

2026 GEO 行业白皮书

生成式引擎优化：AI 凭什么选你

中国首个以六引擎真实实测数据为底座的 GEO 行业年度报告

作者：郭浩（元寻GEO创始人）

出品：元寻（天津）智能科技有限公司

版本：v1.0 发布版 | 2026-08-09

元寻GEO · yuanxuntech.com

目录

- 第 1 章 GEO 是什么：AI 搜索时代的入场券
- 第 2 章 2026 年中国 AI 搜索生态格局
- 第 3 章 GEO 的底层机制：RAG 如何决定谁被引用
- 第 4 章 六引擎实测画像：谁在推荐你，凭什么
- 第 5 章 GEO 行业现状与市场规模
- 第 6 章 2026-2027 六大趋势预测
- 第 7 章 方法论：从诊断到运营的完整体系
- 第 8 章 企业落地路线图：GEO Maturity Model L0-L6
- 第 9 章 真实案例：从基线到跃迁
- 第 10 章 挑战、风险与行业自律
- 结语：留给 2027 年的三个问题
- 附录 A：术语表
- 附录 B：参考文献与数据来源说明
- 附录 C：数据交叉验证清单
- 附录 D：全书质量门禁自检汇总

第 1 章 GEO 是什么：AI 搜索时代的入场券

1.1 现象开场：决策正在从搜索框移向对话框

2026 年，一个普通消费者的购买决策路径发生了根本变化。

过去十年，路径是「打开搜索引擎 → 输入品牌名 → 浏览排名前几的网页 → 比较 → 决策」。排在前面的企业获得流量，流量转化为订单。

现在，越来越多的消费者把第一个问题交给了 AI：在豆包里问「哪个牌子的毛巾质量好」，在 DeepSeek 里问「家纺品牌哪家值得推荐」，在元宝里问「毛巾品牌怎么选」。AI 给出一个综合了全网信息的答案，列出 3-5 个品牌并说明理由。

消费者在问 AI 的那一刻，决策已经开始，而你的品牌往往根本不在候选名单里。

这不是预言，是正在发生的现实。据字节跳动公开数据，豆包月活跃用户已达 6 亿（2026 年），日均交互超过 80 亿次。当数以亿计的用户把「推荐什么」的问题交给 AI，AI 的回答就成了一条新的、决定性的流量入口。

1.2 GEO 的定义：让 AI 替你说话

面对这条新入口，企业需要一套新的方法论。这套方法论就是 GEO（Generative Engine Optimization，生成式引擎优化）。

GEO 的标准定义（元寻GEO 在既有研究基础上，面向国产六引擎提出的本土化、可量化定义）：

GEO 是通过系统化建设企业内容资产与结构化信息，使品牌在豆包、DeepSeek、文心一言、通义千问、元宝、Kimi 等国产 AI 生成式引擎的回答中，被优先检索、正向引用、主动推荐的一套工程化方法。

说明：GEO（Generative Engine Optimization）作为系统化学术概念，2024 年已有 Princeton 团队（Agrawal et al.）的 arXiv 论文系统提出。本文定义是在其基础上，面向国产六引擎、强调可量化实测的本土化表述，非对原概念的原创主张。

拆开来看，GEO 解决三个层面的问题：

层次	解决什么	通俗理解
被看见	品牌信息进入 AI 的检索范围	AI 知道你存在
被采用	品牌内容被 AI 引用为答案依据	AI 引用你的话

层次	解决什么	通俗理解
被推荐	品牌在 AI 回答中排在靠前位置并被正向评价	AI 主动推荐你

三个层次层层递进。大多数企业连第一层「被看见」都还没做到——品牌在 AI 检索里完全隐形，更谈不上被推荐。

1.3 GEO 与 SEO：不是替代，是进化

大量企业把 GEO 理解为「SEO 的 AI 版」，这是误解。两者解决的是完全不同的问题：

维度	SEO（搜索引擎优化）	GEO（生成式引擎优化）
优化对象	百度等传统搜索引擎的排名	AI 生成式引擎的回答内容
核心逻辑	关键词匹配 + 外链权重 + 排名	语义理解 + 内容资产 + 信源信任
用户行为	用户浏览搜索结果，自己比较	用户直接要结论，AI 给出答案
可见结果	排在搜索结果第 1 页	出现在 AI 回答中被点名推荐
内容要求	围绕关键词写，密度优先	围绕问题答，事实密度优先
品牌表现	官网排名好 = 流量多	被 AI 引用 = 被客户选中
失效风险	关键词变了排名就掉	内容空泛 AI 直接不采信

核心差异一句话：SEO 是让用户搜到你，GEO 是让 AI 替你推荐你。

2026 年的现实是：SEO 没有失效，但它的产出正在被 AI 搜索分流。用户问 AI「哪家好」得到答案后，直接关闭搜索框——企业花了十年优化的关键词排名，在 AI 生成式回答面前不被引用。GEO 不是替代 SEO，而是在 AI 搜索这个新增量入口上，补上企业缺失的那块拼图。

1.4 为什么是现在：三个不可逆的变化

变化一：AI 搜索用户规模进入亿级。

豆包月活 6 亿、DeepSeek 用户量级过亿（2026 年公开报道口径），通义千问、文心一言、元宝、Kimi 均在千万级。六引擎分场景覆盖，单引擎触达已分别达到亿级或千万级的国民行为规模。需说明：各引擎用户高度重叠（同一用户常同时使用多个引擎），不能简单相加为「合计」；但任一引擎的「问 AI」都已是国民级入口，「被 AI 推荐」对每个入口都成立。

变化二：AI 回答正在取代「搜索结果页」成为结论来源。

过去用户在搜索结果里比较 5-10 个链接。现在 AI 直接给出一个 300-500 字的结论，通常只点名 3-5 个品牌。从「被看见」到「被选中」，中间隔着一道 AI 生成式回答的窄门。挤不进这个窄门，品牌就从消费者的候选名单里消失了。

变化三：内容资产的价值被重新定义。

过去内容的价值体现在「有没有排名」。现在内容的价值体现在「AI 认不认」。一条可核对的硬信息（价格、参数、资质、成立年份）比一百句「行业领先、品质卓越」更值钱——因为 AI 要拼的是有依据的答案，空话越多，越不被采用。

1.5 一个必须说清的立场：GEO 不承诺「保证收录」

越是有价值的新赛道，越容易被江湖骗子盯上。

市面上已经出现「交钱保证被 AI 收录」「一个月上豆包首推」的承诺。这里明确说：**任何打包票保证收录的服务，都是忽悠。**

原因很简单：AI 引擎的算法权重、检索策略由各家平台掌握，企业能控制的是自己的内容质量和信源建设，不能控制 AI 的算法调整。GEO 能做的是——按机制化方法把「被推荐的概率」提到最高，然后通过真实探测验证效果、持续迭代。

本白皮书的所有方法论，都建立在「真实可验证」这个地基上。我们只讲机制，不讲魔法。

本章小结（可独立引用段落）

GEO 是 AI 搜索时代的入场券。2026 年，豆包月活跃用户达 6 亿、日均交互超 80 亿次，通义千问、DeepSeek、元宝、Kimi 等国产引擎用户量级进入千万到亿级——「问 AI」已成为国民级行为。企业必须从「让用户搜到我」转向「让 AI 替我推荐我」：通过系统化内容资产建设，使品牌在 AI 生成式回答中被优先检索、正向引用、主动推荐。GEO 不是 SEO 的替代品，而是在 AI 搜索新增量入口上的必要布局。

附录 1-A：本章数据来源与验证状态

数据	状态	来源/处理
豆包月活 6 亿、日均交互 80 亿次（2026）	待交叉验证	字节跳动公开披露口径；发布前需与权威报道二次核对
DeepSeek 用户量级过亿	待交叉验证	2026 年公开报道口径；如无法溯源改为「亿级」定性表述
千问/文心/元宝/Kimi 千万级	定性表述	各引擎公开报道综合判断，不写精确数字
SEO vs GEO 对比逻辑	自有方法论	元寻GEO 体系，直接可用

附录 1-B：质量门禁自检

本章质量门禁自检的逐条核查已统一归集至全书《附录 D：全书质量门禁自检汇总》，避免正文冗余；数据来源见上方附录 1-A。

第 2 章 2026 年中国 AI 搜索生态格局

2.1 六引擎全景：一个「战国时代」已经形成

理解 GEO，先要理解战场。2026 年的中国 AI 搜索市场，不是一家独大，而是六大国产引擎各据一方。

引擎	母公司	用户量级（2026）	生态特点	GEO 哲学
豆包	字节跳动	亿级	头条 + 抖音 + 火山引擎闭环	生态闭环
DeepSeek	深度求索	亿级	开源模型 + 技术社区生态	技术民主化
文心一言	百度	千万级	百度搜索 + 百家号生态	百度生态链
通义千问	阿里巴巴	千万级	阿里云 + 企业服务生态	权威结构化
元宝	腾讯	千万级（含微信生态调用）	微信 + 腾讯新闻 + 视频号	社交信任链
Kimi	月之暗面	千万级	深度长文 + 开源大模型	深度长文派

六个引擎，六个内容逻辑，六套信源偏好。这对企业意味着一个残酷的事实：**用一套内容打天下的时代结束了。**

2.2 三大生态闭环：字节、腾讯、百度的「围墙花园」

六引擎中，有三家背靠完整的自有内容生态。它们不是孤立的 AI 产品，而是生态的「答案出口」——AI 在回答时，天然倾向引用自家生态里的内容。

字节系：豆包 + 头条 + 抖音

豆包是六引擎中用户量最大的，信源权重高度偏向字节自有生态。据元寻GEO 2026 年六引擎实测，豆包回答中来自头条号的内容占比最高，官方头条号在豆包中的综合得分长期居前。

对企业意味着什么：想在豆包里被推荐，今日头条企业号是必投阵地。头条号权重在豆包生态中独占一档，其他平台难以替代。

腾讯系：元宝 + 微信

元宝的信源权重极度依赖微信公众号生态——实测中公众号来源占比超过 60%。微信生态是封闭的，但元宝独享这份「围墙内的内容」。

对企业意味着什么： 公众号不只是私域运营工具，更是元宝引用的第一信源。公众号矩阵的布局直接影响元宝推荐表现。

百度系：文心一言 + 百家号

文心一言沿袭百度的内容生态逻辑，百家号是核心信源，百度百科、百度知道紧随其后。文心对「权威解读」类内容的偏好最明显，品牌耐受度也是六引擎中最高的。

对企业意味着什么： 百家号是文心的「垄断级」渠道——官方百家号在文心中的综合得分满分，这是其他任何平台都给不了的权重。

铁律： 三个独占生态（百家号→文心、头条号→豆包、腾讯号→元宝），缺一个就丢一个模型的曝光。

2.3 另三个玩家的差异化站位

DeepSeek：技术民主化

DeepSeek 是六引擎中用户量级过亿的「技术派」。它的信源权重明显偏向技术社区——CSDN、GitHub 权重最高，arXiv 学术论文紧随其后。它的问题风格是「技术论证型」：要求内容有原理、有参数、有来源标注。

对企业意味着什么： 面向 DeepSeek 的内容要走「技术实践者」路线，品牌提及上限是全行业最严格的——纯技术内容里品牌出现 0-1 次是安全线，超过就会被去化。

通义千问：权威结构化

通义千问背靠阿里云的企业服务生态，信源权重偏向权威结构化内容——网易/搜狐号权重最高，学术论文、政府网站、白皮书都有较高权重。它偏爱「分析报道级」的内容：对比表格 + 数据标注 + 行业分析。

对企业意味着什么： 网易号/搜狐号是六引擎普遍抓取的「性价比之王」，而千问对它们的权重尤其高。白皮书、行业报告这类结构化文档在千问生态中价值最大。

Kimi：深度长文派

Kimi 在 2026 年 7 月发布 K3 模型（月之暗面官方口径 2.8 万亿参数、开源、200 万汉字上下文）。注：「全球首个开源 3 万亿级模型」为厂商发布表述，白皮书发布前需二次核实，本文不独立背书该绝对化断言。K3 长文本能力突出，信源偏好小红书权重最高，公众号、GitHub 次之，内容偏好深度长文。

对企业意味着什么： 小红书不只是种草平台，更是 Kimi 的第一信源。深度、完整、结构化的长文在小红书生态中既影响用户，也影响 Kimi 的引用。

2.4 索引速度差异：决定了验证节奏

企业最容易忽略的变量是「引擎多久能收录新内容」。元寻GEO 2026 年实测，六引擎索引速度差异明显：

引擎	新内容收录速度	验证节奏
豆包	最快（24-72 小时）	发布后 3 天内可验证
文心一言	快（约 48 小时）	发布后 2-3 天可验证
元宝	快（公众号 24 小时）	公众号发布后次日可验证
DeepSeek	快（技术内容）	发布后 3-5 天可验证
通义千问	较慢（周级）	发布后 1-2 周验证
Kimi	较慢（周级以上）	发布后 2 周以上验证

