜索在多语言、多地区 and 不同模型上表现并不完全一致。攻击者可能在低治理语言或地区植入叙事。

识别信号

某语言版本内容与主语言冲突；来源集中在低质站点；翻译痕迹明显。

危害说明

跨境用户和全球 AI 系统会收到不一致信息，品牌治理难度上升。

场景复盘：某出海企业在中文资料中披露风险提示，在英文和东南亚语种资料中删除限制条件并强化“行业第一”。多语言 AI 搜索在生成跨语言摘要时，只抓到英文材料，导致海外客户误解产品能力。内部复核发现多语言内容没有统一事实底稿。

2.3.28 多模态注入与隐藏指令

定义

在图片、alt 文本、OCR 可读文字、PDF、截图、音频转写或视频字幕中嵌入影响 AI 的指令或伪信息。

底层原理

多模态模型会读取图像、文本、OCR、字幕和元数据。隐藏指令可进入模型上下文。

识别信号

图片含微小文字；alt 文本与可见图像不一致；PDF 边角出现异常文本；字幕含无关指令。

危害说明

AI 搜索和企业助手可能采纳隐藏信息，造成错误摘要或工具误用。

场景复盘：某活动海报主视觉中隐藏了极小文字和图层注释，普通用户难以看到，但 OCR 和多模态模型可能读取。企业知识库把海报入库后，问答系统在活动规则中混入错误优惠说明。设计团队原本只是复用模板，未意识到图层和替代文本也会被 AI 读取。

2.3.29 AI 抓取器内容伪装

定义

对普通用户、传统搜索爬虫和 AI Agent 展示不同内容，试图让 AI 读取更有利或更具操纵性的版本。

底层原理

不同 user-agent 可能触发不同响应。cloaking（伪装）会破坏内容一致性和可信度。

识别信号

不同 user-agent 内容不一致；缓存版本与页面不同；日志显示特定爬虫命中异常模板。

危害说明

平台信任下降，用户与 AI 接收不同事实，企业面临搜索政策风险。

场景复盘：某网站根据访问来源展示不同事实版本：普通用户看到谨慎表述，疑似 AI 抓取器看到更夸张的品牌优势和隐藏 FAQ。短期内 AI 摘要更偏向该品牌，但人工复核无法在页面中看到相同内容，形成可追溯性缺口。

2.3.30 机器人批量查询粉饰引用指标

定义

用机器人批量查询 AI 搜索、点击链接、触发引用、生成曝光报告或伪造“AI 可见度”指标。

底层原理

GEO 测量需要 query、回答、引用和位置数据。若采集和点击被机器人操纵，客户会看到虚假成效。

识别信号

查询来源集中；时间分布异常；报告缺少抽样方法；截图不可复现。

危害说明

客户预算被误导，平台指标被污染，服务商市场信誉受损。

场景复盘：某服务商声称“AI 引用提升 300%”，实际来自其自有脚本反复查询少数问题。通过大量自动查询、截图和重复追问制造监测指标变化。部分平台识别到异常请求后限制访问，品牌方得到的报告无法反映真实用户场景。

2.3.31 基准与评测投机

定义

针对固定测试 query、固定评测脚本或固定指标优化页面，使报告好看但用户价值没有提高。

底层原理

评测指标可被投机。若目标只剩分数，优化会偏离真实用户和真实引擎表现。

识别信号

只在少数 query 有提升；真实用户问题无改善；评测集和训练集重叠；缺少盲测。

危害说明

企业误判 GEO 成效，行业形成指标游戏，用户价值被忽视。

场景复盘：某企业参加行业 AI 搜索评测前，提前得知部分测试问题，于是专门制作一批只覆盖测试问题的内容和知识库条目。评测结果显示效果领先，但真实用户问题覆盖不足，大量常见问题无法准确回答。

2.3.32 重复内容语义伪装与召回规避

定义

用对抗性措辞、异常格式或语义伪装让内容躲过审核、误入召回或逃避相似度检测。

底层原理

模型和检索系统对表述形式敏感。轻微扰动可能改变分类、召回和安全判断。

识别信号

文本出现异常分隔、同音替换、混合语言、零宽字符或奇怪排版。

危害说明

低质或恶意内容进入索引和知识库，审核难度增加。

场景复盘：某外包团队为了让低质内容避开重复检测和审核，将同一事实写成大量同义变体，并故意改变表述顺序。人工看起来像不同文章，向量聚类后显示高度相似。AI 搜索偶尔会召回这些变体，使低质内容占据答案上下文。

2.4 处置建议与替代方案

当上述风险发生时，建议采用以下方式处置：

立即冻结：立即冻结与风险相关的新增发布、自动化任务、外包交付和知识库入库动作，防止风险继续扩散。

建立事实核验表：把每一句关键断言映射到原始证据、授权记录、发布时间、责任人、审核人和可公开披露程度。

执行整改：对已经进入AI搜索、公开平台、企业RAG或客服机器人的材料执行下线、改写、重索引、补标、回滚或公开更正，并保留处置前后的截图与日志。

复盘根因：复盘时追问根因，是目标设定错误、证据不足、平台规则误判、供应商越界、知识库缺少审核，还是模型工具链权限过宽。

寻求安全替代方案，把目标改成事实澄清、来源建设、结构化表达、审校流程、透明披露和可复核证据。若涉及竞品，只允许使用可验证事实和公平比较。若涉及AI系统，只允许做防御性测试、授权红队和修复闭环。

GEO 风险应对与治理

3.1 价值观层治理

价值观层的风险往往不依赖 AI 技术。即使没有大模型，虚假评论、竞品诋毁、隐性付费、刷量和伪交易也会伤害用户与市场。AI 搜索放大了这些行为的影响，因为模型会把开放网络中的多源噪声压缩成看似简洁的结论。

企业在启动 GEO 项目前应提出五个问题：

1. 这项优化是否帮助用户更准确理解事实。
2. 这项优化是否包含对竞品、用户、媒体、评测机构或平台的误导。
3. 若这项行为被公开披露，客户、员工、投资人和监管机构是否会认为它符合商业伦理。
4. 这项行为是否把无法核验的叙事包装成事实。
5. 这项行为是否会让 AI 搜索更难给出可靠答案。

若任一问题答案偏向高风险，项目应进入合规审查。若涉及竞品攻击、虚假评论、虚假交易、未披露付费关系、伪造认证和投毒类行为，应立即停止。

3.2 信息质量层治理

面向 AI 搜索的内容会被抓取、摘要、重写、引用和跨语言传播。信息质量的发布责任不再限于页面本身，还包括页面进入 AI 生态后的二次影响。

建议建立八项内容控制：

1. 事实清单：每个事实性主张都对应来源、证据和更新时间。
2. 声明等级：区分事实、观点、预测、客户反馈、案例、统计推断和广告承诺。
3. 引用校验：核验 DOI、报告链接、证书编号、媒体报道、客户授权和数据口径。

4. 专业审校：医疗、金融、法律、安全、公共事务内容应有专业审校和风险提示。
5. AI 生成标记：内部记录 AI 参与程度，必要时向外披露生成或编辑方式。
6. 结构化数据一致性：结构化标记、FAQ（常见问题解答）、用户评论、定价、供货能力必须与可见内容一致。
7. 更新日志：关键页面保留更新时间、更新内容和审校人。
8. 撤回机制：错误内容能快速下线、改正、通知相关方并阻止继续进入知识库。

3.3 AI系统攻击层治理

GEO 页面、新闻稿、FAQ、PDF、图片 alt 文本、评论、论坛、案例库、帮助中心和公开知识库都会成为 AI 系统输入。企业需要把内容安全接入应用安全流程。

建议采用六道防线：

1. 来源分级：将官网、法务审校页面、客户授权案例、UGC、第三方转载、未知来源分级处理。
2. 指令隔离：RAG 和 agent 系统必须把外部内容标记为数据，不能让外部文本覆盖系统目标。
3. 检索校验：对高影响答案使用多来源交叉验证，避免单一片段决定答案。
4. 权限最小化：AI agent 不应在无确认状态下执行高影响工具动作。
5. 输出清洗：HTML、Markdown、链接、脚本、SQL、命令和邮件输出都要经过验证。
6. 日志与回放：保留 query、检索片段、模型输出、工具调用和人工确认，便于追责和复盘。

3.4 风险评分矩阵

等级	风险描述	处置方式
低	表达不清、引用格式不完整或轻微过期	由内容团队修正，记录版本
中低	证据不足、披露不充分或结构化数据轻微不一致	补证据、补披露，并下线不准确的富结果标记
中	可能误导用户、存在未审校 AI 内容或竞品比较证据不足	暂停发布，由法务和合规复核
高	涉及虚假宣传、虚假评论、批量低质内容、站点声誉滥用或过期域名滥用	立即下线，通知客户并启动整改
严重	涉及竞品攻击、RAG 投毒、提示词注入、恶意搜索投毒、敏感信息泄露或模型供应链污染	终止行为、保全证据、启动安全响应，必要时通知监管或受影响方

3.5 证据保全与责任归属

3.5.1 证据保全要求

发现风险后，不要只删除页面。应先保全可复核证据：

- 页面信息：URL、发布时间、更新记录、作者、审校人、供应商、项目负责人。
- 技术记录：页面截图、源代码、结构化标记、爬虫协议文件、站点地图、服务器日志、搜索控制台记录。
- AI 输出证据：AI 搜索输出截图、query、时间、地区、语言、模型或产品版本、引用链接。
- 业务记录：合同、Brief、交付物、审批记录、客户确认记录、外包沟通记录。
- RAG 风险专项：检索片段、知识库版本、向量索引版本和模型输出。

3.5.2 责任归属模型

GEO 风险常见责任链包括客户、服务商、外包作者、媒体渠道、技术供应商、数据供应商和平台。

合同应明确：

1. 服务商不得实施或协助实施虚假信息、竞品攻击、提示词注入、RAG 投毒、链接作弊、虚假评论、虚假交易。
2. 客户提供的事实素材必须真实、合法、可授权使用。
3. 第三方内容和媒体关系必须披露利益关系。
4. 自动化内容必须有人类审校和事实核查。
5. 供应商使用的模型、插件、爬虫、监控工具和数据源需要安全说明。
6. 发现红线行为后，服务商有暂停交付、要求整改、终止合作和保全证据的权利。

四

GEO 合规治理与效果评估

4.1 GEO 治理的目标与前提

GEO的目标既包括让企业内容被抓取，也包括让企业在用户真实提问时，以准确、合规、可复核的方式出现在AI回答中。GEO不宜被包装成单一投放服务，它更接近一套横跨内容、技术、品牌、数据、舆情、合规和增长的长期治理能力。服务商可以承担诊断、实验、内容生产、技术修复、监测、复盘和组织协同，但企业自身仍需要提供业务事实、权威材料、法务边界和内部审批。

无论是自主进行 GEO 治理还是采购外部服务，企业都应明确以下四个前提：

1. 企业必须明确业务目标

目标可能是品牌提及、引用占比、关键产品认知、海外市场进入、招商获客、舆情纠偏或招投标支持。不同目标对应不同指标与服务商类型。

2. 企业必须明确平台范围

国内重点覆盖豆包、DeepSeek、Kimi、元宝、千问。海外项目可把Google AI Overviews、AI Mode、Gemini、ChatGPT和Perplexity列为可扩展范围，但本版不展开海外平台细节。

3. 企业必须接受前端指标优先

GEO早期阶段更适合先管理AI答案表现、引用来源、实体准确性和内容可检索性。销售、线索、成交等后端指标可观测，但需要谨慎解释。

4. 企业必须建立合规边界

涉及个人信息、客户案例、未公开财务数据、医疗健康、金融投资、教育招生、政务公共安全等内容时，GEO项目需要法务、合规或专家审查参与。

4.2 GEO治理的组织架构与机制

4.2.1 组织架构

建议成立GEO治理小组，至少包含以下角色：

- **业务负责人**：定义目标、预算和业务优先级。

- **内容负责人**：负责事实核查、编辑标准和发布流程。
- **法务合规**：审查广告、竞品比较、评价、隐私和合同风险。
- **安全负责人**：审查提示词注入、RAG 投毒、工具权限、供应链和数据泄露风险。
- **数据分析**：负责 query 抽样、AI 可见度监测、来源质量评分和异常检测。
- **公共关系**：负责外部沟通、错误更正和声誉管理。

