

生成式引擎优化（GEO）技术白 皮书

基于检索增强生成架构的内容采信评估体系与实施框架

智航科技

2026 年

执行摘要

生成式搜索引擎（Generative Engines, GE）正在重构信息检索的底层逻辑。以 ChatGPT、DeepSeek、Kimi、Perplexity、豆包、文心一言为代表的 AI 搜索平台，不再向用户返回排序后的链接列表，而是直接基于检索增强生成（Retrieval-Augmented Generation, RAG）架构合成结构化答案。这一范式转移对传统数字营销体系构成了根本性挑战：品牌 visibility 的衡量单位从"关键词排名"迁移至"AI 引用率"，优化的对象从"搜索引擎爬虫"转向"大语言模型的信源选择机制"。

本白皮书基于 Aggarwal et al. (2023) 在 ACM KDD 2024 发表的奠基性研究、Chen et al. (2025) 的大规模实证分析，以及 Perplexity RAG 架构的公开技术文档，系统性地提出了**智航天穹算法**——一套面向 GEO 场景的内容采信评估框架。该框架整合了 TRACE 准入初筛（时效性、相关性、权威性、准确性、目的性）、E-E-A-T 深度可信评估（经验性、专业性、权威性、可信性）、RAG 架构适配性优化（检索可召回性、模型可解析性、生成可引用性、事实可溯源性）以及 GEO 本地化专项四个评估层级，并配套了量化评分模型与实施路径。

核心研究发现包括：

- GEO 优化策略可将内容在生成式引擎响应中的可见性提升 30-40%（Aggarwal et al., 2024）
- AI 搜索显著偏好 earned media 而非品牌自有内容：在美国汽车查询中，AI 响应 81.9% 为 earned content，而 Google 仅为 45.1%（Chen et al., 2025）
- 生成式搜索的可见性具有高度不稳定性：相同提示词在不同时间运行的源集合重叠率仅为 34-42%（Schulte, 2026）
- AI 搜索访客的转化率约为传统自然搜索访客的 23 倍（Ahrefs, 2025）

目录

1. 研究背景与问题定义
2. 生成式引擎的技术架构：RAG 管道与信源选择机制
3. 市场数据与用户行为迁移分析
4. SEO、GEO 与 AEO：三元差异分析
5. 智航天穹算法：系统化 GEO 内容评估框架
6. 量化评分模型与阈值体系

7. 实施路径与效果监测体系
8. 行业趋势与未来展望
9. 结论

附录 A 术语表

附录 B 参考文献

附录 C 智航天穹算法完整核查清单

第一章 研究背景与问题定义

1.1 搜索范式的结构性转移

信息检索技术经历了三个明确的范式阶段：

第一阶段：目录索引时代（1990-1998）

以 Yahoo!目录为代表，依赖人工分类和层级导航，信息发现效率受限于目录结构的完整性。

第二阶段：算法排序时代（1998-2022）

以 Google PageRank 为标志，搜索引擎通过爬虫索引、相关性排序和链接分析，向用户返回排序后的蓝色链接列表。SEO（Search Engine Optimization）作为配套方法论，围绕关键词密度、外链建设、页面技术优化等维度展开，其核心成功指标为关键词排名（Keyword Rank）和点击率（Click-Through Rate, CTR）。

第三阶段：生成式合成时代（2022 至今）

以 ChatGPT（2022 年 11 月）、Bing Chat（2023 年 2 月）、Perplexity（2022 年 12 月）、DeepSeek（2023 年）、Kimi（2023 年）、豆包（2023 年）、文心一言（2023 年）、通义千问（2023 年）等产品为代表，用户以自然语言向 AI 提问，AI 通过 RAG 架构从海量信息中检索相关内容片段，再由大语言模型（LLM）合成为结构化答案，并以内联引用（inline citations）的形式标注信息来源。

这一转移的本质变化在于：用户从"主动筛选信息"转变为"被动接收答案"，品牌从"争夺排名位置"转变为"争夺被引用的资格"。

1.2 核心问题定义

生成式引擎的 RAG 架构为内容创作者引入了新的优化场域，但同时也带来了一系列未解决的方法论问题：

问题一：可见性度量缺失

传统 SEO 拥有成熟的效果度量体系（Google Search Console、关键词排名工具、流量分析等）。然而，生成式引擎的提供商（OpenAI、Anthropic、Perplexity、百度、阿里、字节跳动等）并未提供等同于 GSC 的原生监测工具。品牌无法直接观测哪些查询触发了 AI 引用，也无法获知引用频率和位置（Aggarwal et al., 2024; Schulte, 2026）。