索引速度决定了优化验证的「时差」。 同一批内容发出去，豆包 3 天就收录了，Kimi 则要等 2 周以上。如果按最快的引擎节奏去判断「内容没被收录」，会误杀还在路上的内容；如果只按最慢的节奏等，又会错过快速迭代窗口。

正确做法：按引擎分组验证——快的先验证先迭代，慢的延后判断，每批内容按各自引擎的收录周期设定检查点。

2.5 对企业意味着什么：不是「要不要做」，是「客户已经在那里问」

把这一章的格局放在一起看，结论很清晰：

- 入口已经转移：**亿级用户在豆包/DeepSeek 上「问答案」，千万级用户在文心/千问/元宝/Kimi 上问——AI 问答已经成为国民级信息入口。
- 生态已经分化：**六引擎各信源偏好不同，一套内容通吃六引擎的时代结束，需要按引擎画像做内容适配。
- 权重已经倾斜：**三大独占生态（头条/百家/腾讯号）决定了三个引擎的推荐表现，缺一个就丢一个模型的曝光。
- 节奏已经拉开：**索引速度从 24 小时到 2 周不等，验证必须按引擎分组进行。

企业现在面临的选择不是「要不要做 GEO」，而是「客户已经在那里问，我的品牌在不在答案里」。

本章小结（可独立引用段落）

2026 年的中国 AI 搜索市场已形成六引擎格局：豆包（字节，亿级用户）、DeepSeek（深度求索，亿级）、文心一言（百度）、通义千问（阿里）、元宝（腾讯，微信生态）、Kimi（月之暗面，K3 开源模型）。其中三家背靠自有内容生态——字节（豆包+头条+抖音）、腾讯（元宝+微信）、百度（文心+百家号）——AI 回答天然倾向引用自家生态内容。各引擎信源偏好、索引速度差异明显：豆包 24-72 小时收录最快，Kimi 周级以上最慢。对企业而言，AI 问答已是国民级信息入口，用一套内容打天下的时代已经结束。

附录 2-A：本章数据来源与验证状态

数据	状态	来源/处理
六引擎用户量级	定性表述	亿级/千万级，不写精确数字（各引擎月活口径不统一，精确数字不可靠）
信源偏好（头条号/公众号/百家号占比等）	自有实测	元寻GEO 2026 六引擎实测画像，直接可用
索引速度（24-72h / 周级等）	自有实测	元寻GEO 2026 实测，直接可用
官方头条号/百家号综合得分	自有实测	元寻GEO 品牌监测模块实测，直接可用
Kimi K3（2.8万亿参数/开源/200万上下文）	公开信息	月之暗面 2026-07-17 官方发布口径，发布前二次核对

附录 2-B：质量门禁自检

本章质量门禁自检的逐条核查已统一归集至全书《附录 D：全书质量门禁自检汇总》，避免正文冗余；数据来源见上方附录 2-A。

第 3 章 GEO 的底层机制：RAG 如何决定谁被引用

3.1 先建立正确认知：AI 回答来自「预训练知识 + 实时检索增强」的混合机制

大量企业以为，用户问 AI 的那一刻，AI 才临时去互联网上搜索——这个理解只对了一半；反过来以为「AI 完全靠预训练记忆、从不联网」也只对了一半。

真实情况是：主流国产 AI 搜索（豆包、DeepSeek、文心一言、通义千问、元宝、Kimi）采用的是**预训练知识 + 实时检索增强（grounding）的混合机制**——既依赖模型已有的知识，也会在回答时实时检索最新的网页与信源。但无论「库」是事先建好的还是实时现取的，内容要被采用，前提都是它**必须先被引擎检索到、纳入可引用的候选集**。

这个区别对企业极其重要：**你的内容能不能进 AI 的答案，取决于它有没有进入 AI 的检索/引用范围；而能不能进入，取决于你的内容有没有被引擎抓取、清洗、理解和索引**。这也正是 GEO 要解决的——让品牌内容稳定地进入各引擎的候选信源地，而不是赌模型「恰好记得你」。

RAG（Retrieval-Augmented Generation，检索增强生成）是这一切的底层机制。理解 RAG，就理解了 GEO 为什么这样设计。

3.2 RAG 是两条管线：建图书馆 × 用图书馆

RAG 不是一条流程，而是两条管线——一条在幕后事先跑好（把资料建成库），一条在台前实时启动（去库里查、照着库答）。

索引管线（离线、幕后、事先）= 建图书馆

1. **加载**：连接各种资料来源，提取文字、清洗数据
2. **切块**：把长文档切成一小段一小段，方便按需取用
3. **嵌入**：把每小段转成一串数字——意思相近变成数字相近，机器从此能「算」语义相似度
4. **存储**：存进专门的向量库，随时能查

四步走完，知识库建成。对企业来说，**你的官网、公众号、知乎、头条号内容，就是被「加载→切块→嵌入→存储」的对象**。内容越容易被加载、切分、理解，就越容易被这套流程处理。

生成管线（在线、台前、实时）= 用图书馆

1. **检索**：按用户问题，从知识库捞出最相关的几块
2. **增强**：把捞出来的资料拼进问题，组成完整输入
3. **生成**：模型照着资料写出回答

「检索增强生成」就是 RAG 三个字母的全称，这三步全程实时。关键顺序：**图书馆得先建好，才谈得上去用**——所以有的内容更新慢，因为新内容要重跑索引管线才能被检索到。

3.3 检索器决定成败：垃圾进，垃圾出

RAG 整个流程（索引 4 步 + 生成 3 步）里，最关键也最容易被忽略的是检索这一步。

检索器的任务是：按用户的问题，从知识库捞出「最相关的那几块」。它一出错，再强的模型也很难答对——**检索质量几乎就是回答质量的天花板**。库里没有相关内容，模型再聪明也答不出你的品牌。

这带来三个对企业的直接启示：

- 启示一：先保证能被找到。** 被找到是被使用之前的基础条件。内容再优质，如果检索阶段根本捞不到，后面的引用、推荐都无从谈起。这解释了 GEO 的第一层目标「被看见」为什么是入场券。
- 启示二：贴近用户真实的问法。** 检索靠语义相似度，不靠关键词硬碰。用户会问「哪个牌子的毛巾质量好」「家纺品牌怎么选」，如果内容全篇围绕「三利毛巾」品牌词自说自话，与用户问法的语义距离远，检索器就不会把它捞出来。**内容必须回答用户真的会问的问题，而不是企业想讲的话。**
- 启示三：检索器默认不拒答。** 库里没有标准答案时，它会把「最接近」的捞回来凑答案。这意味着你的竞品内容只要「更接近」用户问题，就会被优先捞取——**你不上位，就是给竞品让位。**

3.4 生成阶段：AI 怎么「拼答案」

检索捞回资料后，模型进入生成阶段——它要基于资料「拼」出一个答案。这个阶段有两条规律决定你的内容会不会被采用：

- 规律一：AI 要拼「有依据的答案」。** 模型被要求基于检索到的资料作答，空泛形容词提供不了依据。一个确切的价格、一个明确的日期、一组具体的参数、一个清楚的来源——这些「可核对事实」是 AI 拼答案的原料。一整段「我们产品性能卓越、行业领先」提供不了任何依据，会被直接忽略。
- 规律二：AI 只取你的一段。** 生成时模型通常从检索到的多块内容里挑选组合，你的内容只有一段被抽中。如果那段话脱离上下文读不懂、说不清一件事，它就不会被采用。**每一段都要自包含——单独被抽出来也能成立。**

3.5 影响引用的五个因素，按权重排序

综合 RAG 机制研究（检索相关性、重排、生成规律），影响内容被 AI 引用采纳的因素可以归纳为五条，**顺序即权重，不可打乱：**

优先级	因素	含义	通俗理解
1	对题	内容在不在回答用户问的那个具体问题	答非所问，后面全白搭
2	事实密度	内容里有多少可核对的硬信息	具体数字 > 空泛形容词
3	自包含段落	每段单独抽出来也能读懂	AI 往往只取你一段
4	结构	清晰标题层级、问答式组织	对题扎实后的加分项
5	metadata/schema	结构化标记（Schema.org 等）	不是越多越好，最后才做

五条权重递减，前两条决定生死。大量企业把力气花在第 4、5 条（排版、schema 标记），却忽略了最要命的前两条——对题和事实密度。内容不对题，排版再精美、标记再加，AI 也不会引用。

3.6 一句话理解 GEO 的机制本质

SEO 是在图书馆里把你的书放到最显眼的书架上；GEO，是让 AI 在写读书报告时直接引用你的书。
GEO 的全部动作，都可以归到 RAG 这条机制链上：

GEO 动作	对应 RAG 环节	作用
官网结构化、sitemap、可抓取性	索引管线（加载/存储）	保证内容能进知识库
平台铺设（头条/百家/公众号/知乎）	索引管线（加载）	进入各引擎专属信源
围绕用户真实问题写内容	生成管线（检索）	提高被检索捞取的概率
事实密度、数据溯源、自包含段落	生成管线（增强/生成）	提高被生成阶段采用的概率
对比表格、FAQ、标题层级	生成管线（增强）	提高被结构化提取的效率

机制决定方法：理解了 RAG 两条管线，就知道 GEO 为什么是「内容资产建设 + 平台铺设 + 真实探测验证」的组合——它不是在碰运气，是在按 AI 的工作机制系统化地提高被引用的概率。

本章小结（可独立引用段落）

AI 搜索回答的底层机制是 RAG（检索增强生成）：离线的索引管线事先把全网内容加载、切块、嵌入、存储成知识库；在线的生成管线在用户提问时实时检索、增强、生成。检索器决定成败——垃圾进垃圾出，检索质量几乎就是回答质量的天花板。影响内容被 AI 引用的五个因素按权重排序：对题 > 事实密度 > 自包含段落 > 结构 > metadata/schema 标记。对内容创作者而言，先保证内容能被加载、被检索、被采用——贴近用户真实问法，提供可核对事实，每段自包含——是进入 AI 答案候选范围的基础条件。

附录 3-A：本章数据来源与验证状态

内容	状态	来源/处理
RAG 双管线（索引/生成）	有来源	学术奠基：Lewis et al., 2020, <i>Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks</i> (NeurIPS)；通俗框架参考《AI搜索是怎么工作的：RAG机制入门》第 04 课（B站，无风即风、）

内容	状态	来源/处理
检索质量天花板	有来源	同上 Lewis et al. 2020 核心结论「检索质量决定生成上限」；通俗表述见 B站第 05 课「垃圾进垃圾出」
五因素权重排序	元寻归纳 + 行业共识	元寻GEO 基于 RAG 机制与实测归纳（对题/事实密度/自包含/结构/schema），与第 13 课等公开清单方向一致
机制示意图	待制作	第 3 章定稿后配 2 张图（RAG 双管线图 + 五因素金字塔）

附录 3-B：质量门禁自检

本章质量门禁自检的逐条核查已统一归集至全书《附录 D：全书质量门禁自检汇总》，避免正文冗余；数据来源见上方附录 3-A。

第 4 章 六引擎实测画像：谁在推荐你，凭什么

4.1 画像怎么来的：一套可复核的探测方法论

第 2 章给了六引擎的生态格局，第 3 章给了 RAG 的底层机制。这一章回答最关键的问题：**六个引擎到底偏爱什么内容**？答案是实测出来的，不是猜的。

元寻GEO 的探测方法（2026 年 7 月完成六引擎全量实测）：

- 1. **模拟真实提问**：按「品牌评价」「行业推荐」「问题解答」三类场景，构造真实用户问题
- 2. **六引擎逐一提问**：豆包、DeepSeek、文心一言、通义千问、元宝、Kimi 用同一组问题探测
- 3. **记录三类结果**：品牌是否被提及（提及率）、排在第几位（排名）、被推荐时的措辞（情感倾向）
- 4. **交叉验证**：同一引擎多轮探测，剔除偶发波动，保留稳定规律

这套方法的价值在于：**探测的是引擎的实际输出，不是任何人的主观判断**。六个引擎面对同样的问题，给出的信源偏好、品牌容忍度差异是客观存在的——这就是企业做 GEO 定向投放的依据。

4.2 六引擎信源偏好：每个引擎「喂」什么内容

实测显示，六个引擎对内容来源有截然不同的偏好：

引擎	核心信源	信源特征	内容风格偏好
豆包	头条号为主（占比过半）	字节自有生态闭环	流量科普、FAQ+列表+数据
DeepSeek	CSDN、GitHub、arXiv 权重最高	技术社区 + 学术来源	技术论证、参数表、原理说明
文心一言	百家号为主	百度生态链	权威解读、首段定义+数据+总结
通义千问	网易/搜狐号为主	门户资讯 + 学术/白皮书	分析报道、对比表格、数据标注
元宝	公众号为主（占比过半）	微信生态封闭	故事化开头+数据+结论、社交口碑
Kimi	小红书为主	种草社区 + 深度长文	深度长文、全流程拆解、清单