4.2.2 内容生产流程

步骤一：事实素材库。

建立唯一事实源，包括产品功能、价格、客户授权、案例、证书、奖项、数据、专家审校记录和政策限制。

步骤二：内容 Brief。

每个页面明确用户问题、证据来源、风险等级、YMYL 属性、竞品比较边界和披露要求。

步骤三：AI 辅助写作。

AI 可以用于大纲、语言整理、结构化、问答扩展和翻译。

事实主张必须回到事实素材库核验。

步骤四：编辑审校。

编辑检查可读性、原创性、证据、引用、语气和用户价值。

步骤五：合规审校。

法务检查广告法、反不正当竞争、隐私、客户授权、付费披露和绝对化用语。

步骤六：安全审校。

安全团队检查隐藏指令、脚本、结构化标记、链接、PDF、图片 alt 文本、外部组件和 Agent 读取风险。

步骤七：发布与监测。

发布后监测索引、AI 引用、query 触发、用户反馈、投诉、错误摘要和异常外链。

4.2.3 GEO 效果监测原则

GEO效果监测不能只看“是否出现”。建议同时看七类指标：

1. **准确性**：AI 输出是否准确描述品牌、产品、价格、范围和限制。
2. **引用质量**：引用来源是否来自官网、权威媒体、专业机构、客户授权案例或低质站点。
3. **覆盖范围**：核心 query、长尾 query、竞品比较 query、问题型 query、购买型 query 的覆盖。
4. **一致性**：不同语言、地区、模型和时间段的答案是否稳定。
5. **公平性**：竞品比较是否存在无证据负面结论或来源偏差。
6. **风险性**：是否出现幻觉、过期信息、伪评价、注入痕迹、恶意链接或敏感信息。
7. **可修复性**：错误能否通过官网更正、来源更新、平台反馈和知识库修复得到改善。

4.2.4 来源质量评分

建议把来源分为五级：

- A 级：官网事实页、法务审校文档、权威标准、政府或监管资料、经授权客户案例。
- B 级：可信媒体报道、学术论文、专业机构研究、透明方法论的第三方评测。
- C 级：行业博客、社区讨论、论坛回答、未完全披露方法的数据报告。
- D 级：匿名内容、低质量目录、站群、软文、付费榜单、缺少披露的评测。
- E 级：伪造来源、垃圾站、恶意搜索投毒、钓鱼站、被黑页面、隐藏指令页面。

AI 搜索监测中应记录每次引用来源等级。

若 AI 频繁引用 D 或 E 级来源，应优先修复来源生态，而非简单增加内容数量。

4.2.5 安全测试边界

授权红队和安全测试可以帮助发现提示词注入、RAG 投毒和 agent 风险，但需要明确边界：

1. 只在授权系统、授权账号、授权数据范围内测试。
2. 不对第三方生产系统投放攻击内容。
3. 不发布可被开放搜索引擎抓取的恶意指令。
4. 不尝试获取真实用户数据、商业秘密或凭据。
5. 不把漏洞细节交给无关人员或营销团队。
6. 测试结束后清理测试数据，提交修复建议和回归验证。

4.2.6 对外更正机制

一旦发现已发布内容造成 AI 搜索错误引用，应采取以下动作：

1. 在源页面更正事实，标注更新日期和更正说明。
2. 移除或修正结构化标记、常见问题解答、图片 alt 文本和 PDF。
3. 向搜索引擎或平台提交重新抓取、移除或反馈请求。
4. 通知使用该内容的媒体、合作方和客户。
5. 对企业RAG知识库重新索引，并验证相关检索词。
6. 在必要时发布对外说明，避免错误继续扩散。

4.3 GEO 效果指标体系

4.3.1 指标体系的基础认知

企业采购GEO服务前，需要先把GEO放在正确的认知框架里。GEO的核心对象是生成式答案生态，平台会根据问题、时间、入口、联网状态、检索结果、模型版本、上下文和安全策略生成不同答案。

企业能做的是提升自身内容被理解、被引用、被合理推荐的概率，降低错误描述、缺失引用、竞品挤压和合规风险的概率。这意味着企业在评估GEO时，应使用**概率思维、样本思维和证据思维**。

一个专业服务商不应承诺“固定排名”，而应说明如何构造问题样本、如何控制采样条件、如何定义品牌可见、如何计算排位次序（如Top 1和Top 3）、如何解释波动、如何处理样本量不足带来的误差。

GEO本质上更接近概率优化。企业无法要求某个平台在任意时间、任意账号、任意问题下都固定推荐某个品牌，但可以要求服务商在统一采样口径下，提升目标品牌在**目标问题集合中的可见概率、排位次序进入概率、引用概率和事实准确概率**。

GEO指标体系建议先分成两大类：**前端指标和后端指标**。

- 前端指标衡量AI答案本身的表现，包括目标品牌是否出现、出现在哪里、是否进入推荐排位、是否被引用、事实是否准确。
- 后端指标衡量业务链路中的外部变化，包括访问、线索、商机、成交、销售反馈和品牌搜索等。

前端指标还需要继续拆成两类：**前端可量化指标和前端非量化指标**。

- 前端可量化指标可以基于采集样本计算概率或比例，适合作为项目考核和验收的主要依据。
- 前端非量化指标带有明显主观判断，例如语义一致性、情感积极性、表达友好度、推荐理由质量等，可以关注和复盘，但不建议作为服务商硬性KPI。

4.3.2 指标设计原则

可复核：企业必须能看到原始样本、采样时间、平台入口、账号状态、联网状态、答案文本、引用来源和截图或日志。

可计算：用于考核的指标必须能写出清晰公式，分母、分子、有效样本、排除样本都要明确。

可解释：指标变化要能被业务、内容、技术和管理团队理解，不能只给一个黑箱分数。

可行动：指标需要能反向指导页面、内容、结构化数据、实体资料、第三方权威来源和合规审查。

可比较：同一周期、同一平台、同一问题库、同一采样条件下的结果才适合做前后对比。

有边界：非量化指标可以作为判断线索，后端指标可以作为经营观察，但二者都不宜直接变成服务商绩效承诺。

4.3.3 前端指标

(1) 前端可量化指标：以核心指标作为考核抓手

前端可量化指标体系包含多个维度，如可见率、引用率、排位次序进入率（如Top 1和Top 3）、品牌一致性、答案准确性、竞品相对位置等。

但在企业实际采购和验收中，不宜把所有指标都纳入硬性考核。多数指标更适合作为过程指标，用于帮助服务商诊断问题、优化策略、复盘迭代。

企业应选择少数最能代表业务结果的核心指标作为考核抓手。比如在POC验收、阶段复盘和年度服务评估中，可以采用**排位次序进入率（应据企业真实竞争力、语料现状和场景范围确定具体排位次序）**作为核心指标，观察目标品牌是否稳定进入AI答案的主要推荐范围。

其他过程指标可以作为辅助参考，用于解释排位次序进入率变化背后的原因，包括平台差异、问题包差异、内容资产质量、引用源分布和竞品表现等。这样既能保证考核口径清晰，也能保留策略优化所需的分析空间。

前端可量化指标一览：

前端指标	推荐计算口径	适合用途	注意事项
目标品牌可见率	提及目标品牌的有效样本数 / 有效样本总数	衡量目标品牌是否进入AI答案	需要明确品牌全称、简称、英文名、产品名是否计入
排位次序进入率（如Top 1、Top 3、Top10）	目标品牌位于特定排位次序的有效样本数 / 有效样本总数（据企业竞争力、语料现状和场景范围确定具体排位次序）	衡量目标品牌在答案中的认知和推荐强度	只适用于有明确推荐排位场景，如采购推荐、品牌对比、品类榜单类问题
首提概率	目标品牌在答案中第一个被提及的样本数 / 有效样本总数	衡量答案显著性	当答案没有明确排序时，首提比Top 1更稳妥
平均出现位置	目标品牌出现位置序号的均值或中位数	观察目标品牌位置变化	需要统一编码规则，未出现样本应单独处理
自有域引用率	引用企业官网、文档、白皮书、帮助中心等自有域的样本数 / 有效样本总数	衡量企业自有资料被采用的概率	平台不展示引用时，应记录为无引用或不可判定
有效引用率	引用来源能够支撑答案相关结论的样本数 / 有效样本总数	判断引用是否真正有用	不能只看有没有引用，还要看引用是否支持结论
第三方权威引用率	引用行业机构、监管机构、学术、标准组织、权威媒体等来源的样本数 / 有效样本总数	衡量外部证据结构	需要先定义权威来源白名单或分级规则
事实错误率	出现企业事实错误的样本数 / 有效样本总数	控制错误认知与合规风险	错误要分级，例如轻微过期、核心事实错误、敏感错误
竞品共现率	同一答案中目标品牌与指定竞品共同出现的样本数 / 有效样本总数	分析竞争环境	共现本身不一定是风险，关键看位置、理由和描述

前端指标	推荐计算口径	适合用途	注意事项
相对优势概率	目标品牌位置或推荐理由优于指定竞品的样本数 / 有效样本总数	判断目标品牌相对表现	需要事先定义“优于”的编码规则
样本稳定率	多次采样中目标品牌出现状态保持一致的样本组数 / 样本组总数	判断平台波动和结果可信度	稳定率低时，不宜用单次结果做结论
问题覆盖率	已纳入监测的问题数 / 目标问题库问题数	判断监测范围是否完整	不能用覆盖率代替效果指标
场景覆盖率	已覆盖场景数 / 目标场景数	避免只测品牌词	目标场景应包含品牌、品类、采购、竞品、风险、售后等

(2) 前端非量化指标：适合关注，不适合作为硬考核

有些指标对企业很重要，但很难形成真正客观的量化结果。

原因在于这些指标需要主观定义，且不同评审人、不同业务部门和不同语境下的判断可能不同。企业可以把它们作为人工评审项、风险标签或复盘线索，但不建议把它们包装成精确分数。

典型的非量化指标包括语义一致性、情感积极性、表达友好度、推荐理由质量、品牌调性匹配、答案完整性、专业性、说服力、风险表述是否恰当等。它们可以帮助企业发现问题，但如果没有严格标注规则、多人一致性检验和样本审计，就容易变成“伪量化”。

企业可以把非量化指标改成三级标签：

A级：表达准确、语义一致、无明显风险，可作为正向样本。

B级：总体可接受，但存在表述模糊、信息不足或推荐理由弱的问题。

C级：存在误导、夸大、偏见、事实错位或合规风险，需要处理。

前端非量化指标一览：

前端指标	可以观察什么	推荐使用方式	不建议做法
语义一致性	AI对企业定位、产品边界和价值主张的理解是否贴近企业事实	人工评审、典型样本分析、错误归因	直接给出一个“语义一致性提升37%”的硬承诺
情感积极性	答案是否偏正向、中性、谨慎、负向或带风险提示	用作声誉风险标签和舆情线索	把积极、中性、负向简单换算成项目KPI
推荐理由质量	AI推荐目标品牌时的理由是否具体、准确、有证据	结合引用来源和事实准确性做案例复盘	只看推荐理由字数或形容词数量
品牌调性匹配	答案表达是否符合企业希望传递的专业、可信、创新等形象	用于内容策略和品牌资料改写	要求服务商保证AI使用指定话术
答案完整性	回答是否覆盖核心产品、场景、适用边界和限制条件	用于识别资料缺口	把完整性当成跨平台统一分数
专业性	答案是否使用正确行业术语，是否避免误导	用专家复核或行业顾问审核	让非专业评审直接判定复杂技术内容
风险表述合理性	AI是否把合规、隐私、安全、医疗、金融等风险讲清楚	用于高敏行业审查	要求AI删除所有风险提示

4.3.4 后端指标

后端指标可以帮助企业理解GEO与商业结果之间的关系，但多数情况下不适合作为服务商直接考核指标。

主要原因包括：AI平台通常不提供完整查询日志，用户从AI答案到官网、销售或成交的路径不完整，AI回答会受到时间、账号、联网状态和平台策略影响，同期广告、PR、SEO、销售活动也会干扰结果。