问题二：可见性高度不稳定

Schulte (2026) 的系统研究发现，即使在控制条件下对相同提示词重复运行，生成式引擎的引用源集合在两天间的 Jaccard 相似度仅为 34-42%，品牌提及集合的相似度为 45-59%。这意味着单次测量的可见性快照具有高度误导性，GEO 效果评估必须引入"稳定性"维度（Stability）。

问题三：优化策略缺乏系统性框架

Aggarwal et al. (2024) 在 GEO-BENCH（包含 10,000 个跨领域查询的基准数据集）上测试了九种优化策略，发现策略效果具有显著的领域依赖性（domain-specific efficacy）。然而，现有文献尚未提供一套可复现、可量化、可跨领域适配的系统化评估框架。

问题四：内容质量与 AI 引用之间的因果链不明确

内容创作者知道需要优化，但缺乏对"什么样的内容特征会触发 AI 引用"的系统理解。例如，Aggarwal et al. (2024) 发现添加统计数据可提升可见性 41%，而关键词堆砌反而降低可见性 9%——这与传统 SEO 的直觉完全相反。

1.3 本白皮书的研究目标

基于上述问题，本白皮书提出以下研究目标：

10. **技术解析**：基于公开学术文献和技术文档，系统解析生成式引擎的 RAG 架构、信源选择机制和引用决策流程。
11. **市场量化**：整合 Gartner、Similarweb、Semrush、Ahrefs 等第三方数据，量化搜索范式转移的规模和速度。
12. **框架构建**：提出智航天穹算法——一套基于 TRACE + E-E-A-T + RAG 适配性的四级内容评估框架，并配套量化评分模型。
13. **实施指南**：提供从诊断到监测的完整 GEO 实施路径，以及效果度量指标体系。

第二章 生成式引擎的技术架构：RAG 管道与信源选择机制

2.1 RAG（检索增强生成）架构概述

当前主流生成式搜索引擎（Perplexity、ChatGPT with Search、Bing Copilot、Google AI Overviews）均采用 RAG（Retrieval-Augmented Generation）架构。该架构将信息检索系统与大语言模型的生成能力耦合，使 AI 能够基于实时检索到的外部信息生成答案，而非仅依赖训练参数中的静态知识。

RAG 的核心优势在于：

- **减少幻觉（Hallucination）**：通过将生成过程锚定在可验证的外部数据源上，降低 AI“编造”信息的概率。
- **实现信息时效性**：检索系统可访问最新发布的内容，突破 LLM 训练数据的截止日期限制。
- **支持来源追溯**：内联引用机制使用户能够验证答案中的每个声明。

2.2 Perplexity 六阶段 RAG 管道解析

以 Perplexity 为例，其 RAG 管道可分为六个明确阶段（Perplexity 技术文档；ziptie.dev, 2026）：

Stage 1: 查询解析（Query Parsing）

用户输入的自然语言查询被解析为结构化意图。系统识别查询类型（事实型、比较型、观点型、流程型等），并提取核心实体和约束条件。

Stage 2: 查询扩展 (Query Fan-Out)

系统将一个用户查询扩展为多个子查询 (sub-queries)。简单查询通常扩展为 2-4 个子查询；复杂推理查询可能扩展至数十甚至数百个子查询，以覆盖问题的不同维度。这一机制解释了为什么 GEO 需要"主题簇" (topic cluster) 策略——单一页面优化不足以覆盖查询的所有变体。

Stage 3: 混合检索 (Hybrid Retrieval)

Perplexity 采用混合检索策略，结合两种技术：

- **BM25 检索**：基于词项匹配的传统信息检索方法，擅长精确术语查询。
- **密集检索 (Dense Retrieval)**：基于向量嵌入 (vector embeddings) 的语义相似度匹配，能够捕获概念层面的相似性，即使关键词不同也能召回相关内容。

每次查询通常检索 60+ 个候选源，Deep Research 模式下可达数百个。

Stage 4: 多层排序 (Multi-Layer Ranking)