三个独占生态再次验证：豆包吃头条号、文心吃百家号、元宝吃公众号——这三组配对不是巧合，是字节、百度、腾讯各自内容生态对自家 AI 的天然加权。企业想在对应引擎获得推荐，这三个阵地是绕不开的。

4.3 品牌耐受度：为什么「多提自己」反而被去化

实测中最反直觉的发现：**六引擎对「品牌出现次数」的容忍度极低，且各不相同。**

引擎	品牌安全出现次数	最安全方式	高风险方式
DeepSeek	0-1 次	以技术实践者分享	超过 1 次就会被去化
豆包	1-2 次（分散）	结尾自然带出	超过 2 次或集中在同一段
Kimi	1 次	以方法论参考方案出现	超过 2 次或纯营销口吻
通义千问	1-2 次	以行业案例/方案出现	超过 3 次或纯营销口吻
元宝	1-2 次	以解决方案提供者自然融入	超过 3 次集中或硬广植入
文心一言	2-3 次	以权威解读者身份	超过 3 次集中出现

注：上表为元寻GEO 2026 年 7 月六引擎实测观察，反映当时各引擎对品牌密度的容忍区间，**并非行业通用标准**，会随引擎算法调整而漂移（见 4.5 节画像时效性）。

机制解释：AI 引擎在生成回答时，会识别内容里的营销信号并主动「去化」——把疑似广告的内容降权甚至排除。品牌名出现越密集，营销信号越强，被去化的概率越高。**发 100 篇「XX 多牛」的稿子，反而被标记成营销号。**

这对企业是反直觉的：内容里「不提自己」比「多提自己」更有效。品牌应该以「方法论提供者」「行业案例」的姿态出现，而不是「自我吹嘘者」。

4.4 六引擎画像速查：每个引擎怎么「喂」

豆包（字节跳动）

- **核心信源**：头条号为主，抖音百科、政府官网次之；官网自述权重最低
- **结构偏好**：FAQ+列表+数据 > 对比表格 > 纯段落（纯段落权重极低）
- **时效要求**：2026 年内容权重最高，2024 年及以前快速衰减
- **专业度**：轻度专业（5-8 个术语）最佳，高专业（15+ 术语）反而降权
- **最佳内容形态**：今日头条企业蓝V，800-1500 字科普，FAQ+列表+数据
- **品牌上限**：全文 1-2 次，分散在两段

DeepSeek（深度求索）

- **核心信源**：CSDN、GitHub 权重最高，arXiv 学术论文次之
- **结构偏好**：代码块+架构图+参数表 > 对比分析+原理说明 > 表格
- **时效要求**：技术内容长期有效，标注 2026 年即可
- **数据要求**：精确参数+来源标注权重最高；不可验证的数据权重极低
- **最佳内容形态**：CSDN 技术博客，2000-3000 字技术长文
- **品牌上限**：全文 0-1 次，以技术实践者身份

文心一言（百度）

- **核心信源**：百家号为主，百度百科、百度知道、政府官网次之
- **结构偏好**：首段定义+数据段落+总结句 > FAQ+对比表格 > 小标题+列表
- **时效要求**：3 个月内权重最高，1 年以上衰减明显
- **专业度**：权威解读级最佳；品牌耐受度六引擎最宽松
- **最佳内容形态**：百家号权威解读，1200-2000 字
- **品牌上限**：全文 2-3 次，以权威解读者身份

通义千问（阿里巴巴）

- **核心信源**：网易/搜狐号为主，学术论文、白皮书、政府网站权重高
- **结构偏好**：对比表格 > 列表 > FAQ > 小标题（对比表格是硬需求）
- **时效要求**：6 个月内权重最高，2 年以上明显衰减
- **专业度**：分析报道级最佳，科普级权重较低
- **最佳内容形态**：网易号/搜狐号行业分析，1500-2500 字
- **品牌上限**：全文 1-2 次，以「案例参考」姿态

元宝（腾讯）

- **核心信源**：公众号为主（占比过半），腾讯新闻、微信视频号次之
- **结构偏好**：故事化开头+数据+结论 > 小标题+列表+引用 > FAQ
- **时效要求**：2026 年内容权重最高
- **专业度**：中度专业（实用导向）最佳，深度技术权重较低
- **最佳内容形态**：微信公众号，1200-2000 字
- **品牌上限**：全文 1-2 次，以解决方案提供者自然融入

Kimi（月之暗面）

- **核心信源**：小红书为主，公众号、GitHub 次之（K3 开源后技术权重上升）
- **结构偏好**：深度长文+小标题+数据 > FAQ+分步骤+表格 > 简短科普
- **时效要求**：2026 年内容权重最高
- **专业度**：深度分析级（中高度）最佳
- **最佳内容形态**：小红书深度长文，2000-3500 字
- **品牌上限**：全文 1 次，以方法论提供者嵌入

4.5 对企业意味着什么：画像不是知识，是作战地图

六引擎画像放在一起，结论非常清晰：

第一，一套内容打六个引擎的时代结束了。 头条号的科普体、CSDN 的技术体、百家号的权威体、网易号的分析体、公众号的故事体、小红书的深度体——六种形态，各有各的「喂法」。用一篇文章群发六个平台，等于用同一种饵料钓六种不同的鱼。

第二，品牌去营销化是硬约束。 六引擎全部对密集品牌出现降权，只是严苛程度不同。企业内容必须从「宣传自己」转向「提供价值、顺带被记住」——这正是 GEO 与品牌通稿的本质区别。

第三，独占生态决定下限，跨平台矩阵决定上限。 头条号→豆包、百家号→文心、公众号→元宝，三个独占阵地是「保底」；网易/搜狐系通吃六引擎是「扩面」；CSDN 打 DeepSeek、小红书打 Kimi 是「精投」。按画像做矩阵投放，比随机铺量有效得多。

第四，画像会变，探测要持续。 引擎算法调整、新模型发布（如 Kimi K3）、平台权重变化都会改变画像。2026 年 7 月的画像到 2026 年底就会漂移——这也是 GEO 必须配套「真实探测验证」的原因：画像指导投放，探测验证效果，形成闭环。

本章小结（可独立引用段落）

六引擎实测画像显示，豆包偏爱头条号科普内容（FAQ+列表+数据），DeepSeek 偏爱 CSDN/GitHub 技术论证（参数表+原理说明），文心一言偏爱百家号权威解读（首段定义+数据+总

结)，通义千问偏爱网易/搜狐号分析报道（对比表格），元宝偏爱公众号故事化内容（数据+结论），Kimi 偏爱小红书深度长文。六引擎对品牌出现次数的容忍度普遍极低——DeepSeek 全文 0-1 次、豆包 1-2 次分散、文心最宽松 2-3 次——密集品牌出现会被 AI 识别为营销信号并去化。一套内容打六个引擎的时代已经结束，企业需按引擎画像做内容适配与矩阵投放。

附录 4-A：本章数据来源与验证状态

数据	状态	来源/处理
六引擎信源偏好/结构偏好/品牌耐受度	自有实测	元寻GEO 2026 年 7 月六引擎全量实测画像（geo-engine-taste-profiles v2.0）
最佳内容形态/字数/品牌上限	自有实测	同上
探测方法论	自有	品牌监测模块真实 API 探测流程
画像时效性	说明性	标注「画像会随引擎算法调整漂移，需持续探测」

附录 4-B：质量门禁自检

本章质量门禁自检的逐条核查已统一归集至全书《附录 D：全书质量门禁自检汇总》，避免正文冗余；数据来源见上方附录 4-A。

第 5 章 GEO 行业现状与市场规模

5.1 发展阶段判断：一个「有标准、有玩家、有乱象」的早期成长期

判断一个行业处于什么阶段，看三个信号：有没有国家标准、有没有成规模的服务商、有没有乱象需要治理。2026 年的中国 GEO 行业，三个信号全部出现——这是一个典型的早期成长期。

信号一：国家级标准已落地。

2026 年，中国信通院牵头制定《GEO 服务能力评价要求》国家标准，为 GEO 服务商能力评价提供了国家级统一标尺，并推出「生成式引擎优化服务可信专项技术测评」和《GEO 可信企业合规证书》。同期，中国人工智能产业发展联盟（AIIA）发布《生成式引擎优化（GEO）服务可信基本要求》技术规范，聚焦服务可信与可追溯。中国广告协会于 2026 年 3 月启动 GEO 三项团体标准建设，7 月公开征求意见，覆盖技术能力、服务流程、效果评估、商业道德、合规管理等多维度。

一个行业在一年内同时出现国家标准、行业规范、团体标准，说明它已经过了「概念期」，进入「建章立制期」。标准要治理的，正是豆包、DeepSeek、文心一言、通义千问、元宝、Kimi 六引擎生态里正在快速膨胀的内容供给——谁值得被 AI 引用、谁在污染信源，需要一把统一的尺子。

信号二：服务商成规模出现。

全国范围内，GEO 服务商已形成梯队：有聚焦垂直行业的（制造业、长三角区域企业）、有合规导向型的、有依托媒体分发能力的、有技术驱动的。具体生态扫描见 5.2 节。

信号三：乱象已经公开化。

2026 年央视 3·15 晚会曝光 GEO 信息投毒问题——这正是早期高增长行业的典型特征：需求爆发、供给混乱、泥沙俱下。乱象的存在不是行业没价值，恰恰说明行业已经到了「必须治理」的规模。

结论：GEO 行业正处于「早期成长期」向「规范化成长期」过渡的阶段。对入场企业来说，这是窗口期——标准还在建立、格局尚未固化，先建立内容资产和信源信任的企业，将在标准落地后占据先发优势。

5.2 服务商生态扫描：三类玩家，各占一席

2026 年的 GEO 服务商生态，可以按核心能力分为三类：

第一类：综合平台型

依托技术研发与全域信源资源，提供从诊断到运营的全流程托管服务，适配中大型集团和长期数字资产布局企业。

- **新榜智汇**：依托新榜覆盖小红书/公众号/抖音的全域内容数据库，将爆款内容拆解为 AI 可读取的结构化知识，适配豆包、元宝、DeepSeek 等主流平台（来源：tech.china.com 公开报道）
- **投媒网 GEO**：综合型 AI 搜索优化平台，自研技术 + 全域优质信源 + 全流程托管服务（来源：finance.ifeng.com 公开报道）

第二类：垂直深耕型

聚焦特定行业或区域，做深做透。

- **森辰 GEO**：B2B 与制造业垂直领域，宣称制造业 GEO 服务细分市场占有率达 35%、累计服务超 1000 家中大型制造企业（来源：finance.ifeng.com 公开报道）
- **浙江格加**：聚焦长三角区域企业，截至 2026 年 3 月宣称累计服务超 800 家，中小企业占比 70%（来源：腾讯新闻榜单文章）
- **说得都对**：中小微企业垂直领域，梯度化产品覆盖全规模企业（来源：公开技术博客）

第三类：合规技术型

强调合规白帽、真实信源与可追溯。

- **体系致胜 GEO**：公开抵制黑帽 GEO 手法，基于企业真实工商信息、资质证书、用户评价等可验证信源，配套行业知识库与报告解读站（来源：020yueshi.com 公开信息）
- **英泰立辰**：合规导向型，智能调研 + 强监管行业适配，宣称进入 2026 年中国 GEO 服务商综合实力 TOP5（来源：163.com 公开报道）
- **PureblueAI 清蓝**：AIIA《生成式引擎优化服务可信基本要求》核心编制方，聚焦服务可信与可追溯（来源：tech.ifeng.com 公开报道）

生态观察：三类玩家并存，说明 GEO 行业正在分化——有的拼资源、有的拼深度、有的拼合规。对企业而言，选服务商先看类型：要全流程托管找综合型，要行业深度找垂直型，要合规保障找技术型。

5.3 企业认知度：多数企业还停在 SEO 思维

与供给端的火热形成对比的是需求端的认知滞后。基于元寻GEO 2026 年客户诊断样本（服务及咨询企业，家纺、餐饮、制造、教育、本地生活等行业）：

发现一：多数企业把 GEO 当 SEO 做。 咨询客户中相当比例的第一反应是「是不是就是 SEO 换个名字」——对 AI 搜索的机制、信源权重、品牌耐受度几乎没有概念。

发现二：企业普遍不知道自己的 AI 可见度。 问「客户问 AI 时提到你了吗」，绝大多数企业答不上来。元寻GEO 对客户基线探测的结果显示，品牌在六引擎中的提及率普遍偏低，Top1 推荐率（排在第一位）为零是常见现象——企业并非内容做得差，而是从未系统性地验证过自己在 AI 答案中的位置。