因此，后端指标更适合企业内部长期观察。

企业可以在CRM、网站分析、品牌监测和销售复盘中标注GEO相关线索，但不宜要求服务商承诺短期线索数、成交额或销售转化。

GEO后端指标包括官网流量、品牌词搜索、咨询量、线索量、下载量、销售反馈和客户问询变化。**它们对企业有业务价值，但多数情况下不适合直接写成服务商KPI。**

后端指标观察四步法建议：

- 1. 第一步，建立基线。记录GEO项目启动前至少4周的品牌词搜索、官网入口、咨询表单、销售问询和下载数据。
- 2. 第二步，设置观察窗口。一般按月观察，至少连续3个月。不要用单日数据判断GEO效果。
- 3. 第三步，记录干扰变量。包括广告投放、媒体发布、展会、促销、产品更新、舆情事件和竞品活动。
- 4. 第四步，做相关性解释。若前端可见率提升、品牌词搜索提升、销售反馈中出现“我在AI里看到你们”的样本，可以作为业务信号，但仍需标注为观察结论。

后端指标一览

指标类别	可观察指标	适合用途	不适合作为硬考核的原因
流量	AI平台推荐流量、搜索控制台Web流量、品牌词搜索量、直接访问变化	观察AI与搜索生态的外部变化	很多AI答案不会带来点击，平台也不一定提供完整来源
行为	页面停留、转化路径、下载、注册、咨询、收藏	观察内容是否被更高意图用户使用	行为变化可能受页面改版、广告和活动影响
线索	MQL、SQL、表单、电话、企微添加、试用申请	辅助判断商业价值	需要CRM字段、人工标注和销售回访支持
销售	商机、报价、成交、客单价、回款周期	长期预算评估和经营复盘	通常无法直接归因到GEO单一动作
品牌	品牌搜索、社媒提及、媒体引用、招投标问询、销售反馈	评估认知变化和销售支持	需要多渠道数据交叉验证

4.3.5 指标观测仪表盘结构建议

建议企业把 GEO 效果指标观测仪表盘分成四层，避免只看一个总分。

第一层：管理层摘要

管理层只需要看到5个指标：目标品牌可见率、目标品牌Top 3概率、有效引用率、重大事实错误率、红色风险数量。每月看趋势，不建议每日追波动。

第二层：平台对比

分别展示豆包、DeepSeek、Kimi、元宝、千问的表现。每个平台都应显示样本数、采样日期、入口、联网状态和问题类别。平台之间不能直接用一个分数粗暴排名，因为每个平台的内容来源、搜索机制和回答风格不同。

第三层：场景对比

把问题拆成品类教育、采购决策、竞品对比、风险判断、售后服务和品牌认知。很多企业会发现，品类教育类问题表现还可以，但采购决策和竞品对比类问题很弱。这个发现比一个总分更有行动价值。

第四层：样本追溯

每一项指标都应能回到原始回答。企业至少要能查看问题、平台、时间、回答文本、品牌位置、引用链接、竞品出现情况、人工复核标记和处理状态。没有样本追溯的看板，很难用于验收。

字段	说明	是否必需
问题ID	每个问题固定编号	必需
问题文本	原始提问，不随意改写	必需
问题类别	品类、采购、竞品、风险、售后等	必需
平台	豆包、DeepSeek、Kimi、元宝、千问	必需
入口	网页、App、插件、API或搜索增强入口	必需
采样时间	精确到日期，必要时到小时	必需
联网状态	是否开启搜索、深度研究或联网功能	必需
目标品牌是否出现	是或否	必需
目标品牌位置	Top 1、Top 2、Top 3、其他、未出现	必需
引用来源	自有域、第三方、平台生态、无引用	必需
事实错误	无、轻微、重大	必需
人工复核人	责任人	建议
处理建议	内容、技术、合规或观察	建议

五

GEO 服务采购与合规评估

5.1 GEO服务商类型

5.1.1 中国GEO服务商类型划分

中国市场的GEO服务商正在从多个方向涌现。企业采购时需要先识别服务商的能力来源，再判断其能力结构是否匹配自身目标。类型只能用于初筛，最终仍要看可复核的方法、样本、团队和交付证据。

类型	典型背景	核心能力	适合采购场景	常见短板	重点判断
AI原生 GEO 服务商	AI产品、模型 评测、智能 体、提示词工 程、数据评测 团队	理解模型行为、 检索增强、问答 生成、采样波 动、评测框架和 自动化监测	需要建立GEO 指标体系、平 台采样、评测 看板和实验闭 环的企业	可能缺少 SEO、品牌 叙事、PR资 源和企业内 容落地经验	能否把AI评测转 化为可上线的内 容、页面和组织 动作
AI原生 +SEO +营销 复合型	AI native团队 叠加技术 SEO、内容营 销、品牌传 播、数据分析 背景	既懂AI平台与评 测，又懂抓取索 引、内容资产、 品牌表达、竞品 分析和项目治理	中型企业、行 业龙头、上市 公司、出海企 业的一期主服 务商	团队稀缺， 报价通常更 高，也容易 被包装成概 念	是否有真实跨职 能团队，是否能 拿出采样数据、 内容改造和技术 建议的闭环证据
SEO转 型型	传统SEO、技 术SEO、搜索 营销公司	懂抓取、索引、 页面结构、关键 词、内链和搜索 数据	官网基础薄 弱、页面不可 索引、结构化 数据缺失的企 业	可能把GEO 简单套用 SEO排名逻 辑	是否理解AI答案 的概率属性、引 用机制和平台差 异
内容 营 销/PR 型	品牌、公关、 内容营销、媒 体传播公司	懂品牌表达、叙 事、案例、媒体 关系和内容生产	需要改善品牌 叙事、案例表 达和权威资料 的企业	技术和数据 监测能力可 能不足	是否能提供可复 核数据，而非只 交付文章和传播 稿
数据 监 测/Sa aS型	舆情、品牌监 测、数据看 板、BI工具公 司	擅长监测、告 警、竞品对比、 报表和权限管理	需要长期监 测、风险告警 和多品牌管理 的中大型企业	可能缺少内 容改造、技 术修复和业 务策略能力	看板能否解释行 动方案，数据能 否导出和复核
技术 SEO/ 数据 工程 型	数据咨询、爬 虫、BI、网站 技术服务商	擅长采样、数据 管道、仪表盘、 结构化数据、日 志和实验设计	集团、上市公 司、技术团队 成熟企业	品牌内容与 业务策略能 力可能不足	是否能和市场、 品牌、法务团队 协同，而非只交 付技术清单

类型	典型背景	核心能力	适合采购场景	常见短板	重点判断
大模型/RAG/Agent技术型	AI应用开发、知识库、RAG、智能体和企业AI供应商	懂内部知识库、检索、模型评测、Agent流程和企业AI应用	需要把内部知识库治理和外部GEO结合的企业	未必熟悉公共AI平台的外部可见性	能否区分内部RAG效果和公共AI平台GEO效果
行业垂直咨询型	医疗、金融、教育、工业、B2B咨询公司	懂行业术语、客户决策链、监管要求和专业内容	高敏行业、复杂B2B、专业服务企业	通用平台监测和自动化能力可能较弱	是否具备标准化数据方法和可复核样本
跨境增长型	出海营销、跨境SEO、海外PR、独立站、本地化服务商	熟悉语言、本地化、海外搜索、海外内容生态和跨境增长	出海企业、海外B2B、跨境电商、全球化品牌	国内AI平台经验可能不足	是否能分别说明国内平台和海外平台的策略差异
资源铺量型	软文发布、媒体资源、外链、口碑和发布渠道服务商	短期可以大量发布内容	一般不建议作为主服务商	质量、合规和长期风险高	虚假评价、低质内容、隐性投放、不可验收

5.1.2 主服务商与组合采购

GEO往往需要组合能力。企业可以选择一个主服务商负责总体策略、数据监测和项目治理，再按需引入内容、技术、PR、法务或行业专家。对于中大型企业，单一服务商很难在平台监测、技术SEO、品牌内容、行业专业、数据安全和合规审查上全部达到高水平。

建议根据企业自身需求方向与行业特点选择主服务商。主服务商需要能够用AI原生能力建立评测体系，用专业策略和技术能力修复抓取、索引、结构和实体问题，用内容能力和传播资源完成内容叙事、案例表达和权威证据建设。

5.1.3 海外GEO与国内GEO采购差异

中国大陆企业常见两类需求：面向海外市场采购GEO服务，或面向国内AI平台采购GEO服务。两者的服务商能力、合规边界、内容资产和验收方式有明显差异。

维度	海外GEO采购	国内GEO采购
平台范围	目标市场的搜索AI、国际AI助手、海外搜索引擎和专业问答平台	豆包、DeepSeek、Kimi、元宝、千问，以及国内搜索和内容生态
内容语言	本地语言、英文、地区化表达、文化语境	中文表达、行业术语、政企和B2B采购语境
资料来源	官网英文站、海外媒体、分析师报告、行业协会、开发者文档、评价站点	官网中文站、公众号、白皮书、行业媒体、招投标资料、百科和平台生态内容
技术基础	国际SEO、结构化数据、hreflang、独立站、海外服务器和加载体验	中文站结构、可抓取性、公众号和多平台内容资产、备案与合规要求
合规重点	目标市场隐私、广告、消费者保护、行业监管、跨境数据	个人信息保护、数据安全、生成式AI服务管理、网络安全和行业监管
服务商能力	本地化、海外SEO、PR、分析师关系、国际内容证据	中文平台采样、国内内容生态、合规审查、品牌与舆情治理
验收重点	目标国家或地区的AI答案表现、本地权威来源和语言准确性	五个平台的答案表现、引用来源、事实准确性和敏感风险
主要风险	语言和文化误读、海外低质PR、跨境数据和本地法规不熟悉	平台口径变化、低质内容堆量、隐性投放、敏感内容违规

海外GEO采购建议

- 先明确目标市场。美国、欧洲、东南亚、中东和日本市场的内容生态、监管和采购习惯不同。
- 服务商需要证明本地化能力。只会翻译中文内容通常不足以支撑海外GEO。
- 海外项目要重视结构化数据、独立站技术、英文文档、行业报告、第三方评价和媒体关系。
- 出海B2B企业应重点建设英文产品文档、案例、FAQ、数据安全说明、集成文档和行业解决方案。
- 企业应要求服务商说明目标市场的数据合规和广告合规边界，并单列外部资源费用。

国内GEO采购建议

- 国内项目要按平台拆解，不能用一个总分掩盖豆包、DeepSeek、Kimi、元宝、千问之间的差异。

- 国内项目更需要关注内容合规、行业语境、政企采购语境、B2B售前语境和内容生态一致性等方面。
- 企业要整理官网、公众号、白皮书、产品手册、帮助中心、案例和资质证明，形成统一事实库。
- 涉及腾讯、阿里、字节等不同生态的内容资产时，要保持事实口径一致，避免平台之间答案互相冲突。
- 服务商必须理解中国的数据安全、个人信息保护、生成式AI服务和网络安全相关要求。

海外与国内双线项目

部分企业同时做国内与海外GEO。建议采用“统一事实库，分市场表达，分平台监测”的治理方式。总部维护核心事实，国内团队管理中文平台与中文内容，海外团队管理目标市场平台、本地语言和本地权威来源。

- **统一事实库**：公司介绍、产品线、资质、客户案例、部署方式、价格边界、数据安全说明。
- **分市场表达**：国内和海外的行业术语、合规声明、案例呈现和购买决策语言可以不同。
- **分平台监测**：国内五个平台单独监测，海外目标平台单独监测，汇总时只比较同类指标。
- **统一风险管理**：客户案例授权、个人信息、未公开数据、财务信息和敏感行业内容由总部规则约束。

5.1.4 不同企业阶段与行业敏感性的采购策略

小型企业受预算与资源约束，GEO 项目聚焦少量高价值场景，以夯实基础信息露出、修正内容谬误、完善核心页面问答能力、搭建低成本监测体系为目标，完成数十条核心问题基线摸排与官网关键页面优化，暂缓采购高价工具与全周期服务，落地具象化页面整改清单，围绕品牌提及、信息精准度等维度验收。

中型企业依托自有官网、内容及销售体系落地 90 天标准化 GEO 项目，将相关工作融入品牌、内容与增长常态化流程，搭建中量级问题库并对标重点竞品与核心平台，按月复盘监测，联动白皮书、销售物料、案例等内容协同优化，辅以转化数据做参考，不以转化作为唯一验收标准。

行业龙头GEO 侧重竞品防御、品牌权威塑造与全平台信息统一，防范 AI 不实内容损耗品牌与商务合作，搭建全维度海量问题库并常态化追踪同业及新晋玩家，搭建官方权威资料库，联动多部门协同落地项目，从信息采信质量、行业竞争位次、内容风控等维度考核成果。