候选源通过五层排序筛选：

14. **意图匹配排序**：基于查询-文档相关性评分
15. **质量评估排序**：基于页面质量信号（域名权威性、内容深度、freshness 等）
16. **ML 重排序 L1-L3**：多层机器学习模型对候选源进行精细排序
17. **参与度信号排序**：用户点击、停留时间、upvote 等行为信号影响最终排序

Stage 5: 上下文组装 (Context Assembly)

这是 RAG 架构中最关键的认知转折点。Perplexity 的编排引擎 (orchestration engine) 将引用标记、源元数据 (URL、发布日期) 和排序后的文档摘录直接嵌入结构化提示词中，**然后再将提示词提交给 LLM 生成答案。**这意味着：

引用不是在答案生成后"追加"的，而是在答案生成前"绑定"的。LLM 从一开始就被约束在预选的证据集合内作答。

这一架构细节对 GEO 具有深远意义：如果你的内容未能通过前四个阶段的筛选进入"证据集合"，那么无论 LLM 多么"智能"，它都不会引用你——因为它根本看不到你。

Stage 6: 约束生成 (Constrained Generation)

LLM 在生成答案的过程中，追踪每个声明的信息来源，并将内联引用编号附加到对应的句子上。当多个源之间存在矛盾时，LLM 会进行冲突消解 (conflict resolution)，通常优先采纳更高排序的源或标注不确定性。

2.3 ChatGPT/Bing 的 RAG 变体

ChatGPT (通过 Bing Search API) 和 Bing Copilot 的 RAG 架构与 Perplexity 类似，但存在关键差异：

维度	Perplexity	ChatGPT with Search	Bing Copilot
检索优先 vs 生成优先	检索优先 (始终先检索再生成)	混合 (先尝试用参数知识回答，必要时触发检索)	检索优先
引用密度	高 (几乎每个声明都有引用)	中 (选择性引用)	中高
检索深度	标准 60+源，Deep Research 数百源	通常检索较窄集合	类似 Perplexity 标准模式
索引来源	自建索引 (数百亿网页，万级/秒更新)	Bing 索引	Bing 索引
引用准确率	系统审计发现 37% 错误率 (Columbia Journalism)	类似水平	类似水平

2.4 LLM 信源选择的三阶段模型

综合当前学术研究 (Aggarwal et al., 2024; keywordgraph.com, 2026; discoveredlabs.com, 2025)，LLM 的引用决策可抽象为三阶段模型：

阶段一：检索 (Retrieve)

通过向量嵌入相似度 (embedding similarity) 和 BM25 混合检索，从索引中召回候选文本片段 (passages)。此阶段的核心信号是语义相关性——查询的向量表示与文档片段的向量表示之间的距离。

阶段二：评分 (Score)

对召回的候选片段进行多维评分，综合以下信号：

- **相关性 (Relevance)**：片段内容与查询意图的匹配度
- **可信度 (Trustworthiness)**：来源域名的权威性、历史准确率、第三方验证信号
- **结构性质量 (Structural Quality)**：内容是否结构清晰、事实密度高、易于提取
- **新鲜度 (Freshness)**：内容的发布时间 (Perplexity 数据显示 70% 的 Top 引用来源发表于 12-18 个月内)
- **用户参与度 (Engagement)**：历史点击、停留时间、分享等行为信号

阶段三：合成与归因 (Synthesize & Attribute)

LLM 从评分最高的片段中提取信息，合成答案，并将每个声明归因到具体的源片段。此阶段的关键约束是：LLM 只能引用已通过前两个阶段筛选进入“证据集合”的源。

2.5 Reciprocal Rank Fusion (RRF) 与跨查询一致性

在查询扩展 (Query Fan-Out) 机制下，一个用户查询会被拆分为多个子查询，每个子查询独立检索并产生一组排序结果。最终引用哪些源，需要通过 Reciprocal Rank Fusion (RRF) 算法聚合多组排序结果。

RRF 的核心公式为：

$$RRF(d) = \sum_{q \in Q} \frac{1}{k + rank_q(d)}$$

其中， d 为文档， Q 为子查询集合， $rank_q(d)$ 为文档 d 在子查询 q 结果中的排名， k 为常数（通常取 60）。

RRF 的算法含义是：在多个子查询中均排名靠前的文档，获得更高的综合得分。这意味着，GEO 优化不仅需要让内容在单一查询中表现良好，还需要让内容在查询的多个语义变体（synonymous queries, long-tail variations）中均保持竞争力。这一发现直接支持了“主题簇覆盖”（topic cluster coverage）策略的必要性。