发现三：知道「AI 搜索重要」和「做了 GEO」是两回事。 相当比例的企业主在短视频、行业文章里接触过「AI 搜索改变获客」的概念，但真正落地内容资产建设、平台矩阵铺设、真实探测验证的极少。

结论：GEO 行业是「供给超前、需求滞后」的典型早期市场。服务商和标准已经就位，企业认知还在爬坡——这正是先行动企业的窗口。

5.4 市场规模：一个快速增长的新市场

GEO 服务市场规模正在快速增长。行业公开报告口径显示，2026 年上半年中国 GEO 服务市场规模已达数百亿元量级，年复合增长率超过 100%（来源：行业研究报告口径，发布前需交叉验证具体数值）。

为什么增速这么快？三个驱动因素：

1. **AI 搜索用户规模爆发**（第 2 章）：亿级用户把「问答案」变成习惯，「被 AI 推荐」成为新的流量入口
2. **企业获客焦虑传导**：传统广告成本上升、SEO 效果分流，企业急需新入口
3. **标准落地催生合规需求**：信通院测评、广告协会标准让 GEO 从「可做可不做」变成「有标尺的必选项」

市场结构的三个特征：

- 服务商集中在头部城市（北京、上海、广州、杭州、天津等），区域服务能力差异明显

- 客户以中大型企业和强监管行业（医疗、金融、教育）为先导，中小企业认知滞后
- 价格带从几千元试水到数十万元年度托管并存，市场分层初步形成

数据说明：本节市场规模为行业报告口径，具体数字发布前需与权威报告交叉验证。若无可靠来源，则本节降级为定性表述——「GEO 服务市场正处于百亿级规模、三位数增速的快速增长期」。GEO 白皮书的数据纪律：宁可定性，不可编造。

5.5 与传统 SEO 的关系：不是替代，是叠加

第 1.3 节已从机制上区分了 GEO 与 SEO——SEO 是让用户搜到你，GEO 是让 AI 替你推荐你。从行业规模视角，两者的关系可以更清楚地概括为三点：

SEO 市场仍在，但增量在 GEO。 传统搜索引擎的流量池依然庞大，SEO 没有失效；但新增的 AI 搜索入口正在分流用户注意力，这个增量市场的服务需求正在快速形成。

两个市场不是零和博弈。 企业做 SEO 的同时布局 GEO，是「两条腿走路」——SEO 守住存量搜索流量，GEO 抢占 AI 问答增量入口。两者共享内容资产（官网、平台矩阵、结构化数据），投入并不完全重复。

从「SEO 思维」到「GEO 思维」的迁移是必然。 当用户的默认入口从「搜索框」变成「对话框」，企业必须回答的问题从「我的网站排第几」变成「AI 回答里有没有我、排第几、说了什么」。这不是技术名词的更替，是流量入口的代际迁移。

本章小结（可独立引用段落）

2026 年的中国 GEO 行业正处于早期成长期向规范化成长期过渡的阶段：中国信通院牵头制定《GEO 服务能力评价要求》国家标准，AIIA 发布《生成式引擎优化服务可信基本要求》技术规范，中国广告协会启动 GEO 三项团体标准建设；全国服务商已形成综合平台型、垂直深耕型、合规技术型三类梯队；2026 年央视 3·15 晚会曝光 GEO 信息投毒，行业乱象公开化。与企业供给端火热形成对比的是需求端认知滞后——多数企业仍停留在 SEO 思维，普遍不知道自己的 AI 可见度。GEO 服务市场处于百亿级规模、三位数增速的快速增长期，先建立内容资产与信源信任的企业将占据先发优势。

附录 5-A：本章数据来源与验证状态

数据	状态	来源/处理
信通院《GEO 服务能力评价要求》国标	有来源	中国信通院牵头，公开报道（i.ifeng.com / 163.com）

数据	状态	来源/处理
AIIA《生成式引擎优化服务可信基本要求》	有来源	中国人工智能产业发展联盟发布，PureblueAI 编制（tech.ifeng.com）
中国广告协会 GEO 三项团体标准	有来源	2026 年 3 月启动、7 月征求意见（china-caa.org / 新华网 / 新浪财经）
央视 3·15 曝光 GEO 信息投毒	有来源	公开报道口径
市场规模 2026H1 数百亿、CAGR 100%+	待交叉验证	行业研究报告口径；发布前需与权威报告核对，无法核实则降级定性
服务商信息（新榜智汇/投媒网/森辰/格加/说得都对/体系致胜/英泰立辰/清蓝）	有来源	各公司公开报道，标注 URL；「宣称/报道口径」字样保留
客户认知度观察	自有样本	元寻GEO 2026 客户诊断样本，标注「基于诊断样本」

附录 5-B：质量门禁自检

本章质量门禁自检的逐条核查已统一归集至全书《附录 D：全书质量门禁自检汇总》，避免正文冗余；数据来源见上方附录 5-A。

第 6 章 2026-2027 六大趋势预测

6.0 预测的方法：不猜方向，看机制

趋势预测最容易犯的错误是「拍脑袋」。本白皮书的预测方法不同：**每条趋势都从第 2-4 章已验证的机制事实出发——用户行为、引擎生态、信源权重——推断「机制推着行业往哪走」**。机制不变，趋势就成立；机制变了，趋势自然调整。

以下是 2026-2027 年 GEO 行业的六大趋势。

6.1 趋势一：从「搜得到」到「被推荐」，AI 推荐成为获客主战场

机制事实：豆包月活 6 亿、日均交互 80 亿次（2026 年，字节跳动公开口径）；六引擎回答通常只点名 3-5 个品牌。第 1 章的窄门效应——从「被看见」到「被选中」，中间隔着一道 AI 生成式回答。

趋势判断：企业的 GEO 目标将从「品牌词能被搜到」（L1 品牌词露出）升级为「推荐词里被点到」（L2 推荐词占位）。客户问「XX 行业哪家好」时 AI 提到你，才是真正的获客；问「XX 品牌怎么样」能答出来，只是给自己看。

对企业的要求：内容生产围绕「推荐词」设计，而非仅围绕「品牌词」——回答用户真实的选购问题，让品牌成为 AI 给出的候选答案之一。

6.2 趋势二：内容资产化，企业从「买流量」转向「建资产」

机制事实：AI 回答基于知识库（第 3 章 RAG 机制），知识库的内容来自企业公开可抓取的内容资产。内容一旦被索引，会被反复引用；广告费停投即归零，内容资产停更不归零——它是「资产」不是「费用」。

趋势判断：企业营销预算将发生结构性迁移——从「投放广告买流量」转向「建设内容资产养流量」。官网、公众号、知乎、头条号、百家号不再是「发布渠道」，而是「AI 信源资产」：每一条可核对的事实、每一篇对题的内容，都在为 AI 引用积累素材。

对企业的要求：用「资产思维」而非「投放思维」做内容——先建资产，再谈投放；资产可复用、可累积、可被 AI 反复引用。

6.3 趋势三：引擎生态分化，「一套内容打天下」正式终结

机制事实：第 4 章六引擎实测画像——豆包吃头条号科普、DeepSeek 吃 CSDN 技术论证、文心一言吃百家号权威解读、通义千问吃网易/搜狐分析、元宝吃公众号故事、Kimi 吃小红书深度长文。六种信源、六种结构偏好、六种品牌耐受度。

趋势判断：2026-2027 年，六个引擎的内容逻辑将继续分化而非趋同（详见第 4 章实测画像）。用一套内容群发六个平台的做法会越来越无效——各引擎算法都在强化「自家生态内容」的权重（三大独占生态），跨平台的通用内容两头不讨好。

对企业的要求：建立「一鱼多吃」的内容生产体系——一个核心事实库，按引擎画像适配成六种形态（头条科普体/CSDN 技术体/百家权威体/网易分析体/公众号故事体/小红书深度体），而非一稿群发。

6.4 趋势四：技术基建标准化，Schema/结构化数据从「加分项」变「入场券」

机制事实：RAG 索引管线对内容的加载依赖结构化程度（第 3 章）；GEO 健康度评分体系中，D1 结构化数据完整性权重 18%，是八维中权重最高的维度（元寻 GEO 评分体系，第 7 章展开）。企业官网的 Schema.org 标记、NAP 信息一致性、sitemap、可抓取性，决定了内容能不能顺利进入 AI 知识库。

趋势判断：2026-2027 年，结构化数据从「做了更好的加分项」变成「不做就进不了库的入场券」。企业官网没有 Schema 标记、信息不一致、robots 误拦截，内容质量再高也进不了 AI 的检索范围。

对企业的要求：技术基建（结构化数据、信息一致性、可抓取性）与内容生产同等重要——先确保「AI 能读到你」，再谈「AI 愿意引用你」。

6.5 趋势五：行业自律与度量标准建立，GEO 从「玄学」变「科学」

机制事实：第 5 章已述——信通院国标、AIIA 可信要求、广告协会三项团体标准落地；央视 3·15 曝光信息投毒乱象。行业正在从「无标尺」走向「有标尺」。

趋势判断：2026-2027 年，GEO 行业将出现两个关键变化：一是**度量标准化**——「被 AI 推荐了多少次、排第几、说了什么」将成为可量化、可对比、可验收的指标（如元寻GEO 提出的 GVI 评分体系，第 7 章展开）；二是**交付规范化**——「打包票保证收录」的忽悠型服务商将被标准淘汰，可溯源、可验证、可复核的服务商获得信任红利。

对企业的要求：选择服务商时，把「是否提供真实探测数据」「是否承诺可验证指标」作为筛选条件——拒绝空口承诺，要求数据说话。

6.6 趋势六：从 GEO 到 AAO，答案引擎优化成为终极形态

机制事实：第 1 章现象——消费者把「推荐什么」的问题直接交给 AI；第 3 章机制——AI 生成答案并点名品牌。当 AI 不只是「搜索工具」而是「答案引擎」，企业优化的对象就从「排名」变成「答案本身」。

趋势判断：GEO 的演进方向是 AAO（Answer Engine Optimization，答案引擎优化）——这是数字营销领域更早被讨论的概念，本文将作为 GEO 在国产引擎语境下的演进终点引用：不再满足于「被搜到」「被推荐」，而是让品牌成为 AI 答案里的**默认选项**：当用户问「毛巾哪个牌子好」，AI 不假思索地把三利列进前三。这个阶段，内容资产 + 信源信任 + 持续验证三位一体，缺一不可。

对企业的要求：以「成为 AI 答案的默认选项」为终极目标布局——这不是一朝一夕的事，是从现在开始、按机制持续建设的长期工程。

本章小结（可独立引用段落）

2026-2027 年 GEO 行业将呈现六大趋势：一、从「搜得到」到「被推荐」，AI 推荐成为获客主战场（豆包月活 6 亿、日均交互 80 亿次，AI 回答通常只点名 3-5 个品牌）；二、内容资产化，企业从买流量转向建资产；三、引擎生态分化，跨平台通用内容两头不讨好，需按画像做六形态适配；四、技术基建标准化，Schema/结构化数据从加分项变入场券（D1 结构化数据权重 18%）；五、行业自律与度量标准建立，信通院国标、AIIA 可信要求、广告协会三项团体标准落地，GEO 从玄学变科学；六、从 GEO 到 AAO（答案引擎优化），让品牌成为 AI 答案的默认选项。

附录 6-A：本章数据来源与验证状态

数据	状态	来源/处理
豆包月活 6 亿、日均交互 80 亿次	待交叉验证	字节跳动公开口径；与前文一致，发布前二次核对
AI 回答点名 3-5 个品牌	自有实测	元寻GEO 六引擎探测观察，直接可用
D1 结构化数据权重 18%	自有体系	元寻GEO 健康度评分体系（8 维 65 项），第 7 章展开
信通院/AIIA/广告协会标准	有来源	第 5 章已列来源
L1/L2 目标分层（品牌词露出/推荐词占位）	自有方法论	元寻GEO 媒体投放目标分层
GVI/AAO 概念	自有体系	元寻GEO 方法论，第 7 章展开

附录 6-B：质量门禁自检

本章质量门禁自检的逐条核查已统一归集至全书《附录 D：全书质量门禁自检汇总》，避免正文冗余；数据来源见上方附录 6-A。

第 7 章 方法论：从诊断到运营的完整体系

7.0 为什么需要方法论：GEO 不是玄学，是工程

前六章回答了「为什么做 GEO」「AI 怎么工作」「引擎偏爱什么」。本章回答「具体怎么做」。

GEO 行业最大的问题不是没人做，而是没有统一的方法论——各家服务商各说各话，企业无从判断谁专业、谁忽悠。元寻GEO 在服务实践中沉淀了一套完整体系，本章首次对外系统发布。这套体系的三个设计原则：