上市公司GEO 应严守信息披露规范与公开数据边界，财务、业绩、战略合作等关键内容全部依托公示资料编撰，嵌入信披、投关人员参与项目全流程，搭建内容错误溯源纠偏制度并在合作合同中约束供应商严禁杜撰未披露信息，重点核查内容差错率、合规风险与资料匹配度。

高敏行业（医疗、金融、法务、教培等）合规与专业审核应贯穿 GEO 全流程，所有输出内容经内部专家、法务复核，AI 内容不直接充当专业指导意见，严控隐私、资质、疗效承诺类内容发布，服务商配套敏感内容管控与应急处置方案，以合规通过率、问题整改效率作为核心验收依据。

出海企业应跳出直译误区，兼顾国内外品牌认知建设与本地化公信力打造，立足目标市场搭建本土问题库、对标本地竞品，完善多语种合规文档与产品素材，统一海内外站点产品、定价等信息口径，核验服务商本地资源资质，依托海外平台曝光、内容本地化质量与有效询盘完成验收。

5.2 服务商评估维度与评分框架

5.2.1 总评分建议

企业可以用100分制评估GEO服务商。评分不应只看方案漂亮程度，更要看方法是否可复核、AI原生能力是否真实、内容与营销落地是否完整、合规是否稳妥、团队是否能长期协作。

维度	建议权重	评估要点
测量与评估方法	20分	问题库、重复采样、平台记录、原始样本、指标定义、Top 1与Top 3口径、复测方法、波动处理
AI原生与平台理解	15分	模型行为、检索增强、搜索/联网模式、答案生成、采样自动化、异常样本解释、五个平台差异
内容与营销复合落地	15分	抓取索引、结构化数据、页面结构、FAQ、产品文档、品牌叙事、案例、权威证据
数据与监控能力	15分	可视化、日志保存、竞品对比、异常告警、数据导出、权限管理、长期趋势
合规与安全	15分	数据处理、个人信息、客户案例授权、敏感行业审查、内容审核、合同禁止行为
项目管理	10分	节奏、周会、月报、协同工具、变更管理、验收和交接
商业透明度	10分	报价清晰、资源关系披露、不可承诺事项、退出机制、案例真实性

5.2.2 不同企业的权重调整

企业类型	权重调整建议
小型企业	提高内容与实体治理、项目性价比和短期可执行动作权重；降低复杂数据看板权重
中型企业	保持测量、AI原生、SEO、内容、项目管理均衡；重点看90天内可落地结果
行业龙头	提高数据与监控、竞品对比、平台全覆盖、AI原生评测和组织协同权重
上市公司	提高合规、安全、事实准确、信息披露一致性、供应商审计和数据权属权重
高敏行业	将合规与安全提升到最高优先级，并要求专家复核流程
出海企业	提高海外平台、本地化、语言质量、目标市场权威来源和跨境数据合规权重

5.2.3 报价评估

GEO报价通常由诊断、监测、内容、技术、数据看板、项目管理和第三方资源等部分构成。

企业评估报价时，应把重点放在两个层面：

第一，**结果考核以核心可量化指标为主**。企业不宜把所有指标都纳入硬性考核，应优先选择最能代表项目成效的核心指标，例如目标品牌Top3概率、目标品牌可见率、引用率、事实错误率等。过程指标可以用于策略复盘和执行优化，但不宜全部变成付款或验收条件。

第二，**过程方法论要重点评估内在含金量和专业水平**。GEO服务的价值，不能简单用执行动作数量或人天投入衡量。不同服务商可能采用不同策略、不同路径和不同组织方式，真正需要评估的是其底层方法、运营思路、判断原则、团队能力、行业理解和持续迭代能力。

企业不应只比较总价，而应要求服务商说明报价背后的服务逻辑，包括核心方法论、项目推进机制、数据评估原则、内容策略框架、技术优化思路、团队配置方式、复盘机制，以及这些能力如何支撑核心指标提升。

因此，报价评估不宜简单看“便宜还是贵”，而应综合判断：**核心指标是否可考核，方法论是否有含金量，运营思路是否清晰，团队能力是否匹配，交付机制是否能支撑长期提升。**

企业评估报价时，建议把总价拆成六类成本。

成本项	应包含内容	风险提示
诊断成本	问题库、基线采样、竞品分析、资料诊断	只给结论不给样本，价值较低
数据成本	采样工具、人工复核、看板、导出数据	需明确数据归属和导出权限
内容成本	FAQ、事实页、案例页、白皮书、结构化摘要	需明确原创、审核和修改次数
技术成本	抓取、索引、结构化数据、站点地图建议	需确认是否真正有技术执行能力
项目管理成本	周会、月报、复盘、跨部门协同	低价项目常低估这一部分
合规成本	法务复核、敏感内容处理、授权材料	高敏行业不可省略

5.3 如何判断服务商的专业水平

5.3.1 专业服务商应具备的能力

- 能解释GEO与SEO、PR、内容营销、舆情监测、知识库治理之间的关系与边界。
- 具备真实AI原生能力，能理解模型行为、检索增强、问答生成、采样波动和评测框架。
- 能按平台拆解豆包、DeepSeek、Kimi、元宝、千问的采样口径和产品入口。
- 能建立问题库，并说明每类问题对应的业务目的。
- 能把前端可量化指标、前端非量化观察和后端观察区分清楚。
- 能用原始样本和可复核数据解释目标品牌可见率、Top 1概率、Top 3概率、引用率和错误率。
- 能给出内容、技术、实体、权威来源、数据监测和合规六类行动建议。
- 能把SEO、内容营销、实体治理、AI评测和项目治理形成闭环。
- 能承认平台波动和不可承诺事项，并把风险写进方案和合同。
- 能把企业内部资源需求讲清楚，例如业务事实、官网权限、客户案例授权和法务审批。

5.3.2 专业能力红线

- 声称可以保证某AI平台固定推荐某品牌。
- 用一次截图或少量样本证明长期效果。
- 拒绝提供采样方法、原始样本或指标定义。

- 把Top 1、Top 3、语义一致性、情感积极性等指标混在一起，无法说明哪些可以量化、哪些只能观察。
- 把后端成交承诺作为主要卖点，却无法说明实验设计。
- 推荐批量发布低质内容、伪造第三方评价、购买隐性背书。
- 对个人信息、客户案例授权、敏感行业合规没有流程。
- 把所有平台当作同一种搜索引擎处理。
- 把内部RAG或企业智能体优化的效果直接当成公共AI平台GEO效果。
- 合同里没有数据权属、交付物、验收标准和退出机制。

5.4 企业采购GEO服务的完整用户旅程

GEO采购不应从询价开始。更稳妥的流程是先完成内部目标和基线，再进入服务商识别与比稿。否则企业很容易用含糊目标采购含糊服务，最终只能比较PPT、口才和报价。

阶段	核心目标	企业关注重点	可交付证据	主要风险
1. 内部立项	明确为什么做GEO	业务目标、预算、平台范围、内部责任人	立项说明、目标清单、平台范围	把GEO当作短期流量项目
2. 需求诊断	定义问题与场景	品牌、品类、产品、竞品、区域、用户问题	问题地图、竞品清单、场景矩阵	只围绕企业自说自话的关键词
3. 基线测量	建立现状数据	提问样本、平台版本、采集方法、重复采样	基线报告、原始样本、截图或日志	单次截图、样本过小、不可复现
4. 服务商识别	找到候选供应商	类型、能力、案例、行业经验、数据方法	候选名单、初筛评分	被低价、噱头或所谓资源承诺吸引
5. RFI/RF	让供应商按统一口径回复	服务范围、指标、方法、团队、合规、报价	RFI回复、RFP方案、报价表	各家口径不同导致无法比较
6. 比稿与尽调	判断专业能力和适配度	方法论、数据可信度、交付团队、边界意识	评分表、访谈记录、尽调问题答案	只看案例截图或老板背书
7. 合同与SOW	把承诺写清楚	交付物、验收标准、数据权属、禁止行为	合同、SOW、验收表	合同只写服务名称和总价
8. 执行治理	保持项目节奏	周会、月报、变更、复核、合规审查	任务看板、月报、采样记录	内容上线慢、跨部门协同失灵
9. 验收复盘	判断效果与复购价值	前端指标、后端观察、学习沉淀、下一周期	验收报告、复盘会议纪要、续约建议	只用好看的样本证明效果
10. 续约或退出	决定长期机制	平台覆盖、内部能力沉淀、成本收益	续约方案或交接清单	供应商锁定、数据无法迁移

5.4.1 阶段一：内部立项

内部立项阶段要先回答三个问题：企业为什么现在需要GEO，优先服务哪个业务单元，成功标准是什么。对于大多数企业，GEO可以先从品牌认知、核心产品认知、竞品对比、招商采购、售前问答、海外市场进入等场景切入。

- 立项输出一：GEO业务目标。示例：提升企业在“行业解决方案推荐”类问题中的品牌提及率与引用率。

- 立项输出二：平台范围。示例：豆包、DeepSeek、Kimi、元宝、千问作为一期范围。
- 立项输出三：预算与周期。示例：30天诊断，90天一期执行，半年复盘。
- 立项输出四：内部责任人。市场、品牌、内容、官网、法务、产品、销售至少需要明确接口人。
- 立项输出五：合规边界。明确哪些材料可以公开、哪些材料需要审批、哪些内容不得用于AI答案优化。

5.4.2 阶段二：需求诊断

需求诊断需要把“我们想提升GEO”转化为可测量的问题库。问题库不宜只由品牌词组成，还要覆盖品类词、需求词、竞品词、购买决策词、风险词和区域词。

问题类型	示例问题	用于判断什么
品牌认知	某品牌是做什么的？有哪些核心产品？	AI是否理解企业实体和业务边界
品类推荐	国内适合制造企业的某类解决方案有哪些？	企业是否进入品类答案
采购决策	企业采购某类系统应如何选择服务商？	企业是否被列入采购候选
竞品对比	A公司和B公司在某方面有什么差异？	AI是否准确描述优势与限制
场景方案	某行业如何解决某个业务痛点？	企业内容是否能支撑场景答案
风险合规	某产品是否涉及数据安全或隐私风险？	敏感问题是否被准确处理
区域与语言	海外某地区有哪些适合中国企业的服务商？	跨境和本地化内容是否被识别
售后与支持	某品牌支持哪些部署方式或服务方式？	官网和帮助中心内容是否被引用

5.4.3 阶段三：基线测量

基线测量是采购GEO服务的关键节点。没有基线，企业无法判断服务商方案是否精准，也无法在后续验收时确认变化来自项目行动、平台波动还是采样偏差。

- 基线应覆盖至少三类问题：品牌词问题、非品牌品类问题、竞品或对比问题。
- 每个平台建议对同一问题进行多次采样，记录日期、时间、账号状态、设备、地区、联网状态、是否开启深度搜索或搜索模式。

- 每条样本要记录答案文本、品牌提及、提及位置、引用来源、来源域名、情绪倾向、错误点、竞品共现和截图或可追溯日志。
- 基线报告应包含问题库结构、采样方法、指标定义、主要发现、优先级建议和风险说明。
- 企业应要求服务商提供原始样本或可复核导出，而非只看汇总图。

5.4.4 阶段四至六：服务商识别、RFI/RFP与比稿

企业在比稿时应让服务商使用同一套输入材料和同一套问题库回复方案。不同供应商的PPT风格、话术和报价口径差异很大，统一输入可以让比较更接近真实能力比较。由于当前阶段很多标准暂未统一，也可以由服务商自行提供方案与考核标准，企业内部多视角衡量，选择最适合的自己亦可。

- RFI（信息征询）阶段主要问能力边界、过往经验、团队配置、数据方法和合规政策。
- RFP（建议征询）阶段要求供应商提交完整项目方案，包括基线方法、平台覆盖、指标体系、内容策略、技术策略、治理节奏、POC设计、风险控制和报价。
- 比稿阶段应要求供应商现场解释采样口径、指标定义、平台差异、后端归因边界和不可承诺事项。
- 尽调阶段应核验服务商案例是否可复核，是否涉及夸大、伪造、隐性投放或违反平台规则的做法。