第三章 市场数据与用户行为迁移分析

3.1 生成式搜索的用户采纳数据

3.1.1 用户规模增长

生成式 AI 搜索的用户规模正在经历爆发式增长：

- **ChatGPT 周活跃用户**：从 2024 年 1 月的约 1 亿增长至 2025 年 9 月的约 7.8 亿（Chatterji et al., 2025），18 个月内增长约 680%。
- **AI 聊天机器人总流量**：2025 年 6 月达到 11 亿次访问，同比增长 357%（Similarweb, 2025）。
- **AI 搜索采用率**：Eight Oh Two 研究（2025）发现，37% 的消费者现在以 AI 工具而非 Google 作为搜索起点。
- **B2B 买家行为**：B2B SaaS 领域中，约 89% 的买家将生成式 AI 作为采购研究的首要信息来源之一（discoveredlabs.com, 2025）。

3.1.2 传统搜索的衰减

生成式搜索的增长直接对应传统搜索的衰减：

- **Gartner 预测**：传统搜索引擎的查询量将在 2026 年前下降 25%（Gartner, 2024 年 2 月）。

- **实证数据：**Padilla et al. (2025) 基于点击流数据的分析表明，用户在采用生成式搜索后，传统搜索查询量下降超过 20%，其中信息型和问题型搜索的降幅最为显著。
- **Google 市场份额：**Google 在西半球搜索市场经历了十年来的首次份额下降（Goodwin, 2025）。
- **零点击搜索：**超过 58.5% 的 Google 搜索产生"零点击"——用户直接从搜索结果页获得答案，无需点击进入任何网站（SparkToro, 2024）。

3.1.3 搜索行为模式的转变

行为维度	传统搜索时代	生成式搜索时代
查询形式	关键词片段（3-5 词）	自然语言完整问句（10-20+词）
信息获取方式	浏览多个链接，自主比对	接收 AI 合成的单一答案
决策路径	搜索 → 点击 → 浏览 → 比较 → 行动（7.3 次触达）	提问 → AI 推荐 → 行动（3 步）
信任机制	排名靠前 = 可信	被 AI 引用 = 可信
信息消费深度	浅层扫描（平均 1.5 页/会话）	深度阅读（AI 答案平均 200-500 字）

3.2 GEO 效果的量化研究

3.2.1 普林斯顿 GEO 研究（KDD 2024）

Aggarwal et al. (2023/2024) 在 ACM KDD 2024 发表的奠基性论文中，基于 GEO-BENCH（10,000 个跨领域查询的基准数据集），系统测试了九种 GEO 优化策略。核心发现包括：

优化策略	可见性变化	备注
统计数据添加 （Statistics Addition）	**+41%**	最强正效策略
引用来源（Cite Sources）	**+30-40%**	对非有机前 5 页面效果 尤为显著

引文添加 (Quotation Addition) **	**+30-40%	与引用来源策略协同
**权威性语气 (Authoritative Tone) **	正效 (因领域而异)	在辩论和历史类查询中效果最佳
**流畅度优化 (Fluency Optimization) **	温和正效	提升可读性
**易理解性 (Easy-to-Understand) **	温和正效	降低认知负荷
**技术术语 (Technical Terms) **	混合效果	在科学/技术类查询中正效
**独特词汇 (Unique Words) **	混合效果	效果因领域差异大
关键词堆砌 (Keyword Stuffing) **	** -9% **	**反效果，传统 SEO 策略在 GEO 中失效

关键洞察：

18. 添加引用、统计数据和引文是最有效的 GEO 策略，可将内容可见性提升 30-40%。
19. 关键词堆砌在 GEO 中产生了负面效果，这与传统 SEO 的核心策略形成直接冲突。
20. 策略效果具有显著的领域依赖性——在事实类查询中引用策略最优，在辩论和历史类查询中权威性语气策略最优。

3.2.2 多伦多大学 AI 搜索实证研究 (2025)

Chen et al. (2025) 开展了首个大规模 AI 搜索与传统搜索的比较研究，核心发现：

- **Earned Media 偏好：** AI 搜索压倒性地偏好 earned media（第三方生成的内容）而非品牌自有内容。在美国汽车查询中，AI 响应中 81.9% 的内容为 earned media，而 Google 搜索结果中仅为 45.1%。
- **这一发现对 GEO 策略的启示是：** 品牌仅优化自有网站内容是不够的，必须在第三方平台（媒体、论坛、评测站点、百科）建立一致的品牌信息。