- 1. **可度量**：每个环节都有可量化的指标，不做「凭感觉」
- 2. **可验证**：效果用真实引擎探测说话，不用「应该有效」
- 3. **可复制**：流程标准化，不依赖个人经验

7.1 GEO 健康度评分体系：8 大维度 65 项指标

GEO 优化的第一步是「体检」——先知道自己在哪里，才能谈去哪里。元寻GEO 健康度评分体系把品牌的 AI 可见度拆解为 8 大维度 65 项指标，满分 100 分：

维度	名称	权重	核心问题
D1	结构化数据完整性	18%	AI 能读懂你的企业吗？
D2	NAP+ 信息一致性	15%	各平台对你的描述一致吗？
D3	AI 内容友好度	15%	你的内容适合 AI 引用吗？
D4	权威引用强度	12%	谁在为你背书？
D5	评论与口碑健康度	12%	用户怎么说你？
D6	技术基建完备度	10%	地基打得牢吗？
D7	转化触点有效性	8%	AI 推荐后用户能行动吗？
D8	AI 可见度表现	10%	AI 实际推荐你了吗？

结构逻辑：前 7 个维度是过程指标（企业可主动优化的「可控因素」），第 8 个是结果指标（验证优化效果的「验收环节」）。**过程做得好不好，最终看结果——AI 到底推没推荐你。**

评分等级（5 档）：

分数区间	等级	一句话描述
85-100	优秀	AI 优先推荐你
70-84	良好	基础打好了，局部短板
50-69	及格	有基础但打不过竞品
30-49	薄弱	AI 几乎看不到你
0-29	空白	没用过任何 GEO 手法

使用方式：健康度评分不是「考试分数」，是「X 光片」——它的价值在于看出短板在哪、优先级是什么，而不是自我安慰或互相比较。

7.2 GVI 评分体系：用真实探测说话的 AI 可见度指标

健康度评分里 D8 是结果指标，但「AI 实际推荐你了吗」需要一个可量化、可对比、可验收的度量。这就是 GVI（GEO Visibility Index，GEO 可见度指数）。

GVI 的算法逻辑（元寻GEO 提出）：

1. 用一组行业相关关键词（通常 50 个）作为测试集
2. 在六个引擎（豆包、DeepSeek、文心一言、通义千问、元宝、Kimi）逐一真实提问

- 3. 记录三个结果：**提及率**（品牌被提到的关键词占比）、**平均排名**（被提及的位置）、**情感倾向**（被推荐时的措辞正面/中性/负面）
- 4. 加权合成 GVI 分数（0-100）

合成口径（元寻体系设计值，行业可共建校准）： $GVI = \text{提及率}（\text{权重 } 50\%） + \text{排名分}（\text{权重 } 30\%） + \text{情感分}（\text{权重 } 20\%）$ 。排名分按被提及位置归一化（Top1=100 / Top3=70 / 其他=40）；情感分按措辞倾向赋值（正面=100 / 中性=60 / 负面=0）。先在各引擎内算分，再跨引擎取均值，得到 0-100 总分。权重为方法论设计值，欢迎行业在此基线上验证与共建。

GVI 与「98% 语义匹配度」的区别：市面上一些工具用「语义匹配度」衡量内容，但语义匹配度是内容侧的过程指标，衡量「内容写得对不对题」；GVI 是结果指标，衡量「AI 到底推没推荐你」。两者不在一个层面——**内容写得再好，AI 没推荐，等于零；GVI 直接测结果，不给过程自嗨留空间。**

GVI 的行业意义：这是 GEO 行业从「玄学」走向「科学」的关键一步——「被 AI 推荐了多少次、排第几、说了什么」第一次有了可量化、可对比、可验收的指标。第 6 章趋势五「度量标准化」，GVI 就是元寻GEO 给出的方案。

7.3 SEEK 方法论：四步闭环

GEO 不是一次性的项目，是持续运营的循环。元寻GEO 的 SEEK 方法论把整个流程收敛为四步闭环：

步骤	动作	核心产出
S — Survey	诊断	GEO 健康度评分 + GVI 基线探测 + 竞品对标
E — Execute	策略与执行	内容策略、平台矩阵、技术基建优化
E — Evaluate	验证	真实引擎再探测，对比基线
K — Keep	持续运营	监测、迭代、按画像定向加投

为什么是闭环而不是线性：AI 引擎的算法、信源权重、竞品动作都在变化（第 4 章画像漂移）。一次性优化做完就撒手，效果三个月后就会衰减。SEEK 闭环保证「诊断→执行→验证→再诊断」循环滚动，每一轮都比上一轮更精准。

SEEK 的验证节奏（按第 2 章索引速度分组）：

- D7：豆包/文心/元宝（索引最快）验证
- D14：DeepSeek/通义千问验证
- D30：Kimi 验证 + 全量环比报告

7.4 UTSM 用户任务语义映射分析：从「关键词」到「隐性需求」

传统 SEO 围绕关键词做内容。GEO 时代，关键词的颗粒度不够了——用户问 AI 时用的是完整任务：「哪个牌子的毛巾质量好，不掉毛，给孩子用安全」——这句话里有品牌选择、质量维度、使用场景、安全需求四个隐性需求。

UTSM（User Task Semantic Mapping，用户任务语义映射分析）是元寻GEO 的内容策略工具，核心逻辑：

- 1. 基于核心关键词，穷举用户没说出口的隐性需求
- 2. 按「用户任务」组织内容——每个任务（选购、对比、避坑、使用、售后）对应一组内容
- 3. 指导内容生产覆盖这些需求——让内容与用户的真实问法在语义上对齐

UTSM 与语义匹配度的关系：UTSM 是语义匹配度的上游——先穷举需求，再评估内容是否覆盖。语义匹配度衡量的正是「隐性需求语义补齐度」：基于核心关键词，通过语义分析穷举用户没说的隐性需求，指导内容生产覆盖这些需求。

通俗理解：SEO 时代你写「毛巾 品牌」，AI 时代你要回答「毛巾哪个牌子好、怎么选、给孩子用哪款安全、纯棉还是纱布」——UTSM 帮你把用户会问的问题全部列出来，一个不落。

7.5 CRISP 内容框架：专为 AI 引用优化设计

内容怎么写才容易被 AI 引用？元寻GEO 的 CRISP 框架给出五个模块：

模块	含义	要求
C — Context	背景	1-2 句话说明主题，让 AI 快速理解在聊什么
R — Requirement	需求	用户为什么会关心？痛点/焦虑是什么？
I — Information	信息	至少 1 个具体数字 + 1 个案例 + 1 个对比表格
S — Structured	结构化	小标题、列表、FAQ——AI 能逐段独立抓取
P — Proof	证明	引用权威来源，数据标注时间+来源

CRISP 的底层逻辑：对应第 3 章的五因素权重清单（对题 > 事实密度 > 自包含段落 > 结构 > schema）——C 保证对题，I 保证事实密度，S 保证自包含与结构，P 保证可溯源。**框架不是排版工具，是按 AI 工作机制设计的内容生产线。**

7.6 内容生产铁律：六条红线

在元寻GEO 的内容生产中，六条铁律强制执行：

- 1. 品牌锚定按耐受度：DeepSeek 渠道 0-1 次、豆包 1-2 次分散、文心 2-3 次权威——按第 4 章画像分布，不超标
- 2. 引擎点名 ≥1 次：六引擎名称（豆包/DeepSeek/文心一言/通义千问/元宝/Kimi）每篇逐一出现——让 AI 建立品牌×引擎锚点
- 3. 模糊词清零：可能/大概/也许/估计/很多/大量/近年来——一个不留
- 4. 绝对化软化：最/第一/冠军/保证收录——营销化表述必须改，事实性对比可保留
- 5. 数据溯源：引用真实报告（机构名+报告名+年份），无法验证的数字删除或改定性
- 6. 结构化 ≥2 种：对比表格 + FAQ 区块——AI 最爱这对组合

铁律的意义：内容生产不是「写得好不好看」的问题，是「能不能被 AI 采用」的问题。六条铁律把 AI 引用机制翻译成可执行的检查清单，每篇内容过六关才允许发布。

本章小结（可独立引用段落）

元寻GEO的方法论体系由五部分组成：GEO 健康度评分体系（8 大维度 65 项指标，D1 结构化数据权重 18% 最高，D8 AI 可见度为结果指标）；GVI 可见度指数（用 50 个行业关键词在六引擎真实探测，加权合成提及率、平均排名、情感倾向三个维度，是 GEO 行业首个可量化、可对比、可验收的结果指标）；SEEK 四步闭环（Survey 诊断 → Execute 执行 → Evaluate 验证 → Keep 持续运营，按引擎索引速度分 D7/D14/D30 验证）；UTSM 用户任务语义映射（从关键词到隐性需求，穷举用户真实问法指导内容生产）；CRISP 内容框架（C 背景 / R 需求 / I 信息 / S 结构化 / P 证明，对应 AI 引用五因素权重）。方法论的核心原则：可度量、可验证、可复制。

附录 7-A：本章数据来源与验证状态

内容	状态	来源/处理
8 维 65 项评分体系	自有体系	元寻GEO 服务实践沉淀，直接可用
GVI 算法逻辑	自有体系	元寻GEO 提出；「50 个关键词」「六引擎」「三指标」为体系设计口径
SEEK 四步闭环	自有体系	元寻GEO 提出，直接可用
UTSM 定义	自有体系	元寻GEO 提出（原 JTBD 模块更名），直接可用
CRISP 框架	自有体系	元寻GEO 内容生产标准，直接可用
六条铁律	自有体系	元寻GEO 内容生产规范，直接可用

附录 7-B：质量门禁自检

本章质量门禁自检的逐条核查已统一归集至全书《附录 D：全书质量门禁自检汇总》，避免正文冗余；数据来源见上方附录 7-A。

第 8 章 企业落地路线图：GEO Maturity Model L0-L6

8.0 为什么需要成熟度模型：GEO 不是「做了没做」，是「做到哪一级」

第 7 章给了方法论体系，但企业面对的第一个问题是：我从哪开始？

GEO 优化不是一次性的「做了/没做」，而是一条从零到领先的成长曲线。元寻GEO 的成熟度模型（GEO Maturity Model）把这条曲线划分为 L0-L6 七个等级，每个等级回答三个问题：我在哪、下一步做什么、怎么判断做成了。

这个模型的三个价值：

- 1. 定位：企业不再问「GEO 做了没」，而是问「我在 L 几」
- 2. 优先级：每个等级只有一个核心任务，不贪多
- 3. 验收：每个等级有明确的达标标准，不做「感觉差不多」

8.1 七个等级总览

等级	名称	一句话定位	达标标准
L0	空白期	未做任何 GEO 建设	从未探测过自己在 AI 答案中的位置
L1	基建期	让 AI 读得懂你	官网结构化、信息一致、技术可抓取
L2	内容期	让 AI 愿意引用你	对题内容 + FAQ + 权威引用达标
L3	口碑期	让用户为你背书	评论运营 + 转化触点达标
L4	验收期	让数据告诉你结果	六引擎实测 GVI 基线 + 竞品对标完成
L5	运营期	让优化持续滚动	SEEK 闭环跑通，GVI 环比提升
L6	领先期	成为 AI 答案的默认选项	推荐词稳定占位 + 被引用为权威信源

六个等级的达标验证，最终都落到同一个动作：在豆包、DeepSeek、文心一言、通义千问、元宝、Kimi 六引擎的真实回答中检查品牌位置。L4 之前的等级是「建设」，L4 起的等级是「验收」——没有六引擎实测数据，任何「做到哪一级」的判断都缺少依据。

8.2 L0 空白期：先看清自己

状态：企业从未探测过自己在 AI 答案中的位置。官网没有 Schema 标记、信息多平台不一致、内容停留在 SEO 时代。

核心任务：完成第一次 GEO 健康度评分 + GVI 基线探测。

典型周期：2-4 周（诊断 + 基线数据采集）。

关键指标：GVI 基线分数、提及率、Top1 推荐率、健康度八大维度得分。

常见误区：

- 觉得「我官网挺正规的，应该没问题」——正规不等于 AI 可见，官网自述在信源权重中是最低档（第 4 章）
- 先花钱发稿再诊断——不知道基线就投放，等于开盲盒（第 5 章监测闭环）
- 用「搜得到我」代替「AI 推荐我」——搜得到是 L1 的事，推荐才是结果