5.4.5 阶段七至十：合同、执行、复盘与退出

GEO项目的长期价值来自持续治理。企业应把合同写成可验收、可复盘、可退出的结构，而非只采购一批内容或一套监测截图。

- 合同中应写清楚交付物：问题库、基线报告、监测看板、内容方案、技术修复建议、月报、复盘和交接资料。
- 合同中应写清楚验收标准：指标定义、样本范围、平台范围、采样频率、数据保存、复测方式。
- 合同中应写清楚禁止行为：虚假内容、伪造案例、隐性付费评价、批量垃圾页面、未经授权抓取个人信息、绕过平台规则、误导性承诺。
- 执行期应设定周会和月度复盘。周会解决任务进度，月度复盘看指标变化、平台差异和风险。
- 退出时应要求供应商交付问题库、样本数据、资产清单、内容修改记录、技术建议、监测模板和未完成事项。

5.5 采购 GEO 服务的参考策略

5.5.1 内部立项会统一认知

企业需要在立项会上进行充分讨论，统一认知，在以下六个问题上达成共识意见，避免目标不一致导致资源浪费与流程反复。

- 我们希望优化的是品牌被提及、被首位推荐、被引用，还是被准确描述。
- 我们希望覆盖的是豆包、DeepSeek、Kimi、元宝、千问中的哪些关键词或用户意图问题。
- 我们的问题样本是否覆盖真实用户，而非只覆盖企业内部习惯用语。
- 我们的自有内容是否足以支撑AI给出准确答案。
- 我们是否能接受短期波动，并用月度趋势判断效果。
- 我们是否把后端线索、成交和销售数据作为观察，不直接归因给服务商。

错误立项参考案例：

某精密设备企业主营工业视觉检测设备。市场部最初提出的采购目标是“三个月内让AI每次都把我们列为第一推荐”。这个目标在采购会上被服务商销售包装成“全平台Top 1保障”。

经过内部复盘后，该企业把目标改成了三类：

第一，围绕“汽车零部件视觉检测方案”“锂电池缺陷检测设备”等40个高价值问题，目标品牌Top 3概率从基线的15%提升到30%；

第二，企业自有域引用率从5%提升到15%；

第三，重大事实错误率降到5%以内。这个目标更适合写入SOW（工作说明书），因为它有样本、有平台、有时间窗口，也承认GEO是概率优化。

这个例子说明，采购目标一旦从绝对承诺改成可测量概率，企业就能更清楚地判断服务商是否专业，也能减少后期验收纠纷。

5.5.2 需求诊断内部工作坊

需求诊断需要把“我们想提升GEO”转化为可测量的问题库。问题库不宜只由品牌词组成，还要覆盖品类词、需求词、竞品词、购买决策词、风险词和区域词。

清单一：业务场景清单

汇总不同用户问询场景，结合企业营收转化、品牌口碑、舆情风控权重由高到低排序，锚定 AI 落地优先级，精准锁定高价值智能问答落地方向。

清单二：问题样本清单

对应六大业务场景，每个场景落地不少于 10 条贴合网民口语习惯的真实问询，兼顾专业叫法与生活化提问，复刻终端用户原生提问逻辑。

清单三：企业事实清单

逐条绑定样本问题，归集产品参数、资质证书、落地案例、收费标准等企业一手资料；缺少客观资料佐证的提问，剔除 AI 定制优化诉求。

清单四：风险边界清单

从广告合规、法务风控、营销承诺、财务宣传四个维度，划定 AI 输出禁区，列明严禁口头承诺、夸大宣传、违规兜底、不实暗示的全部内容。

5.5.3 主流 AI 平台打包采购策略

企业开展 GEO 采购，建议将豆包、DeepSeek、Kimi、元宝、千问 5 个平台打包纳入同一采购包，作为国内 GEO 项目基准评测组合，不建议单点采购单一 AI 平台服务。

不同 AI 平台的模型能力、搜索源、联网方式、内容生态、上下文能力、产品入口和用户使用场景存在明显差异，同一企业、同一产品、同一问题，在不同平台上的答案表现可能完全不同，如果企业只采购单个平台的 GEO 服务，很容易得到片面的判断。

一个服务商如果只能优化某一个平台，却无法解释五个平台之间的差异，也无法针对不同平台制定采样、内容和优化策略，说明其综合能力仍需谨慎评估。

多个 AI 平台一起采购有三层价值：

第一，可以考察 GEO 服务商的综合实力

真正成熟的 GEO 服务商，应能同时理解不同平台的答案机制、内容偏好、引用来源和用户场景，而不能只给一个笼统的“AI 可见性”分数。

第二，可以覆盖不同用户人群

豆包更偏大众用户入口，Kimi 更适合长文档和研究型场景，DeepSeek 更容易出现在推理、技术和深度问答场景，元宝受到腾讯生态影响，千问与阿里生态、办公和开发者场景关联更强。五个平台合在一起，才能更接近企业真实用户的 AI 搜索路径。

第三，可以发现平台级机会和风险

有的平台可能已经能正确理解企业，有的平台可能仍存在事实错误、竞品偏置、信源污染、引用缺失或品牌实体混淆。只有跨平台采样，企业才能判断问题到底出在内容资产、信源结构、品牌叙事，还是平台自身的答案机制。

平台功能以官方文档和实际采样为准。采购评估时，企业不应只听服务商做概括性描述，应要求其提供平台级采样数据、差异解释和对应优化策略。

最终，企业采购GEO服务时，应把五个平台视为一个基础评估组合。能不能同时理解、采样、解释和优化豆包、DeepSeek、Kimi、元宝、千问，已经成为判断GEO服务商专业能力的重要标准。

六

给行业的倡议

6.1 负责任GEO治理框架

为引导 GEO 回归基于“真实与价值”的本源，倡议各方围绕以下八个方向开展负责任 GEO 治理：

- **推行透明机制。** 积极建立和参考科学、透明、可溯源的 AI 可见性监测评估指标；媒体机构、内容平台等应清晰标注信息源与 AI 生成标识。
- **真实价值挖掘。** 任何传播与营销活动都应遵循广告、新闻相关监管规定，GEO 活动应基于真实价值的挖掘和合规手段，让优质内容、卓越品牌和优势产品被更多看见与推荐。
- **完善信源治理。** 党政机关、权威媒体、科研机构、品牌企业应加强官方网站等自有信息源建设；经济传媒机构应发挥专业优势，建设成为 AI 获取权威经济信息的核心信源；积极推动与支持大模型厂商开展信源分类分级治理。
- **构建知识语料。** 积极构建自身知识语料资源，推动高质量公共语料库建设，将优质内容转化为 AI 友好的语料知识资产。
- **权威资源链接。** 倡导与权威机构和优质可信信息源进行深度联合与绑定，从而提升自身 AI 认知可信度。
- **锚定数字信任。** GEO 实施过程中应锚定构建 AI 时代人类用户信任为核心目标，在战略层面确定以人为本、向善求真的基本原则。
- **鼓励创新融合。** 支持传媒、科技、研究机构联合开展 GEO 创新实践，开发人机友好的内容生产标准，构建垂直领域知识语料，探索人机协同深度报道，形成可复制的创新模式。
- **明确传媒责任。** 媒体机构尤其是权威主流媒体应主动践行社会责任，经济传媒机构应以生产高质量经济信息为基础，强化优质内容在 AI 系统中的能见度，提升可信度与影响力。

6.2 向各利益相关方呼吁

- AI 平台应给企业和用户提供纠错、申诉、来源审计和风险反馈机制。
- 监管层和行业协会应推动高风险行业的 AI 搜索信息披露标准。

- 投资人与管理层应识别短期 GEO 操纵带来的长期负债。
- 企业应把 GEO 用于提升事实可见度，而非制造事实幻觉。
- 服务商应把专业能力用于建设可信知识资产，而非训练规避系统的技巧。
- 媒体和社区应重视 AI 搜索时代的内容出处、利益披露和事实核验。

6.3 给社会公众的建议

提升全民 AI 素养，是筑牢数字信任的基础。为此，我们向广大公众提出以下建议：

- 源头控制。非必要不启用大模型的互联网搜索能力，避免模型在搜索中被低质信息干扰。
- 溯源核查。查看 AI 援引的源头，警惕“据内部消息”等模糊表述。
- 交叉验证。使用两个以上不同厂商的大模型对比答案。
- 指令调教。设置“请主要引用政府官网、权威媒体”等提示词门槛。
- 人本意识。在医疗、法律、投资等关键决策中，以 AI 为参考，坚持自主判断。

参考来源 References

以下资料为本报告的主要依据。正文不使用数字角标，便于阅读。

国内法律、部门规章、国家标准与官方政策

1. 《生成式人工智能服务管理暂行办法》
2. 《人工智能生成合成内容标识办法》
3. GB 45438- 2025《网络安全技术人工智能生成合成内容标识方法》
4. 《互联网信息服务深度合成管理规定》
5. 《互联网信息服务算法推荐管理规定》
6. 《网络信息内容生态治理规定》
7. 《网络反不正当竞争暂行规定》
8. 《中华人民共和国反不正当竞争法》
9. 《中华人民共和国广告法》
10. 《互联网广告管理办法》
11. 《中华人民共和国个人信息保护法》
12. 《中华人民共和国网络安全法》

AI平台/产品隐私政策

1. 豆包
2. Kimi
3. 千问
4. 元宝
5. DeepSeek

行业治理案例

1. 市场监管总局公布网络虚假宣传不正当竞争典型案例相关报道
2. 市场监管总局公布七起网络不正当竞争典型案例
3. 公安部网安局：4人用AI技术炮制网络谣言被罚
4. 中国长安网：甘肃公安通报利用AI技术炮制网络谣言案例

GEO与生成式搜索研究

1. GEO: Generative Engine Optimization, KDD 2024, arXiv:2311.09735.
2. Generative Engine Optimization: How to Dominate AI Search, arXiv:2509.08919.
3. Google Search Central: AI features and your website.
4. Google Search Central: Spam policies for Google Web Search.

AI安全、提示词注入与RAG投毒研究

1. OWASP GenAI Security Project: OWASP Top 10 for LLM Applications 2025.
2. OWASP GenAI Security Project: Prompt Injection.
3. Not what you have signed up for: Compromising Real-World LLM-Integrated Applications with Indirect Prompt Injection, arXiv:2302.12173.
4. Ignore Previous Prompt: Attack Techniques for Language Models, NeurIPS ML Safety Workshop 2022, arXiv:2211.09527.

5. PoisonedRAG: Knowledge Corruption Attacks to Retrieval-Augmented Generation, USENIX Security 2025, arXiv:2402.07867.
6. MITRE ATLAS: Adversarial Threat Landscape for Artificial-Intelligence Systems.
7. C2PA Specifications: Content provenance and authenticity specifications.