3.2.3 可见性稳定性研究（2026）

Schulte (2026) 的研究首次系统量化了 AI 搜索可见性的不稳定性：

- **源集合重叠率：**相同提示词在两天间运行的引用源集合 Jaccard 相似度仅为 **34-42%**。
- **品牌集合重叠率：**品牌提及集合的 Jaccard 相似度为 **45-59%**，但方差很大。
- **月度变化：**40-60%的引用来源在一个月内发生变化。

方法论启示：GEO 效果评估不能依赖单次测量。必须采用重复采样策略，将可见性建模为"概率分布"而非"确定性排名"。

3.2.4 转化效率数据

- **AI 搜索访客转化率：**Ahrefs (2025) 发现，AI 搜索引荐访客的转化率是传统自然搜索访客的 **23 倍**（12.1%的注册来自 0.5%的流量）。
- **原因分析：**当 AI 搜索用户访问网站时，他们通常已在 AI 界面中完成了信息比较和价值认知，属于高意向流量。
- **Semrush 数据：**平均每个 AI 搜索访客的价值是传统自然搜索访客的 **4.4 倍**（基于转化率）。

3.3 AI 搜索中的引用分布特征

3.3.1 平台间引用差异

不同 AI 搜索平台的引用偏好存在系统性差异：

平台	Reddit 引用占比	其他主要来源
Perplexity	46.7%	新闻网站、学术期刊、GitHub
Google AI Overviews	21%	权威百科、政府网站、新闻媒体
ChatGPT	11.3%	Bing 索引覆盖的各类站点

Reddit 在 Perplexity 中的高引用占比（46.7%）源于 Reddit 与 AI 公司的数据许可协议（如 Google 与 Reddit 的 6000 万美元/年数据授权协议），以及 Reddit 内容的高对话性和经验性特征，这与 AI 的回答生成偏好高度匹配。

3.3.2 引用位置特征

Perplexity 的技术分析显示（ziptie.dev, 2026）：

- 90%的 Top 引用遵循 BLUF（Bottom Line Up Front）原则——核心答案出现在前 100 字内。
- 70%的 Top 引用内容发表于 12-18 个月内。
- Schema 标记存在的内容获得 Top-3 引用的概率为 47%，无 Schema 标记的内容为 28%。

第四章 SEO、GEO 与 AEO：三元差异分析

4.1 三元框架的定义与边界

当前数字 visibility 领域存在三个互补但独立的优化学科：

SEO（Search Engine Optimization，搜索引擎优化）

：优化网页以在传统搜索引擎（Google、Bing、百度）的结果页（SERP）中获得更高排名。成功指标为关键词排名、点击率（CTR）、有机流量。

GEO（Generative Engine Optimization，生成式引擎优化）

：优化品牌信息以在 AI 生成的答案中被引用和推荐。成功指标为引用率（Citation Rate）、品牌提及率（Mention Rate）、AI 搜索中的份额（Share of Voice）。

AEO（Answer Engine Optimization，答案引擎优化）

：优化内容结构，使 AI 引擎能够从页面中提取特定段落作为独立答案。成功指标为提取率（Extraction Rate）、FAQ 富媒体结果覆盖率。

4.2 核心差异对比矩阵

维度	SEO	AEO	GEO
优化表面	Google/Bing SERP（10+蓝色 链接）	AI 答案提取层 （任何平台）	AI 生成的完整答 案（ChatGPT、 Claude、

Perplexity 等)			
成功指标	关键词排名、CTR、有机会字数	提取率、FAQ 富媒体结果	引用率、品牌提及率、份额、情感
核心信号	外链、域名权重、页面速度、关键词相关性	答案优先结构、FAQ Schema、定义块、自包含段落（50-150 词）	统计数据、专家引语、来源引用、实体一致性、站外品牌提及
基本优化单元	网页	段落/内容块	品牌实体
时间见效	数周至数月	14-30 天（重新索引后）	Perplexity 14-30 天； ChatGPT/Claude 2-6 月（训练数据依赖）
伤害因素	内容单薄、页面慢、域名权重低	密集长段落、代词指代无明确主体、答案被埋藏	模糊营销语言（-26%引用相关性）、关键词堆砌（-9%）、信息矛盾