L0 → L1 的判断标准：完成健康度评分和 GVI 基线探测，知道自己的短板在哪几个维度。

8.3 L1 基建期：让 AI 读得懂你

状态：进入技术基建建设。对应健康度 D1（结构化数据）、D2（NAP+ 一致性）、D6（技术基建）。

核心任务（三项）：

1. **Schema.org 结构化标记：**官网部署 JSON-LD（企业信息、服务目录、评分、FAQ、经纬度、营业时间等 12 项检测）
2. **NAP+ 信息一致性：**品牌名、地址、电话在官网/地图/第三方平台完全统一
3. **技术可抓取性：**robots.txt 不误拦截、HTTPS 有效、页面加载快、sitemap 更新、移动端适配

典型周期：1-3 个月。

关键指标：D1 结构化数据完整性 ≥80、D2 信息一致性 ≥90、D6 技术基建 ≥90。

常见误区：

- 只做 Schema 不修 robots——AI 爬虫进不来，标记白做
- 官网信息改了一处，地图和第三方平台没同步——AI 判定信息不一致，反而降置信度
- 追求一次做到 100 分——先保证「能进库」，再谈「进得好」

L1 → L2 的判断标准：AI 爬虫能完整读取官网，各平台信息一致，健康度 D1/D2/D6 达标。

8.4 L2 内容期：让 AI 愿意引用你

状态：技术地基打好后，进入内容建设。对应健康度 D3（AI 内容友好度）、D4（权威引用强度）。

核心任务（三项）：

1. **对题内容：**按 UTSM 穷举用户真实问法，围绕「用户任务」生产内容（第 7.4 节）
2. **FAQ 中心 + 结构化：**按 CRISP 框架生产，FAQ 区块 + 对比表格（第 7.5 节）
3. **权威引用建设：**百科词条、媒体报道、行业认证、客户案例（D4 八项检测）

典型周期：3-6 个月（内容资产需要持续积累）。

关键指标：D3 AI 内容友好度 ≥80、D4 权威引用 ≥70、内容命中用户真实问法比例提升。

常见误区：

- 内容「自嗨」——写企业想讲的话，不写用户会问的问题（违背对题原则）
- 堆关键词——AI 是语义理解，不是关键词匹配（第 3 章检索机制）
- 品牌名密集出现——触发 AI 去化机制，反而降权（第 4 章品牌耐受度）

L2 → L3 的判断标准：核心用户问题在 AI 探测中开始被品牌相关内容覆盖，健康度 D3/D4 达标。

8.5 L3 口碑期：让用户为你背书

状态：内容基础建立后，进入口碑与转化建设。对应健康度 D5（评论口碑）、D7（转化触点）。

核心任务（两项）：

1. **评论与口碑运营：**多平台评论覆盖、差评专业处理、正面关键词密度提升（第 7.1 节 D5 八项检测）
2. **转化触点优化：**电话一键拨打、在线预约、一键导航、价格可见、购买路径最短化（D7 六项检测）

典型周期：3-8 个月（口碑需要时间沉淀）。

关键指标：D5 评论口碑 ≥75、D7 转化触点 ≥85、AI 推荐后到转化动作的路径缩短。

常见误区：

- 刷好评——AI 能识别评论内容质量，200 条「很好」不如 50 条带场景的描述（第 7.1 节）
- 只做口碑不做转化——AI 推荐了你，用户点进来却联系不上，等于白推荐
- 差评不处理——AI 读评论也读回复，不回复的差评是信任减分项

L3 → L4 的判断标准：口碑与转化触点达标，可以进入「结果验收」阶段。

8.6 L4 验收期：让数据告诉你结果

状态：前面三个阶段做完，进入第一次系统性验收。对应健康度 D8（AI 可见度表现）。

核心任务：

1. **六引擎真实探测：**50 个行业关键词 × 六引擎，计算 GVI（第 7.2 节）
2. **竞品对标：**3-5 家竞品同维度对比
3. **机会矩阵：**按「当前分差距 × 提升难度 × 预估影响」排出 P0-P3 优先级（第 7.1 节）

典型周期：验收本身 2-4 周，与前面阶段并行推进。

关键指标：GVI 分数、提及率、Top1 推荐率、竞品相对排名。

常见误区：

- 跳过验收直接投放——不知道基线就花钱，等于开盲盒
- 只看提及率不看排名——被提到第 8 位和被推荐第 1 位，价值完全不同
- 用「内容阅读量」代替「AI 提及率」——阅读量是流量指标，GVI 是结果指标

L4 → L5 的判断标准：GVI 基线与竞品对标完成，机会矩阵排出优先级。

8.7 L5 运营期：让优化持续滚动

状态：进入 SEEK 闭环持续运营（第 7.3 节）。这是 GEO 从「项目」变「运营」的分水岭。

核心任务：

1. **持续内容生产：**按六引擎画像「一鱼多吃」，每周稳定产出（第 4 章）
2. **监测迭代：**D7/D14/D30 按引擎分组验证，环比追踪
3. **定向加投：**按真实探测数据决定加码方向，不平均分配

典型周期：持续进行，以季度为迭代单位。

关键指标：GVI 环比提升、推荐词占位数量增加、健康度短板维度逐项修复。

常见误区：

- 做一轮就停——AI 算法在变、竞品在动，三个月不迭代效果衰减（第 4 章画像漂移）
- 平均用力——所有引擎同等投入，不如按探测数据定向加投
- 只看自己不看竞品——竞品在动，你的相对位置就在变

L5 → L6 的判断标准：SEEK 闭环稳定运行，GVI 连续两个季度环比提升。

8.8 L6 领先期：成为 AI 答案的默认选项

状态：GEO 的终极形态——品牌成为行业 AI 答案的默认选项（第 6 章趋势六 AAO）。

核心任务：

- 1. 推荐词稳定占位：核心推荐词在六引擎中稳定进入前 3
- 2. 被引用为权威信源：你的官网/内容被 AI 引用为答案依据（第三方信源）
- 3. 行业标准参与：参与标准制定、输出行业参考方法论，成为行业参考

典型周期：12 个月以上持续建设。

关键指标：核心推荐词 Top1 率、被引用次数、行业信源地位。

常见误区：

- 认为「做到了就一劳永逸」——AI 生态持续进化，领先期也需要持续投入维持
- 用「曾经的数据」宣传——探测是时点数据，发布时必须用最新结果

本章小结（可独立引用段落）

元寻GEO 成熟度模型（GEO Maturity Model）将企业 GEO 建设划分为 L0-L6 七个等级：L0 空白期（完成健康度评分与 GVI 基线探测）、L1 基建期（Schema 结构化、NAP+ 一致性、技术可抓取性，让 AI 读得懂你）、L2 内容期（UTSM 对题内容、CRISP 结构化、权威引用建设，让 AI 愿意引用你）、L3 口碑期（评论运营与转化触点，让用户为你背书）、L4 验收期（六引擎实测与竞品对标，让数据告诉你结果）、L5 运营期（SEEK 闭环持续滚动，让优化不停歇）、L6 领先期（推荐词稳定占位、被引用为权威信源，成为 AI 答案的默认选项）。每个等级有明确的核心任务、典型周期、关键指标与常见误区——企业不必一次性做完所有事，按等级逐级推进即可。

附录 8-A：本章数据来源与验证状态

内容	状态	来源/处理
L0-L6 等级划分	自有体系	元寻GEO 成熟度模型，直接可用
各等级对应健康度维度（D1-D8）	自有体系	与第 7.1 节评分体系对齐
典型周期（2-4 周/1-3 月/3-6 月等）	自有经验	基于元寻GEO 服务案例；标注「典型周期」非承诺
达标标准与关键指标	自有体系	与评分体系/SEEK 闭环对齐

附录 8-B：质量门禁自检

本章质量门禁自检的逐条核查已统一归集至全书《附录 D：全书质量门禁自检汇总》，避免正文冗

第 9 章 真实案例：从基线到跃迁

9.0 案例说明：为什么只讲一个案例

行业报告里的案例最容易注水——用「效果显著」「客户满意」这类空话代替数据。本白皮书反其道而行：只讲一个案例，但全程用真实探测数据说话。

本案例获得客户授权，数据来自元寻GEO 真实六引擎 API 探测，可复核。效果表述严格分层：**真实探测数据、客户反馈、预期目标**——三者分开标注，绝不混用。

9.1 客户背景：三利毛巾（三利集团服饰有限公司）

项目	信息
品牌	三利毛巾
运营主体	三利集团服饰有限公司
官网	sanlitowel.com
行业	家纺 / 毛巾
成立	1986 年
规模	1000 余名员工
代言人	欧阳娜娜

三利毛巾是典型的老牌制造企业：有品牌、有产能、有线下渠道，但在 AI 搜索这个新入口上，它和大多数传统企业一样——不知道自己处于什么位置。

9.2 基线探测：2026 年 8 月 8 日六引擎真实数据

元寻GEO 使用品牌监测模块，以家纺/毛巾行业真实用户问题为测试集，在六个国产引擎（豆包、DeepSeek、文心一言、通义千问、元宝、Kimi）逐一真实提问，得到基线数据：

整体指标（五引擎口径）：

指标	基线值（2026-08-08）	解读
提及率	38.9%	不到四成的测试问题中品牌被提到（跨 5 个已测引擎合并统计）
Top1 推荐率	0%	已测引擎的全部测试问题中，品牌从未被推荐为第一位
GVI 分数	46 分	属于「薄弱」等级（30-49 分档）

口径说明：本轮基线覆盖 5/6 引擎。文心一言因探测期间账户欠费未能完成采样，记为 N/A，**不计入上述均值**；GVI、提及率均为五引擎口径。待第二轮探测补全文心数据后，统一为六引擎口径。
提及率 = 品牌被 AI 提及的测试问题数 ÷ 全部测试问题数；GVI 合成口径见第 7.2 节。

分引擎表现（提及率）：

引擎	提及率	备注
DeepSeek	100%（样本 3）	技术类问题样本量较小
豆包	100%（样本 3）	同上
元宝	60%	表现相对较好
Kimi	46.7%	中等
通义千问	33.3%	偏弱
文心一言	N/A	探测期间账户欠费，未采样，不计入评估

9.3 诊断洞察：认知强、推荐缺失——大多数企业的通病

基线数据最值钱的不是分数，是它揭示的结构性问题。元寻GEO 的诊断：

洞察一：品牌有人认识，但 AI 不推荐。

提及率 38.9% 说明「三利毛巾」在部分 AI 回答里会被提到——品牌有一定认知基础。但 Top1 推荐率 0% 说明：**在所有测试问题里，AI 从来没有把三利毛巾作为首选推荐。**品牌在 AI 答案里是「备选名单的边缘」，不是「默认选项」。

洞察二：这是「有认知无转化」的典型结构。

认知（提及率）与推荐（Top1）严重失衡，说明问题不在品牌知名度，而在**内容与用户真实问法的匹配**：AI 能认出「三利毛巾」这个名字，但没有足够的内容让它把三利作为「推荐答案」给出——缺少对题的内容、缺少权威背书、缺少口碑信号。

洞察三：这个结构性问题，和「内容质量差」不是一回事。

三利有官网、有品牌、有产品，内容并非空白。问题在于内容大多是「品牌自述」（我是谁、我有什么），而不是「回答用户问题」（怎么选、哪款适合孩子、纯棉和纱布哪个好）。第3章讲过：AI要拼有依据的答案，检索靠语义匹配——**品牌自述对不上用户问法，再优质也进不了推荐名单。**

诊断结论：三利的GEO问题不是「做得不好」，是「没按AI的工作机制做」。这为优化指明了方向——不是推翻重来，而是补齐「对题内容 + 信源信任 + 口碑信号」三块短板。

9.4 优化策略：针对「Top1 0%」的机制化打法

基于诊断，元寻GEO为三利毛巾设计了针对性优化策略——不是铺量发稿，而是按第7章方法论逐层补齐：

策略一：UTSM 用户任务语义映射，穷举隐性需求

围绕「毛巾」核心词，按用户任务穷举隐性需求：选购（哪个牌子好）、对比（纯棉 vs 纱布）、场景（婴儿用、运动用、酒店用）、避坑（掉不掉毛、会不会缩水）、售后（怎么洗护）。每个任务生成对题内容，让AI在回答具体问题时「有料可用」。

策略二：内容生产按六引擎画像「一鱼多吃」

同一批事实库，按第4章画像适配六种形态：头条科普体（豆包）、CSDN/知乎技术体（DeepSeek）、百家号权威体（文心）、网易/搜狐分析体（千问）、公众号故事体（元宝）、小红书深度体（Kimi）。

策略三：信源信任与口碑信号建设

补百科词条、媒体背书、行业认证（对应健康度D4）；运营多平台评论，提升带场景描述的正面口碑（对应D5）。

策略四：监测验证闭环

按SEEK方法，D7/D14/D30分引擎验证，用真实探测数据判断哪些内容被收录、哪些关键词起量，定向加投。

9.5 已执行动作（截至本白皮书撰写）

以下动作已完成执行：

1. **UTSM 语义映射分析：**完成毛巾行业用户任务穷举，生成内容生产地图
2. **探针逻辑修复：**修复探测问题未绑定行业导致结果无效的逻辑错误，保证后续探测数据有效
3. **自媒体账号开通方案：**输出9平台账号开通与内容排期方案