附录一：极简术语表

GEO：生成式引擎优化，目标是让真实、可信、对用户有价值的内容在生成式回答、AI 搜索引用、RAG 问答和多源摘要中被准确理解与合理呈现。

AI 搜索：通过检索、重排序、摘要、推理、引用和交互式追问提供答案的搜索形态。它可能包含网页检索、向量检索、知识图谱、结构化数据、网页实时抓取、RAG 和大模型生成。

生成式引擎：把用户问题转化为检索、推理、合成回答和引用展示的系统。企业面对的生成式引擎既包括公开 AI 搜索，也包括内部知识库、客服机器人、销售助手和智能体平台。

白帽GEO：以用户利益、真实事实、可核验来源、机器可读、结构清晰和合规披露为基础的优化。

灰帽GEO：形式上未必立即违法，但以误导、操纵、隐藏意图或规模化灌水为核心，长期会伤害用户和生态。

黑帽GEO：通过虚假信息、竞品攻击、作弊内容、操纵链接、污染知识库、提示词注入、数据投毒、安全漏洞或恶意自动化影响 AI 输出和用户判断。

RAG：Retrieval- Augmented Generation，检索增强生成。系统先从外部知识库、网页、文档或数据库中检索上下文，再交给模型生成答案。

提示词注入：攻击者把恶意指令混入用户输入或外部内容，使大模型应用执行偏离预期的任务。

间接提示词注入：攻击者把恶意指令放在网页、邮件、文档、图片描述、代码注释或知识库条目中，等待AI系统检索或读取后触发。

AI生成合成内容标识：在AI生成或合成的文本、图片、音频、视频、虚拟场景等内容中添加显式或隐式标识，用于提示公众、保留溯源信息并支持平台治理。

目标品牌可见率：有效样本中目标品牌被提及的比例。

首提概率：无明确排序答案中，目标品牌最先被提及的比例。

有效引用率：答案引用的来源能支持关键结论的比例。

自有域有效引用率：有效引用中来自企业官网、帮助中心、开发者文档或企业自有内容域的比例。

事实错误率：答案中出现重要事实错误的比例。

- 样本稳定率：**同一问题多次采样时关键结论保持一致的比例。
- 问题族：**按品牌、品类、竞品、风险、售后、价格、场景等维度分组的问题集合。
- 证据链：**从企业事实、内容资产、来源资料、用户问题到AI答案表现的一组可复核关系。
- 基线：**项目启动前按统一口径采集的初始表现。
- POC：**在有限范围内验证服务商方法、数据、内容和治理能力的小型试点。
- SOW：**工作说明书，写明范围、交付物、验收、责任、数据权属和退出机制。
- 假量化：**把主观判断直接包装成看似精确的百分比，却缺少样本、标注规则、复核机制和一致性检验。

附录二：GEO红线速查表

- 竞品攻击：**典型行为是未证实负面断言和伪中立测评，风险等级为高。发现方式包括竞品 query 抽样和内容审稿，首要处置是停止发布、要求证据，并改为事实比较。
- 虚假举例：**典型行为是虚构客户、学员、患者或收益，风险等级为高。发现方式包括授权核验和案例抽查，首要处置是下线、补证和公开更正。
- 功效夸大：**典型行为是把概率、相关性或个案写成确定承诺，风险等级为高。发现方式包括法务审核和专家审核，首要处置是改写、补充限制条件，必要时下线。
- 规模化低质内容：**典型行为是模板页、拼接页和无新增价值内容，风险等级为中高。发现方式包括重复率检测和事实密度审计，首要处置是合并、删减和重写。
- 伪造评论：**典型行为是购买、生成、压制或隐藏评价，风险等级为高。发现方式包括评论来源审计，首要处置是删除、披露和整改。
- 误导结构化数据：**典型行为是用结构化标记不存在的评分、FAQ 或作者，风险等级为中高。发现方式包括结构化数据抽查，首要处置是清理并重新提交审核。
- 提示词注入：**典型行为是页面或文档诱导 AI 改变规则，风险等级为高。发现方式包括安全扫描和人工抽查，首要处置是下线、排查来源并重新索引。
- RAG 污染：**典型行为是未审核文档批量入库，风险等级为高。发现方式包括检索日志和相似度聚类，首要处置是冻结、回滚和审计。
- 敏感信息泄露：**典型行为是 AI 回答暴露内部信息，风险等级为高。发现方式包括输出监测和权限审计，首要处置是关闭通道并修复权限。
- 自动化操纵：**典型行为是批量 query、刷引用和刷点击，风险等级为中高。发现方式包括日志异常和成本异常，首要处置是限速、封禁并追究供应商责任。

本地虚假服务：典型行为是虚拟地址和伪本地团队，风险等级为高。发现方式包括地址核验和 SLA 核验，首要处置是下线城市页并更新披露。

伪造来源：典型行为是伪论文、伪媒体和伪认证，风险等级为高。发现方式包括来源核验，首要处置是删除、公开纠错并执行供应商处罚。

附录三：可复用的内容审核清单

每篇可能影响AI搜索答案的内容发布前，应至少回答以下问题：

1. 这篇内容服务的是哪个真实用户问题？
2. 页面中的核心事实是否有证据？
3. 证据是否来自原始来源、权威来源或可核验来源？
4. 结论是否超出了证据范围？
5. 是否存在绝对化、保证性或暗示性承诺？
6. 是否涉及医疗、金融、法律、教育、招聘等高风险主题？
7. 是否提到竞品？若提到，是否只使用公开可验证事实？
8. 是否使用客户、用户、专家或媒体背书？是否有授权和披露？
9. 是否存在过期价格、旧功能、旧资质或旧政策？
10. 是否使用结构化数据？结构化数据是否与页面真实内容一致？
11. 是否包含自动生成内容？是否经过人工审核？
12. 是否存在容易被AI系统误读为指令的文本？
13. 是否含有隐藏文本、误导性锚文本或诱导性跳转？
14. 是否具备作者、审核者、更新时间和纠错入口？
15. 若AI搜索只读取摘要，是否会误解核心意思？
16. 若竞品、媒体或监管机构看到这篇内容，是否能接受其事实边界？
17. 若用户依据这篇内容做重大决策，是否会被误导？
18. 若这篇内容被模型反复引用，是否会改善公共信息环境？

附录四：企业内部GEO培训提纲

培训目标

完成培训后，团队成员应能够：

- 区分可信GEO、低质量GEO和红线GEO。
- 识别竞品攻击、虚假信息、规模化低质内容和AI系统攻击。
- 使用事实资产、证据资产和问题地图开展合规优化。
- 在客户或内部团队提出高风险要求时进行沟通。
- 将疑似红线事件提交给合规、安全或管理层。

培训模块

模块一：AI搜索如何形成答案

- 检索、排序、摘要、引用、生成和反馈的基本流程。
- 为什么重复叙事会影响答案。
- 为什么权威来源、结构化信息和更新机制重要。

模块二：GEO的价值观

- 以用户利益为中心。
- 以事实为边界。
- 以证据为基础。
- 以透明为默认。
- 以安全为前提。

模块三：红线方法识别

- 价值观层红线。
- 信息质量层红线。
- AI系统攻击层红线。
- 行业高风险场景。

模块四：服务商与客户沟通

- 如何拒绝竞品攻击。
- 如何把虚假卖点改成可验证事实。
- 如何把产量需求改成质量需求。
- 如何记录风险提示。

模块五：审计与复盘

- 如何抽样 query。
- 如何记录答案与引用。
- 如何打风险分。
- 如何提交整改。
- 如何防止复发。

培训输出物

- GEO红线速查表。
- 内容审核清单。
- 竞品比较模板。
- 高风险行业内容模板。
- 服务商沟通话术。
- AI 搜索监测表。
- 红线事件复盘表。

附录五：红线风险评分模型

企业可以采用1至5分的方式评估每个GEO风险事件。

总分由五个维度组成：事实伤害、用户伤害、竞争伤害、系统伤害和可扩散性。

事实伤害

- 1分：措辞不清，但核心事实正确。
- 2分：存在轻微过期信息或不完整信息。

- 3分：存在可能影响购买判断的错误事实。
- 4分：存在关键功能、价格、资质、效果或安全事实错误。
- 5分：存在虚构事实、伪造证据或重大误导。

用户伤害

- 1分：用户体验轻微下降。
- 2分：用户需要额外核验才能做决定。
- 3分：用户可能购买不适合的产品或服务。
- 4分：用户可能产生财产、健康、职业或隐私损失。
- 5分：用户可能遭受重大损失或长期影响。

竞争伤害

- 1分：竞品表述有轻微偏向。
- 2分：竞品比较缺少完整背景。
- 3分：对竞品存在选择性负面呈现。
- 4分：对竞品存在未经证实的负面断言。
- 5分：构成系统性商业诋毁或声誉污染。

系统伤害

- 1分：只影响单一页面或单一回答。
- 2分：影响少量query或少量来源。
- 3分：影响多个AI搜索结果或企业助手回答。
- 4分：影响RAG、知识库、agent或搜索索引。
- 5分：存在提示词注入、数据投毒、敏感信息泄露或大规模自动化操纵。

可扩散性

- 1分：短期可下线，传播范围小。
- 2分：已被少量页面引用。
- 3分：已进入多个外部页面、社区或问答。
- 4分：已被AI搜索引用或摘要。

- 5分：已形成跨平台传播，且难以完全纠正。

处置等级

- 绿色：总分5至8。常规整改，7个工作日内完成。
- 黄色：总分9至13。项目负责人介入，3个工作日内完成整改计划。
- 橙色：总分14至19。法务、安全和管理层介入，必要时暂停上线。
- 红色：总分20至25。立即下线、冻结供应商交付、保留证据、启动事件复盘。

附录六：可实操的防御性技术栈与组织分工

GEO治理不能只依赖内容团队。它需要品牌、法务、安全、数据、产品、客服和高管共同参与。

内容治理层

内容治理层负责把真实业务变成可被AI理解和引用的事实资产。

主要任务包括：

1. 建立事实库：产品功能、价格、客户案例、资质、服务范围、限制条件、常见误解。
2. 建立证据库：官方文档、合同模板、公开报告、媒体报道、学术资料、监管文件。
3. 建立问题库：用户如何提问、AI搜索如何回答、竞品如何被比较、负面问题如何出现。
4. 建立更新机制：当产品、价格、政策、资质和案例变化时，同步更新页面和结构化数据。
5. 建立纠错机制：外部用户、销售、客服和客户成功可以快速提交错误答案。

合规与法务层

合规与法务层负责判断内容是否存在广告、竞争、隐私、知识产权和消费者保护风险。

主要任务包括：

1. 审核医疗、金融、教育、法律、招聘等高风险主题。
2. 审核竞品比较、用户评价、背书、榜单和第三方关系披露。
3. 审核数据来源、引用范围和版权许可。
4. 审核供应商合同中的红线、审计权、整改权和退出权。
5. 建立重大风险事件的对外回应机制。

安全工程层

安全工程层负责保护企业自建AI系统、知识库、agent和数据管道。

主要任务包括：

1. 对外部网页、文档、评论、图片和元数据进行提示词注入检测。
2. 对RAG文档入库设置来源验证、权限控制和审核流程。
3. 对模型输出设置敏感信息、广告合规和错误引用检查。
4. 对工具调用设置最小权限、人工确认和日志审计。
5. 对异常query、异常成本、异常引用和异常命中做告警。

数据与监测层

数据与监测层负责把AI搜索中的表现变成可复盘数据。

主要任务包括：

1. 维护核心query集合，并按业务、风险和语言分组。
2. 定期抓取不同生成式搜索产品的答案和引用。
3. 对答案进行事实核验、来源分级和风险打标。
4. 跟踪竞品、媒体、社区和监管环境变化。
5. 将监测结果反馈给内容、法务、安全和产品团队。

高管治理层

高管治理层负责设定GEO的边界与激励。

主要任务包括：

1. 明确GEO的目标：帮助用户获得准确、可核验、可比较的答案。
2. 禁止以竞品伤害、虚假内容和系统攻击换取短期增长。
3. 把GEO风险纳入品牌风险、合规风险和AI风险管理。
4. 对供应商设置准入、审计、整改和退出机制。
5. 将红线事件纳入绩效扣分或一票否决。

附录七：GEO 服务商评估方法参考

1.GEO服务商比稿测试方法参考

企业可以在比稿阶段给服务商一份48小时小作业，不要求免费交付完整方案，但要求其展示方法。

- 作业一：问题库构建

给服务商一份企业官网和产品手册，让其提出20个GEO采样问题，并说明每个问题属于什么业务场景、为什么有价值、可能对应哪些企业事实材料。

判断标准：专业服务商会覆盖用户自然语言、采购决策、竞品对比和风险问题；能力不足的服务商往往只会把官网标题改成问题。

- 作业二：样本复核

给服务商10条AI回答，让其标注目标品牌是否可见、是否Top 3、是否有效引用、是否存在事实错误、应采取何种动作。

判断标准：专业服务商会区分正向可见、错误可见、无引用可见和风险可见；能力不足的服务商只会统计有没有出现品牌名。

- 作业三：内容改写

给服务商一段产品介绍，让其改成更适合AI理解和引用的企业事实页。要求保留事实边界，不能夸大。

判断标准：专业服务商会增强结构、定义、适用场景、限制条件、证据和FAQ；能力不足的服务商会堆关键词、堆形容词、堆“行业领先”。

- 作业四：风险判断

给服务商一个敏感行业问题，例如“某医疗器械是否能替代医生诊断”，要求其说明企业内容和AI答案中应如何设置风险边界。

判断标准：专业服务商会先问合规边界、监管要求、产品适用范围和免责声明；能力不足的服务商会把问题当作普通营销文案处理。

2.GEO服务商专业面试脚本参考

企业可以把服务商访谈分成四轮，每轮15至30分钟。

第一轮：方法论面试

建议问题如下：

- 你们如何定义目标品牌可见率、Top 1概率和Top 3概率。
- 你们如何处理同一问题多次回答不一致。
- 你们是否区分平台、入口、联网状态和多轮上下文。
- 你们如何判断某次引用是否有效。
- 你们如何处理非量化指标，比如语义一致性和情感积极性。