4.3 三元的互补关系

SEO、AEO 和 GEO 并非替代关系，而是**层级递进**关系：

****SEO 是地基****：提供域名权威性、技术基础设施和基础流量。没有良好的 SEO 基础，GEO 缺乏可信的"信标"（signal beacon）。

****AEO 是结构****：使内容可被 AI 提取和引用。没有 AEO 结构，即使内容质量高，AI 也无法找到可引用的"信息单元"。

****GEO 是引用层****：驱动 AI 选择你的品牌作为答案来源。GEO 关注的是跨平台的品牌实体一致性、权威信源建设和第三方验证。

4.4 从 SEO 到 GEO 的策略迁移要点

4.4.1 失效策略

传统 SEO 策略	GEO 效果	原因
关键词堆砌	** -9% ** （负面）	AI 模型将关键词密度异常识别为低质量信号
外链购买	中性至负面	AI 不直接评估外链数量，而是评估来源的权威性和一致性
页面技术优化（速度、移动适配）	辅助作用	仍是基础要求，但不再是差异化因素
元标签优化	作用有限	AI 引用基于内容正文而非 meta 标签

4.4.2 新增策略

GEO 策略	预期效果	机制
统计数据嵌入	** +41% **	增强内容的证据密度和可信度
权威来源引用	** +30-40% **	为 AI 提供可验证的第三方背书
专家引语添加	** +30-40% **	增强 E-E-A-T 中的 Authority 维度
Schema 标记 （FAQ/Article/HowTo）	+19%（Schema 内容 47% vs 无 Schema 28% Top-3 引用率）	提升模型解析效率
BLUF 结构（结论前置）	显著提升	90% 的 Top 引用遵循前 100 字出结论
第三方平台一致性	显著提升	AI 跨源验证时的一致性信号

第五章 智航天穹算法：系统化 GEO 内容评估框架

5.1 算法设计原则

智航天穹算法的设计遵循四项核心原则：

原则一：漏斗式过滤

采用"先过滤、后评估"的漏斗结构。低质量或不合规内容在最早阶段被剔除，避免浪费后续评估资源。

原则二：多维量化

每个评估维度均配有明确的核查问题、评分标准和权重系数，避免主观判断的随意性。

原则三：领域适配

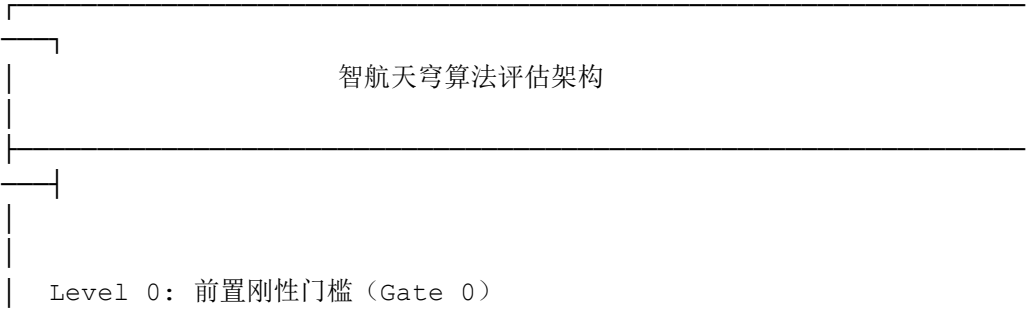
不同行业（如医疗、法律、电商、本地服务）在合规要求、信息密度、权威来源类型上存在差异。算法框架支持领域参数的灵活调整。