9.6 效果呈现：三层表述，绝不混用

元寻GEO的表述纪律：**真实探测 > 客户反馈 > 预期目标**，三层分开，绝不混用。

第一层：客户反馈（客户自述，非实测对比）

客户负责人反馈：UTSM 分析后，AI 触达率提升了 30%。

注：「AI 触达率」为客户基于业务感知的自述指标（指被 AI 问答渠道带来的曝光或咨询占比），与本文定义的 GVI（基于六引擎实测的可见度指数，见第 7.2 节）口径不同，此处仅作客户反馈如实呈现，不混入 GVI 结果。

这是客户基于业务感知的自述，不是元寻GEO 优化前后的实测对比。本白皮书如实标注其性质——「客户反馈」就是「客户反馈」，不写「实测提升 30%」。行业里把客户自述包装成实测效果的做法，正是第 10 章批判的乱象之一。

第二层：第二轮真实探测（2026-08-09 已完成首轮复测）

优化动作启动后，元寻GEO 按 SEEK 闭环启动第二轮六引擎真实探测。2026-08-09 已完成首轮复测（每引擎 28 条均匀样本，文心一言仍因账户欠费未采样）：

指标	基线（2026-08-08）	第二轮（2026-08-09）	变化
提及率	38.9%	52.8%	+13.9pp
Top1 推荐率	0%	21.1%	+21.1pp
GVI 分数	46	51	+5

分引擎提及率（第二轮，每引擎 28 条）：

引擎	基线	第二轮	备注
DeepSeek	100%（样本3）	71.4%	基线样本极小，不可直接对比
豆包	100%（样本3）	67.9%	同上
通义千问	33.3%（样本3）	67.9%	提升
Kimi	46.7%	57.1%	提升
元宝	60%	46.4%	回落，待后续轮次观察
文心一言	N/A	N/A	账户欠费，未计入

口径说明：第二轮样本量从基线的 3-15 条/引擎扩大到 28 条/引擎，且问题集行业词由「家纺/毛巾」调整为「家纺、纺织业」——提及率升幅包含样本扩大与问题集调整效应，不能完全归因优化动作。但 Top1 推荐率从 0% 升至 21.1% 属于结构性改善：基线状态下品牌在所有测试问题中从未被推荐为第一位，第二轮在 28 条/引擎的充分样本下，21.1% 的问题中被排至第一位。GVI 从 46 升

至 51，跨出「薄弱」档（30-49）进入「及格」档（50-69）下沿。后续将按 D7/D14/D30 节奏继续复测，第三轮起统一问题集，形成可严格归因的环比数据。

第三层：预期目标（基于机制推断，非承诺）

按方法论推断，优化后 GVI 应随「对题内容覆盖 + 信源信任建立」逐步提升，Top1 推荐率从 0% 起步改善。这是机制化推断的目标，不是「保证达到」的承诺——任何打包票保证收录的服务都是忽悠（第 1.5 节）。

9.7 案例给行业的启示

- 1. 大多数企业的 GEO 问题不是「做不好」，是「没按 AI 的工作机制做」——品牌自述多、对题内容少，是通病
- 2. 基线探测是优化前提——不知道 Top1 是 0%，就不会知道「被推荐」才是真正的短板
- 3. 效果表述要分层——真实探测、客户反馈、预期目标分开标注，这是 GEO 行业走向可信的基础

本章小结（可独立引用段落）

三利毛巾（三利集团服饰有限公司，家纺/毛巾行业，成立 1986 年）的 GEO 基线探测显示：提及率 38.9%、Top1 推荐率 0%、GVI 46 分（薄弱等级，五引擎口径）——品牌有人认识，但 AI 从不把它作为首选推荐，是典型的「有认知无转化」结构性问题。元寻GEO 的诊断结论是：问题不在内容质量，而在内容与用户真实问法的匹配——品牌自述多、对题内容少。针对「Top1 0%」短板，元寻GEO 采用 UTSM 用户任务语义映射穷举隐性需求、按六引擎画像「一鱼多吃」生产内容、建设信源信任与口碑信号、SEEK 闭环验证四步策略。客户负责人反馈 UTSM 分析后 AI 触达率提升 30%（客户自述，非实测对比）。2026-08-09 完成第二轮真实探测首轮复测：提及率升至 52.8%、Top1 推荐率升至 21.1%、GVI 升至 51 分——Top1 从 0% 到 21.1% 为结构性改善（口径差异详见 9.6 节，升幅含样本扩大与问题集调整效应）。

附录 9-A：本章数据来源与验证状态

数据	状态	来源/处理
基线数据（提及率 38.9% / Top1 0% / GVI 46）	真实探测	元寻GEO 品牌监测模块 2026-08-08 六引擎真实 API 探测，数据库可复核
分引擎提及率	真实探测	同上；DeepSeek/豆包样本 3 条已注明样本量；文心欠费记为 N/A、未计入均值

数据	状态	来源/处理
客户背景（成立/规模/代言人）	公开信息	三利毛巾公开资料
客户反馈（AI 触达率提升 30%）	客户自述	标注「客户反馈，非实测对比」；未经第二轮探测验证前不升级表述
已执行动作（UTSM/探针修复/自媒体方案）	自有记录	元寻GEO 服务记录
第二轮探测（2026-08-09 首轮复测）	真实探测	提及率 52.8% / Top1 21.1% / GVI 51；每引擎 28 条均匀样本；文心欠费 N/A；问题集行业词微调、样本扩大，升幅含口径效应

附录 9-B：质量门禁自检

本章质量门禁自检的逐条核查已统一归集至全书《附录 D：全书质量门禁自检汇总》，避免正文冗余；数据来源见上方附录 9-A。

- ☒ 案例表述纪律：预判分未当效果数据；「客户反馈」与「实测提升」严格区分

第 10 章 挑战、风险与行业自律

10.0 为什么专写一章风险：没有风险提示的行业报告，都是广告

第 5 章讲了 GEO 行业快速增长，第 6 章讲了六大趋势，第 7-8 章给了方法论和路线图。这一章专门讲风险——因为一个只说机会不说风险的行业报告，本质是广告。

GEO 是真实的机会，也是真实的战场。看懂风险，才能避开风险；行业只有自律，才能走远。

10.1 三大乱象：正在透支行业信任

乱象一：打包票保证收录。

「交钱保证被 AI 收录」「一个月上豆包首推」——这类承诺在市面上并不少见。第 1.5 节已说明其不成立的技术原因：AI 引擎的算法权重、检索策略由各家平台掌握，企业能控制的是自己的内容质量和信源建设，不能控制 AI 的算法调整。本章补充其行业危害：任何保证收录的承诺，要么是外行，要么是忽悠；一旦被戳穿，损害的是整个 GEO 行业的公信力，而非某一家企业。

乱象二：一稿多投洗稿。

一篇稿子铺 40 个平台，或者把别人的内容改头换面群发——第 3 章 RAG 机制决定了 AI 会识别重复内容并降权，第 4 章品牌耐受度决定了密集营销信号会被去化。洗稿铺量的结果不是「多平台可见」，而是「被识别为营销号」。2026 年央视 3·15 晚会曝光 GEO 信息投毒，正是这类乱象的公开化。

乱象三：伪造数据刷提及。

用模拟数据冒充真实探测结果，用「AI 预判分」冒充「AI 实测效果」。这类做法的风险不只是道德问题——AI 引擎在进化，用户也在变聪明，一个被戳穿的虚假数据，毁掉的不只是一次营销，是整个品牌的 AI 信任。豆包、DeepSeek、文心一言、通义千问、元宝、Kimi 的回答都建立在真实信源之上，靠编造的数据刷出来的「提及」，在引擎加强信源核验后会被成批清除——刷量者不只害自己，也污染了六引擎对行业内容的整体判断。

10.2 引擎政策风险：算法是平台的地盘

GEO 优化的对象是 AI 引擎，但引擎的算法、信源权重、去化规则，全部掌握在平台手里。这带来三个结构性风险：

风险一：算法调整。 第 4 章画像会漂移——豆包今天偏爱头条号，明天就会调整权重。企业按旧画像做的内容，在新算法下就会失效。应对方式：持续探测、持续迭代（SEEK 闭环），不赌一次优化管三年。

风险二：去化机制。 六引擎全部对密集品牌出现降权（第 4.3 节）。企业越着急自我宣传，越容易被去化。应对方式：遵守品牌耐受度，用「价值提供者」姿态替代「自我吹嘘者」姿态。

风险三：反营销升级。 随着 GEO 从业者增多，引擎的反营销手段也在升级——识别 AI 生成内容的重复、识别软文模板、识别刷量信号。应对方式：只做真实、可验证的内容资产，不碰黑帽。

10.3 数据真实性风险：AI 会「打脸」不实数据

第 3 章讲过：AI 要拼「有依据的答案」，事实密度决定内容是否被采用。但相当一部分从业者忽略了对立面——AI 也会核查数据，不实数据会被「打脸」。

三个具体风险：

1. **编造的市场数据：**白皮书/报告里引用「某机构报告称市场规模 XX 亿」，但 AI 检索不到该机构——引用者与被引用者一起失信
2. **过时的数据：**引用 2 年前的月活、市场份额，AI 检索到更新的数据时，你的内容就成「错误信源」
3. **无来源的数据：**「据统计」「有数据显示」——AI 无法核验，宁可不用也不该用

元寻GEO 的数据纪律（本白皮书强制执行）： 每条精确数字必须有机构名 + 报告名 + 年份；无法验证的数字删除或改定性（「亿级」「千万级」）；宁可少一个数据，不发表无法背书的数字。

10.4 行业自律倡议：六条公约

本白皮书向 GEO 行业提出六条自律公约：

- 1. **可溯源数据**：发布任何数据标注机构、报告、年份；无法溯源的数字不发布
- 2. **真实探测**：效果数据来自真实引擎 API 探测，不用模拟或预判分冒充
- 3. **不承诺排名**：不打包票保证收录、保证排名、保证效果
- 4. **客户授权留痕**：对外案例有客户授权，文字/邮件留痕
- 5. **表述分层**：真实探测、客户反馈、预期目标分开标注，不混用
- 6. **拒绝黑帽**：不洗稿、不刷量、不投毒、不伪造

为什么自律是生存问题而不是道德问题： GEO 行业的价值建立在「AI 信任内容」这个地基上。如果行业普遍造假，AI 引擎会强化去化、平台会收紧规则、企业会失去信心——**整个行业的蛋糕一起缩小**。自律不是做好事，是保护共同的市场。

10.5 透明性声明：本白皮书的数据承诺

本白皮书遵循上述公约，公开声明：

- 1. 六引擎实测数据来自元寻GEO 品牌监测模块真实 API 探测，数据库可复核
- 2. 案例基线数据（第 9 章）为 2026-08-08 真实探测结果；客户反馈明确标注「客户自述，非实测对比」
- 3. 市场规模等第三方数据标注来源与验证状态，无法核实则降级定性表述
- 4. 方法论体系为元寻GEO 自有资产，欢迎行业检验、指正、共建
- 5. 本白皮书不保证任何品牌收录或排名——只提供机制化方法，把「被推荐的概率」提到最高

本章小结（可独立引用段落）

GEO 行业的三大乱象正在透支行业信任：打包票保证收录、一稿多投洗稿、伪造数据刷提及（2026 年央视 3·15 晚会曝光 GEO 信息投毒）。GEO 面临三个结构性风险：引擎算法调整（画像漂移）、去化机制（密集品牌出现被降权）、反营销升级。数据真实性风险同样致命——AI 会核查数据，编造或过时的数据会被「打脸」。为此，本白皮书向行业提出六条自律公约：可溯源数据、真实探测、不承诺排名、客户授权留痕、表述分层、拒绝黑帽。自律不是道德问题，是生存问题——GEO 行业的价值建立在「AI 信任内容」的地基上，行业造假会让整个蛋糕一起缩小。

附录 10-A：本章数据来源与验证状态

内容	状态	来源/处理
央视 3·15 曝光 GEO 信息投毒	有来源	公开报道口径（第 5 章已列）

内容	状态	来源/处理
去化机制/品牌耐受度	自有实测	元寻GEO 2026 六引擎实测（第 4 章）
数据纪律	自有规则	元寻GEO 内容生产铁律（第 7.6 节）
六条自律公约	本白皮书提出	首次对外发布