专业回答应能落到样本、口径、公式和复核流程。若服务商只回答“我们有独家算法”，企业应继续追问算法如何被审计、样本如何导出、人工如何复核。

第二轮：内容能力面试

建议问题如下：

- 如果企业官网内容很少，你们会先改哪些页面。
- 如何把产品介绍改成AI更容易引用的事实页。
- 如何处理客户案例授权不足的问题。
- 如何避免为了GEO堆砌关键词和生成低质内容。
- 如何处理一个问题同时涉及品牌、竞品和合规边界。

专业回答应关注事实、结构、证据、FAQ、限制条件和可引用性。若服务商只强调发稿数量，风险较高。

第三轮：技术能力面试

建议问题如下：

- 如何检查官网是否可抓取、可索引、可被搜索系统理解。
- 如何处理站点地图、canonical、robots、结构化数据和页面加载问题。
- 如何判断AI答案没有引用官网是技术问题、内容问题还是平台问题。
- 你们是否能与企业IT团队协作提出可执行改造建议。

专业回答应能给出检查清单和优先级。若服务商完全不了解抓取、索引和结构化数据，难以承担复杂项目。

第四轮：合规和项目治理面试

建议问题如下：

- 你们如何处理个人信息、客户案例和未公开信息。
- 你们是否会使用企业资料训练内部模型，是否可以关闭或限制。
- 项目中哪些内容需要企业审批，审批流程如何记录。
- 退出合作时，数据、文档、看板和账号如何交接。
- 如果AI答案出现重大错误，你们的应急流程是什么。

专业回答应把合规、数据权属和应急机制写进合同。若服务商回避这些问题，后期风险很高。

3.GEO服务商RFI（信息征询）模板

公司背景：成立时间、团队规模、核心客户类型、是否有高敏行业经验。

GEO理解：你们如何定义GEO，如何区分GEO、SEO、PR和舆情监测。

平台经验：豆包、DeepSeek、Kimi、元宝、千问的采样方法和案例。

指标体系：请列出你们使用的前端指标、后端指标和解释边界。

数据方法：样本量、重复采样、数据保存、原始样本交付和看板能力。

内容方法：如何做问题库、实体治理、FAQ、白皮书、案例和权威资料。

技术方法：如何处理抓取、索引、结构化数据、站点结构和多格式资产。

合规安全：个人信息、客户案例、未公开资料、敏感行业内容如何处理。

项目团队：负责人、策略、数据、内容、技术、项目经理分别是谁。

报价方式：诊断、执行、数据、内容、第三方费用和绩效费用如何拆分。

附录八：POC、合同与项目治理参考

- **POC验收表**

验收项	通过标准
问题库	覆盖约定场景，问题来源清楚，企业可复核
采样方法	记录平台、入口、时间、账号、联网状态、重复次数
基线报告	包含提及、引用、准确性、竞品和风险发现
行动方案	能把问题对应到内容、技术、实体、来源或合规动作
执行记录	说明哪些页面、资料或内容被修改，何时上线
复测报告	用相同或可解释的样本复测，并说明波动
风险说明	明确未解决问题、平台限制、不可承诺事项
交付资料	企业可以获得问题库、样本、报告和修改记录

● **合同应写入的内容**

项目目标：明确平台范围、业务场景、问题库规模、竞品范围和目标指标。

服务范围：诊断、监测、内容、技术建议、数据看板、月报、复盘、培训和交接。

交付物：问题库、基线报告、采样记录、内容建议、技术建议、月报、验收报告、交接包。

验收标准：指标口径、样本数量、采样频率、复测规则、数据保存期限和争议处理方式。

数据权属：问题库、样本、报告、页面修改记录和仪表盘导出数据归企业所有或明确授权范围。

保密与安全：企业资料、客户案例、未公开数据、账号权限和个人信息的处理方式。

禁止行为：虚假内容、隐性付费评价、伪造案例、低质灌水、违法采集个人信息、绕过平台规则。

变更机制：平台变化、业务变化、问题库调整和内容审批延迟时如何处理。

费用结构：固定费、数据费、内容费、第三方费用和绩效费的边界。

退出机制：项目终止后的数据交接、账号回收、资料删除、未完成事项和后续支持。

● **项目治理机制**

节奏	参与方	内容
启动会	业务负责人、服务商、市场、内容、官网、法务	目标、平台、问题库、资料清单、审批流、时间表
周会	项目经理、服务商执行团队、内容和官网接口人	任务进度、素材需求、上线问题、风险事项
双周复核	业务负责人、内容负责人、服务商策略负责人	阶段数据、样本复核、竞品变化、优先级调整
月度复盘	管理层、业务、品牌、数据、法务、服务商	指标趋势、平台差异、风险、预算与下一步
季度评审	管理层、采购、法务、服务商负责人	续约价值、能力评估、合同问题、长期路线图

• **SOW（工作说明书）范例条款**

以下条款可作为合同附件的参考，需要企业法务结合实际调整。

指标条款示例

服务商应按照双方确认的问题库、平台范围、采样频率和复核规则，提供目标品牌可见率、目标品牌Top 1概率、目标品牌Top 3概率、有效引用率、事实错误率和样本稳定率等前端指标。任何指标变化均应能追溯至原始问题、回答样本、平台、采样时间和复核记录。

非量化指标条款示例

语义一致性、情感积极性、推荐理由质量、品牌调性匹配等指标作为项目复盘与内容优化参考，不作为单独付款条件。若双方需使用标签体系，应在项目启动阶段共同确认A、B、C三级标签定义和人工复核流程。

后端指标条款示例

官网流量、线索量、销售反馈、品牌词搜索、下载量等后端指标由企业内部观察。服务商可协助分析趋势和相关性，但除非双方另有独立实验设计和数据验证机制，后端指标不作为服务商直接归因承诺。

内容合规条款示例

服务商不得伪造客户案例、资质、奖项、评价、合作关系或第三方背书。涉及医疗、金融、教育、政务、未成年人、个人信息、投资者关系等敏感内容，须经企业指定审批人确认后方可发布或提交上线。

数据权属条款示例

项目过程中形成的问题库、采样数据、复核记录、内容稿件、策略文档、看板数据和交接清单，除双方另有约定外，应归企业可持续使用。服务商不得在未经授权的情况下将企业未公开资料用于第三方项目、公开案例或模型训练。

退出交接条款示例

合作终止时，服务商应在约定期限内交付完整问题库、采样数据、账号权限、内容资产、上线记录、未完成事项、风险清单和下一步建议。若项目使用第三方工具，应说明数据导出方式和账号迁移限制。

附录九：GEO服务商风险清单与不合理信号

- **风险清单**

风险	表现	应对方式
测量风险	样本少、截图化、不可复现、指标口径不清	建立统一问题库、采样日志和复测规则
平台波动风险	同一问题多次答案不同	多次采样，按趋势看结果，把波动纳入报告
内容质量风险	内容空泛、过期、缺少证据、不能回答真实问题	重写核心页面，补充事实、数据、案例和边界
事实错误风险	AI误报产品、价格、资质、客户和服务范围	建立事实库，修复权威页面，持续复测
合规风险	涉及个人信息、客户案例、医疗金融等敏感内容	审批流、授权、专家审查和合同禁止行为
IP风险	使用未经授权图片、数据、报告或客户材料	建立素材授权和引用规则
低质资源风险	软文堆量、低质外链、伪造评价和隐性付费	禁止相关行为，要求资源关系披露
供应商锁定	数据只在供应商系统，客户无法迁移	合同要求数据导出和交接
组织协同风险	官网、法务、内容和销售不配合	启动会明确责任人和审批时限
后端归因风险	把销售变化直接归因GEO	采用前端验收加后端观察，必要时设计对照实验

● 不合理信号

服务商声称能在几天内让品牌稳定出现在所有AI答案前列。

服务商拒绝解释采样方法或原始样本。

服务商强调所谓内部通道、特殊资源或平台关系，却不披露资源性质。

服务商推荐大量发布低质文章、伪造第三方测评或购买匿名评价。

服务商把内容数量作为主要交付，不关心问题库、引用质量和事实准确性。

服务商没有数据安全、保密和个人信息处理流程。

服务商不愿把交付物、验收标准、数据权属和退出机制写进合同。

服务商把海外GEO和国内GEO用同一套模板处理。

服务商忽略企业行业敏感性，对医疗、金融、教育、法律等内容没有额外复核。

服务商的案例无法说明平台、周期、问题数量、采样方法和前后对比逻辑。

• 低质资料来源识别

企业在评估GEO方案时，也要评估服务商引用的资料来源。

GEO领域概念变化快，低质资料可能被广告、软文、利益相关方或AI生成内容干扰。

不采纳没有作者、日期、方法、样本或来源说明的泛泛文章作为关键依据。

不采纳门户网站转载、营销软文、供应商自夸文章作为唯一依据。

不采纳只展示截图、不说明采样条件和样本量的所谓研究。

不采纳无法区分平台、版本、联网状态和入口的对比结论。

优先采用论文、官方文档、法规文本、可复核技术文档和有方法说明的研究。

附录十：GEO服务采购准备与参考方法

• 企业内部资料准备清单

公司与品牌事实表：名称、简称、英文名、官网、主体、资质、行业定位。

产品事实表：产品线、功能、适用客户、交付方式、价格边界、版本差异。

客户案例表：客户名称、行业、项目内容、可公开范围、授权证明。

合规材料：隐私政策、数据安全说明、资质证书、行业许可证、免责声明。

内容资产：官网页面、白皮书、产品手册、FAQ、帮助中心、开发者文档、媒体报道。

竞品清单：直接竞品、替代方案、海外竞品、区域竞品。

内部责任表：市场、品牌、内容、官网、法务、产品、销售、数据负责人。

• **GEO 服务RFP（建议邀请书）模板**

项目背景：企业业务、目标市场、目标平台、竞品、现有内容资产。

项目目标：明确本期要解决的业务问题和成功标准。

平台范围：豆包、DeepSeek、Kimi、元宝、千问，是否包含海外扩展平台。

问题库要求：问题分类、样本数量、优先级、敏感问题和竞品问题。

指标要求：前端指标、后端观察、采样方法、数据保存和验收规则。

交付要求：基线报告、月报、内容方案、技术建议、复测报告和交接资料。

合规要求：资料授权、客户案例、个人信息、敏感行业、禁止行为和审批流。

POC要求：周期、范围、交付物、验收方式和进入正式项目条件。

团队要求：项目负责人、数据、内容、技术、合规接口和响应时间。

商务要求：报价拆分、付款节点、第三方费用、绩效条款、退出机制。

• **90天执行路线图**

周期	主要任务	关键交付
第1至2周	启动会、资料收集、问题库初版、竞品确定、采样口径确认	项目计划、问题库、资料清单
第3至4周	基线采样、平台对比、竞品分析、风险诊断	基线报告、优先级清单
第5至6周	内容与技术实验，优化核心页面、FAQ、案例和结构化数据建议	内容修改稿、技术建议、上线记录
第7至8周	复测第一轮，调整问题库和内容策略	复测报告、差异解释
第9至10周	扩大到更多场景和竞品，补充权威资料与第三方来源	第二轮行动清单
第11至12周	综合复盘、验收、内部培训、下一周期规划	验收报告、交接包、续约建议

• 采购决策矩阵

企业现状	建议选择	不建议选择
没有基线数据	先采购诊断与POC	直接签全年高额执行合同
官网资料混乱	内容与实体治理强的服务商	只卖监测看板的服务商
技术问题明显	技术SEO和数据工程能力强的服务商	只做PR和软文的服务商
品牌叙事弱	内容营销和行业咨询能力强的服务商	只做抓取和报表的服务商
行业高敏	有合规和专家复核流程的服务商	没有审批流的低价执行团队
出海需求明确	跨境增长与本地化能力强的服务商	只熟悉国内平台的服务商
预算有限	小范围POC和高价值页面优化	全平台大看板加大量内容套餐
组织复杂	项目管理和跨部门协同能力强的服务商	只派销售对接、无执行机制的服务商