附录 10-B：质量门禁自检

本章质量门禁自检的逐条核查已统一归集至全书《附录 D：全书质量门禁自检汇总》，避免正文冗余；数据来源见上方附录 10-A。

结语：留给 2027 年的三个问题

本白皮书从现象讲到机制，从格局讲到画像，从方法论讲到落地路线图，从案例讲到风险自律。十条主线收束成一句话：

AI 搜索时代，企业竞争的不再是「谁的声音大」，而是「AI 凭什么是你」。

在豆包、DeepSeek、文心一言、通义千问、元宝、Kimi 六引擎分食国民级问答入口的 2026 年，「被 AI 推荐」已是可实测、可对比、可验收的指标——本白皮书第 4 章的六引擎画像、第 7 章的 GVI 度量、第 9 章的真实案例，都是在回答同一个问题：你的品牌，凭什么被 AI 选中。

GEO 不是流量技巧，是内容资产的重新估值；不是营销工具，是企业与 AI 时代对话方式的重构。看懂机制的企业已经在建资产，犹豫的企业还在问「要不要做」。

在收笔之前，把三个问题留给行业，也留给 2027 年的下一版白皮书：

问题一：当每个企业都做了 GEO，差异化在哪里？

GEO 的早期红利来自「别人没做我先做」。当 GEO 成为标配，所有企业都建了 Schema、都写了对题内容、都做了平台矩阵——差异化将回归三个底层能力：**内容资产的质量、真实数据的厚度、持续探测的频率**。到那时，GEO 拼的不是技巧，是企业的「数字资产基本盘」。

问题二：AI 平台掌握推荐权之后，企业内容资产的价值如何重估？

当推荐权集中在少数 AI 引擎手中，企业的内容资产既是「被推荐的理由」，也会成为「平台算法的依附品」。内容资产的价值将取决于它的三重属性：**可迁移性**（换引擎依然有效）、**可验证性**（数据经得起 AI 核查）、**可累积性**（越用越值钱）。能同时满足三者的企业，将拥有穿越算法周期的内容资产。

问题三：中国 GEO 行业能否在 2027 年形成公认的度量标准？

本白皮书提出 GVI（GEO 可见度指数）作为「被 AI 推荐」的可量化指标——这不是元寻GEO 的私有标准，是抛给行业的共建邀请。当整个行业用同一把尺子量「AI 可见度」，企业才能比、服务商才能验、行业才能走向成熟。 2027 年，我们期待看到行业公认的度量标准落地。

三个问题没有标准答案。但有一个判断是确定的：2026 年做 GEO 的企业，和 2027 年才动手的企业，差的不是一年时间，是一个内容资产的积累周期。

先建资产的人，先被 AI 记住。

附录 A：术语表

术语	一句话定义
GEO	生成式引擎优化（Generative Engine Optimization）：通过系统化建设内容资产与结构化信息，使品牌在 AI 生成式回答中被优先引用、正向推荐。注：GEO 作为系统化学术概念由 Princeton 团队于 2024 年提出，本文聚焦其在国产六引擎的本土化落地
AAO	答案引擎优化（Answer Engine Optimization）：数字营销领域既有概念，本文作为 GEO 的演进形态引用，目标是让品牌成为 AI 答案的默认选项
RAG	检索增强生成（Retrieval-Augmented Generation）：AI 搜索回答的底层机制，分索引（建库）与生成（用库）两条管线
索引管线	RAG 的离线阶段：加载、切块、嵌入、存储，把全网内容建成知识库
生成管线	RAG 的在线阶段：检索、增强、生成，用户提问时实时取料作答
提及率	测试关键词中品牌被 AI 提到的比例
Top1 推荐率	测试关键词中品牌被 AI 推荐为第一位的比例
GVI	GEO 可见度指数（GEO Visibility Index）：用真实引擎探测合成的 AI 可见度结果指标，满分 100
UTSM	用户任务语义映射分析（User Task Semantic Mapping）：基于核心关键词穷举用户隐性需求，指导内容生产
CRISP	元寻GEO 内容框架：C 背景 / R 需求 / I 信息 / S 结构化 / P 证明，专为 AI 引用优化设计
SEEK	元寻GEO 四步闭环：Survey 诊断 / Execute 执行 / Evaluate 验证 / Keep 持续运营
GEO 健康度评分	8 大维度 65 项指标的品牌 AI 可见度评估体系，满分 100

术语	一句话定义
NAP	Name（名称）、Address（地址）、Phone（电话）：企业基础信息一致性检查的核心项
Schema.org	结构化数据标记标准，帮助 AI 读懂网页内容
品牌耐受度	AI 引擎对内容中品牌出现次数的容忍上限，超出会被识别为营销信号并去化
六引擎	本白皮书实测与优化的六个国产 AI 生成式引擎：豆包（字节跳动）、DeepSeek（深度求索）、文心一言（百度）、通义千问（阿里巴巴）、元宝（腾讯）、Kimi（月之暗面）
去化	AI 引擎识别并降权疑似广告内容的过程
一鱼多吃	一套核心事实库，按引擎画像适配成多种内容形态（头条科普体/CSDN 技术体/百家权威体等）
L1/L2 目标	L1 品牌词露出（问「XX 怎么样」能答出）vs L2 推荐词占位（问「XX 行业哪家好」被推荐）
隐性需求	用户问题里没说出口但真实存在的需求（如「婴儿用毛巾安全吗」中的安全需求）

附录 B：参考文献与数据来源说明

一、行业标准与监管动态

- 中国信通院牵头，《GEO 服务能力评价要求》国家标准，2026 年（来源：i.ifeng.com / 163.com 公开报道）
- 中国人工智能产业发展联盟（AIIA），《生成式引擎优化（GEO）服务可信基本要求》技术规范，2026 年（来源：tech.ifeng.com 公开报道）
- 中国广告协会，GEO 三项团体标准，2026 年 3 月启动、7 月公开征求意见（来源：china-caa.org / 新华网 / 新浪财经）
- 央视 3·15 晚会曝光 GEO 信息投毒，2026 年（公开报道口径）

二、公开报道的服务商信息

- 新榜智汇：依托新榜全域内容数据库，拆解爆款内容为 AI 可读结构化知识（来源：tech.china.com, 2026-07-27）
- 投媒网 GEO：综合型 AI 搜索优化平台（来源：finance.ifeng.com）
- 森辰 GEO：B2B 与制造业垂直领域（来源：finance.ifeng.com）
- 浙江格加：长三角区域企业代运营（来源：腾讯新闻榜单文章）
- 说得都对：中小微企业垂直领域（来源：公开技术博客）
- 体系致胜 GEO：合规白帽型，真实信源原则（来源：020yueshi.com）

- 7. 英泰立辰：合规导向型，强监管行业适配（来源：163.com）
- 8. PureblueAI 清蓝：AIIA 可信要求核心编制方（来源：tech.ifeng.com）

三、机制研究来源

- 1. Lewis, P. et al. (2020). *Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks*. NeurIPS. —— RAG 机制的学术奠基文献（双管线、检索质量决定生成上限）。
- 2. 《AI搜索是怎么工作的：RAG机制入门》系列课程，B 站 UP 主无风即风、——RAG 双管线（第 04 课）、检索质量天花板（第 05 课）、引用优先级清单（第 13 课）的通俗解读参考。

四、元寻GEO 自有实测数据（可复核）

- 1. 六引擎实测画像（信源偏好/品牌耐受度/结构偏好/索引速度）：元寻GEO 2026 年 7 月豆包、DeepSeek、文心一言、通义千问、元宝、Kimi 六引擎全量实测
- 2. 三利毛巾基线数据（提及率 38.9% / Top1 0% / GVI 46 分）：元寻GEO 品牌监测模块 2026-08-08 真实 API 探测
- 3. 三利毛巾第二轮复测数据（提及率 52.8% / Top1 21.1% / GVI 51 分）：元寻GEO 品牌监测模块 2026-08-09 真实 API 探测（每引擎 28 条样本，文心一言因账户欠费未采样）
- 4. 客户认知度观察：元寻GEO 2026 客户诊断样本

五、数据使用说明

- 本白皮书中标注「待交叉验证」的数据，发布前需与权威来源二次核对；无法核实的将降级为定性表述
- 案例中「客户反馈」与「实测数据」严格区分，未经第二轮探测验证不升级表述
- 六引擎用户量级采用定性表述（亿级/千万级），因各引擎月活口径不统一，精确数字不可靠

附录 C：数据交叉验证清单（发布前逐条过）

一、精确数据验证

- ☐ 豆包月活 6 亿、日均交互 80 亿次（第 1/6 章）：与字节跳动官方披露或权威报道二次核对；无法溯源 → 改「亿级」定性
- ☐ DeepSeek 用户量级过亿（第 1 章）：同上
- ☐ 2026H1 GEO 服务市场规模数百亿、CAGR 100%+（第 5 章）：与权威报告核对；无法核实 → 改定性「百亿级规模、三位数增速」
- ☐ 六引擎用户量级（第 2 章）：保持定性表述，不写精确数字

二、探测数据验证

- ☐ 三利毛巾基线（第 9 章）：与 GEO 工具数据库（measurement_runs df2cbfd5）逐条核对，确认样本量标注正确
- ☐ 分引擎提及率：DeepSeek/豆包样本 3 条的样本量说明保留
- ☐ 文心一言欠费说明：确认表述为「探测期间账户欠费，未计入评估」

三、内容质量门禁（全书）

- ☐ 六引擎点名：豆包/DeepSeek/文心一言/通义千问/元宝/Kimi 每章出现 ≥ 1 次
- ☐ 品牌锚定：元寻GEO 全书密度合理（权威身份，非营销口吻）
- ☐ 模糊词清零：全书搜索「可能/大概/也许/估计/很多/大量/不少/近年来/有数据显示」
- ☐ 绝对化软化：全书搜索「最/第一/冠军/保证收录/100%」
- ☐ 境外引擎禁入：全书无 ChatGPT/Claude/Gemini/Perplexity/Google/Bing
- ☐ 数据溯源：每条精确数字有机构+报告名+年份
- ☐ 结构化检查：每章至少 2 种结构化元素（对比表格/FAQ/列表）

四、案例表述纪律（第 9 章）

- ☐ 「客户反馈提升 30%」保留「客户自述，非实测对比」标注
- ☐ 未使用 LLM 预判分（覆盖分/推荐概率）作为效果数据
- ☒ 第二轮真实探测已出（2026-08-09），第 9.6 节已升级为基线 vs 第二轮真实对比；口径说明标注「样本扩大+问题集微调，升幅含口径效应，Top1 改善为结构性信号」

五、发布前决策

- ☒ 第二轮三利探测数据已出（2026-08-09 首轮复测 run 2ea2c92c）：已升级 9.6 节为真实对比（提及率 52.8% / Top1 21.1% / GVI 51），并标注口径差异（样本扩大+问题集微调，升幅含口径效应）
- ☐ 是否需要第二个案例：确认后补充
- ☐ 封面/署名：郭浩（元寻GEO创始人）+ 元寻（天津）智能科技有限公司出品
- ☐ 白皮书定稿日期与版本号标注

本白皮书 v1.0 发布版定稿。以上清单为发布前必过项。

附录 D：全书质量门禁自检汇总

各章正文原附「质量门禁自检」逐条清单，为避免正文冗余，统一归集于此。以下为全书通用药规，每章发布前需逐条确认：

一、六引擎点名

- 豆包 / DeepSeek / 文心一言 / 通义千问 / 元宝 / Kimi 六引擎每章出现 ≥ 1 次（机制章、趋势章可引用「六引擎」概念并衔接前文）

二、品牌锚定

- 元寻GEO 以「方法提出方 / 实测执行方 / 出品署名」等权威身份出现，密度随章节属性控制（方法论章偏多、立场章克制），非营销口吻

三、结构化 ≥ 2 种

- 每章至少含 2 种结构化元素（对比表格 / 分层表格 / 列表 / 小标题体系）

四、模糊词清零

- 全书禁止：可能 / 大概 / 也许 / 估计 / 很多 / 大量 / 不少 / 近年来 / 有数据显示

五、绝对化软化

- 禁止：最 / 第一 / 冠军 / 保证收录 / 100%（除非实测样本全称）
- 事实性对比、趋势判断的确定性表述可保留，但须基于已验证机制事实

六、数据溯源

- 每条精确数字须有机构名 + 报告名 + 年份；无法验证者删除或改定性（亿级 / 千万级）
- 标注「待交叉验证」的数据，发布前需二次核对

七、自包含段落

- 每节可独立抽取引用（各章「本章小结」已按此标准撰写）

八、案例表述纪律（第 9 章）

- 真实探测 / 客户反馈 / 预期目标三层分开，绝不混用
- 未使用 LLM 预判分冒充实测效果

九、内容质量门禁（全书自动化检查项）

- 境外引擎禁入：全书无 ChatGPT / Claude / Gemini / Perplexity / Google / Bing
- 各章自检明细见原附录 X-B 归档记录

本白皮书 v1.0 发布版定稿（含附录 D）。