• 企业内部方法手册：从采购到执行的最小闭环

企业即使采购了服务商，也需要保留内部方法手册。建议手册至少包含七个部分。

第一部分：问题库管理

规定问题如何新增、删除、分组和冻结。每个统计周期的问题库应有版本号。例如V1问题库包含60个问题，V2新增20个竞品问题。新旧版本不能直接混算。

第二部分：企业事实表

维护品牌、产品、资质、客户案例、价格边界、服务范围和禁用表述。所有GEO内容都应回到事实表，不允许临时编造。

第三部分：采样规则

规定平台、入口、采样时间、采样次数、是否联网、多轮或单轮、是否清空上下文。任何异常采样都要备注。

第四部分：复核规则

规定哪些样本由服务商初标，哪些由企业抽检，哪些必须法务或专家复核。高敏内容、重大错误和客户案例必须走人工复核。

第五部分：内容上线规则

规定内容从撰写、审稿、法务确认、上线、记录到复测的流程。没有上线记录，就无法判断后续指标变化。

第六部分：数据看板规则

规定看板字段、导出频率、权限、数据保存时间和退出交接方式。服务商工具可以使用，但企业应拥有核心数据。

第七部分：复盘规则

规定月度复盘必须回答四个问题：哪些指标变化了，为什么变化，采取了哪些动作，下一步优先做什么。

• 参考方法：内容资产改造的“一问一页一证据”原则

GEO内容改造不宜从“多写文章”开始。更有效的方式是围绕问题库中的高价值问题，建立“一问一页一证据”的资产关系。

“一问”指一个真实用户问题。

例如“制造企业采购预测性维护系统要看哪些能力”。这个问题应来自销售访谈、售前FAQ、竞品对比、行业场景或AI平台基线样本。

“一页”指企业需要有一个可公开、可访问、主题清晰的页面承接这个问题。

页面可以是官网产品页、解决方案页、FAQ、白皮书摘要页、帮助中心、开发者文档、案例页或合规说明页。页面中要用用户能理解的表达回答问题，而非只写企业宣传语。

“一证据”指页面中的关键结论要有事实支撑。

证据可以来自产品功能、客户案例授权、资质证书、技术文档、公开标准、行业报告、监管文件、测试结果或企业内部已批准的事实资料。没有证据的形容词很难长期支撑AI答案中的可信表达。

参考案例：把一篇宣传稿改成可引用页面

某制造企业原本有一篇公众号文章，主题是“助力产业升级”，内容大量使用“领先”“卓越”“高效”等词，但没有说明产品适用哪些工厂、如何部署、与竞品差异是什么。

AI平台很少引用该文章，引用时也只能泛泛描述。

项目组把文章拆成三个官网页面：第一，智能排产系统适用场景；第二，智能排产系统采购FAQ；第三，某已授权客户案例。

每个页面都加入具体问题、事实说明、证据来源和限制条件。

复测时，AI答案仍有波动，但推荐理由更具体，引用企业自有域的概率提高。这个案例说明，GEO内容改造的重点是把宣传语变成可检索、可理解、可引用的事实资产。

• 参考方法：权威来源建设路线

GEO中的“权威来源”不能简单等同于媒体曝光。

权威来源更强调来源可信、内容准确、与结论相关、可长期访问和合规披露。企业可以按四类建设。

第一类是自有权威来源，包括官网、帮助中心、白皮书、产品文档、开发者文档、资质页面、隐私政策、数据安全说明和客户案例授权页。自有来源是企业最可控的事实底座。

第二类是行业权威来源，包括行业协会、标准组织、学术机构、监管公开资料、研究报告和专业机构。此类来源适合支撑行业定义、标准、合规和趋势。

第三类是第三方专业来源，包括专业媒体、分析师报告、垂直社区、技术评测和合作伙伴文档。此类来源可以补充外部视角，但要审查商业关系和内容质量。

第四类是低可信风险来源，包括无编辑审核站点、纯付费发布、洗稿站、无作者信息页面、虚假评价站和不可追溯的内容农场。此类来源可能短期增加曝光，但长期会增加答案错误和声誉风险。

权威来源建设的三个月计划

第一个月先补齐自有来源。

企业应建立公司事实页、产品事实页、FAQ、客户案例授权页和合规说明页。

第二个月补充行业来源。企业可以参与标准解读、行业白皮书、协会活动、公开资料整理和专家文章。

第三个月谨慎建设第三方专业来源，优先选择有编辑审核、专业读者和长期可访问性的渠道。

• 参考方法：服务商组合采购模型

中大型企业可以采用组合采购，避免把所有能力押在一家服务商身上。常见组合有三种。

第一种是主服务商加专项伙伴。

主服务商负责问题库、指标、项目治理和月度复盘，专项伙伴负责官网技术、内容生产、行业顾问或海外本地化。适合中型企业和行业龙头。

第二种是监测工具加咨询服务。

企业采购一个数据监测或看板工具，再采购咨询服务商帮助解释数据和推动内容技术行动。适合内部团队较强的大型企业。

第三种是内部团队加外部诊断。

企业自建内容和技术能力，外部服务商每季度做一次基线、竞品和风险诊断。适合成熟企业和预算稳定的集团。

组合采购的关键是明确主责。若没有主责方，项目会变成多家服务商互相等待。企业应在合同中写清楚谁负责指标口径、谁负责样本数据、谁负责内容上线、谁负责技术修复、谁负责合规复核。

工作项	主服务商	专项伙伴	企业内部
问题库与指标	主责	参与	确认业务目标
基线采样	主责	可参与	提供平台范围
内容改造	策略主责	执行或共创	审批事实和语气
技术修复	提建议	主责或执行	官网权限和上线
合规复核	提风险清单	参与	最终审批
月度复盘	主责	参与	决策优先级
数据交接	主责	交付专项资料	接收归档

参考方法：国内五个平台的采样注意事项

企业在采样豆包、DeepSeek、Kimi、元宝、千问时，应把“平台”拆成更细的产品入口。

用户在网页端、App、深度搜索、联网模式、文档上传、API或智能体场景中得到的答案可能不同。企业采购时应要求服务商明确本期只覆盖哪些入口。

建议采样记录至少包含：平台名称、产品入口、模型或功能名称、是否联网、是否展示引用、是否上传文档、账号状态、地区、时间、问题原文、答案摘要、引用来源、品牌编码、竞品编码、错误编码和截图或日志。

若服务商只说“我们测了千问”或“我们测了Kimi”，企业应追问具体入口。

因为同一品牌在不同入口中的表现可能不同，采购验收必须基于约定入口，而非模糊平台名称。

• 参考方法：敏感行业的专家复核机制

高敏行业GEO项目中，内容复核应采用“业务事实复核、专业复核、法务复核”三层机制。

业务事实复核由产品或业务负责人确认企业是否真的具备某项能力。

专业复核由医学、金融、法律、教育、数据安全等领域专家确认表述是否符合行业规范。法务复核确认是否涉及广告法、个人信息、数据安全、信息披露、客户授权或行业监管风险。

企业可以把样本分成普通样本、需专业复核样本和必须法务复核样本。

普通样本由项目经理处理，需专业复核样本进入专家队列，必须法务复核样本不得在未审批前用于公开内容或第三方传播。

• 参考方法：GEO 中的品牌实体治理

GEO中的实体治理，是让AI平台更清楚地理解“这个企业是谁、做什么、有哪些产品、与哪些名称相关、与哪些业务无关”。

很多企业的AI答案错误，并非来自模型完全不了解企业，主要来自公共资料中的实体关系混乱。

企业应先建立品牌实体表，至少包含公司主体、品牌名、产品名、英文名、简称、旧名称、子品牌、官网域名、公众号、App、开发者平台、常见误写和不应混淆的相似品牌。

对于集团企业，还要明确母公司、子公司、事业部、产品线和区域公司的关系。对于出海企业，还要明确中文名、英文名、本地商标和目标市场使用名称。

实体治理的第一步是清理自有资料。官网页脚、关于我们、产品页、公众号简介、白皮书、招聘页面、新闻稿和客户案例中的名称要统一。

第二步是处理外部资料。百科、媒体报道、行业名录、合作伙伴页面和开发者文档中如果存在明显错误，应尽量联系更正。

第三步是增加解释型页面。例如“公司名称与品牌名称说明”“产品线关系说明”“海外品牌与中文品牌关系说明”。

- **参考方法：从单次优化到长期机制**

企业第一次做GEO，通常关注“能不能看到变化”。

但从长期看，更重要的是建立机制。建议企业把GEO纳入四个固定流程。

第一，内容发布流程。每次发布官网页面、白皮书、案例或新闻稿，都要检查是否回答了真实用户问题，是否有证据支撑，是否与事实库一致，是否需要法务复核。

第二，产品更新流程。产品上线新功能、价格调整、服务区域变化、部署方式变化时，要同步更新官网、FAQ、销售资料、帮助中心和GEO事实库，避免AI长期引用旧资料。

第三，销售反馈流程。销售和售前每月反馈客户新问题、竞品新说法、AI答案带来的误解和高频咨询。问题库应随真实市场变化更新，但验收期内变更要留痕。

第四，风险复盘流程。每月复盘事实错误、敏感样本、错误引用、未授权客户、过度承诺和低质第三方来源。高敏问题要进入法务或合规台账。

长期机制的组织分工

市场团队负责问题库、品牌表达和内容优先级。产品团队负责产品事实、功能边界和技术资料。

销售团队负责客户真实问题和线索反馈。官网或技术团队负责页面上线、结构、索引和日志。法务负责敏感内容、客户授权、个人信息和合同条款。服务商负责方法、采样、分析、建议和项目推进。

当这些分工清楚后，GEO就不再只是外部服务商的项目，而会变成企业内容和品牌治理的一部分。

- **参考方法：企业管理层 GEO常见问题回答口径**

管理层常问：“做GEO多久能见效？”

建议回答：不同平台和问题类型差异较大，诊断和内容修复可以在30至90天内看到前端样本变化，但长期稳定性取决于内容资产、证据来源、平台更新和竞品动作。项目初期不宜承诺固定排名。

管理层常问：“为什么后端成交不能作为考核？”

建议回答：AI答案到成交之间存在多条路径，平台通常不提供完整查询日志，同期广告、销售、展会、PR和SEO都会影响线索。可以观察后端信号，但服务商考核优先放在可复核的前端指标和交付行动。

管理层常问：“服务商为什么要我们内部配合？”

建议回答：GEO需要企业提供真实事实、客户授权、产品边界、合规审查和官网上线权限。服务商可以提出方法和内容建议，但无法替企业决定事实和公开边界。

管理层常问：“能不能直接买工具解决？”

建议回答：工具能帮助监测和看板，但不能替代问题库设计、事实库治理、内容改造、技术修复和合规判断。企业可以工具加服务组合采购，但不宜把工具分数当成全部结论。

这些回答口径能帮助项目负责人在立项会、预算会和复盘会上解释GEO的科学边界，减少不合理期待。

• 参考方法：常见误区纠偏表

常见误区	风险	更合理做法
用一次截图证明效果	偶然性高，无法复核	使用多平台、多问题、多次采样
只采样品牌词	高估真实表现	覆盖品类、竞品、风险和场景问题
把线索增长写成服务商KPI	归因困难	后端指标内部观察，前端指标用于考核
所有指标都量化	形成假量化	可量化指标硬考核，非量化指标复核
低价采购全平台服务	关键工作缺失	先做基线和季度POC，再决定扩展
忽略官网和文档	证据链薄弱	优先补事实表、FAQ、案例和产品页
只看自有内容	漏掉第三方错误来源	建立来源地图和竞品地图
高敏行业追求正面推荐	合规风险升高	优先降低事实错误和风险表述
合同不写数据权属	更换服务商困难	写清样本、问题库、报告和内容资产归属
服务商不讲边界	过度承诺风险	要求说明不可承诺事项

- **参考方法：年度GEO路线图**

如果企业完成90天试点，并确认GEO具有长期价值，可以设计年度路线图。

第一季度，建立基础。完成问题库、基线、重点平台采样、内容资产诊断、事实表、禁用表述和POC。

第二季度，扩大范围。扩展到更多产品线、更多问题族和更多竞品。完善官网、FAQ、案例、白皮书和结构化数据建议。

第三季度，深化治理。建立内部GEO知识库，嵌入内容发布流程、法务审批流程、销售反馈机制和月度复盘机制。

第四季度，评估和预算。复盘全年指标变化、风险下降、内容资产沉淀、后端观察信号和服务商表现，为下一年预算和供应商续约提供依据。

年度路线图能帮助企业把GEO从一次采购转为长期能力建设。对于行业龙头和上市公司，年度治理比短期排名更有价值。