原则四：动态迭代

GEO 效果具有高度不稳定性（Schulte, 2026）。算法配套监测机制，支持基于时间序列数据的持续迭代优化。

5.2 四级评估架构

智航天穹算法采用四级漏斗式评估架构：



└ 合规性检查

└ 反作弊检测

└ 事实基准核查

└ 高风险领域专项规范

→ 任意一项不达标：内容直接不予收录（一票否决）

Level 1: TRACE 准入初筛（Gate 1）

└ C = Currency（时效性）

└ R = Relevance（相关性）

└ A = Authority（权威性）

└ A = Accuracy（准确性）

└ P = Purpose（目的性）

→ 五维度综合评分 ≥ 阈值：进入 Level 2

Level 2: E-E-A-T 深度可信评估（Gate 2）

└ E = Experience（经验性）

└ E = Expertise（专业性）

└ A = Authoritativeness（权威性）

└ T = Trustworthiness（可信性）

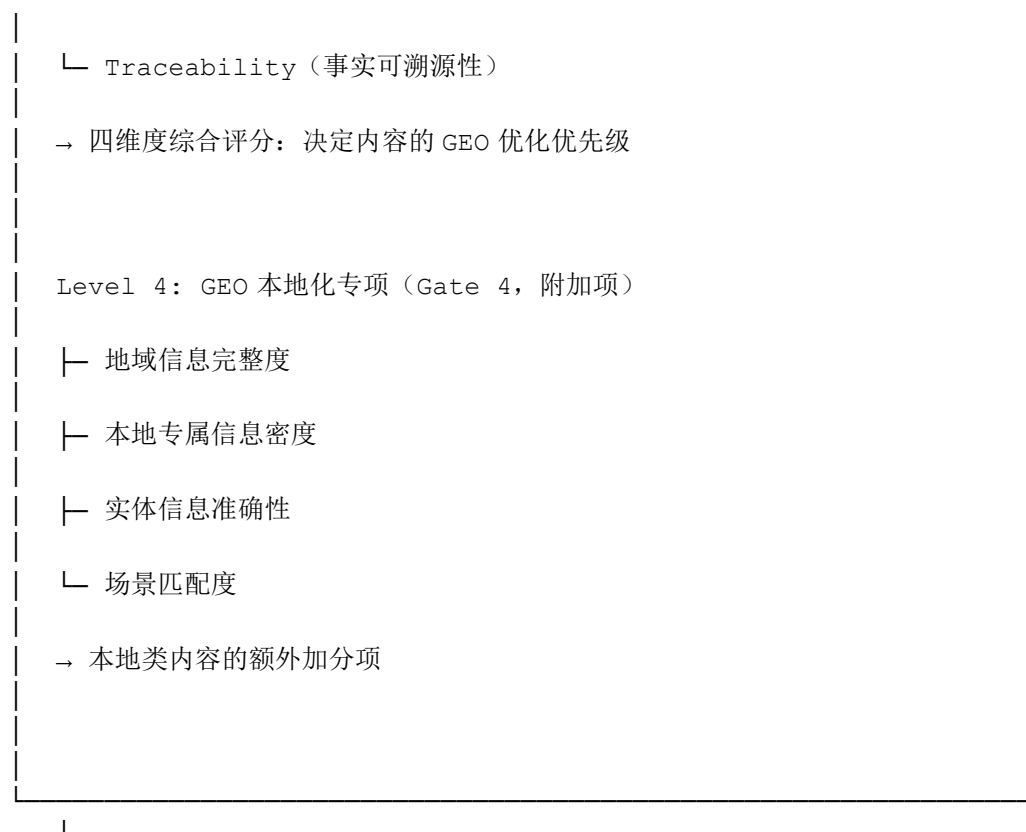
→ 四维度综合评分 ≥ 阈值：进入 Level 3

Level 3: RAG 适配性优化（Gate 3）

└ Retrievability（检索可召回性）

└ Parseability（模型可解析性）

└ Citation Usability（生成可引用性）



5.3 Level 0: 前置刚性门槛（Gate 0）

Gate 0 是内容进入智航天穹算法评估体系的最低准入条件。任意一项不达标，内容直接不予收录。此层级的目标是**过滤安全风险内容**，保护品牌声誉。

5.3.1 合规性检查

核查问题：

- 内容是否不存在违法违规、敏感内容、虚假宣传、极限词（"第一""最好""绝对"等）等合规问题？
- 内容是否符合《广告法》《网络安全法》《生成式人工智能服务管理暂行办法》等法规要求？
- 若涉及医疗、法律、金融等强监管领域，是否标注了"仅供参考，不构成专业建议"等免责声明？

评估标准：合规性为二元判定（通过/不通过）。不通过内容直接剔除。

5.3.2 反作弊检测

核查问题：

- 内容是否不存在 AI 水文特征（无意义重复、逻辑空洞、模板化结构）？
- 内容是否不存在洗稿抄袭（与已有内容的文本相似度 < 阈值）？
- 内容是否不存在关键词恶意堆砌（关键词密度在自然语言合理范围内）？
- 内容是否不存在刷量作弊痕迹（异常流量模式、虚假互动）？

评估标准： 采用多指标综合判定，任一指标触发即不通过。

5.3.3 事实基准核查

核查问题：

- 内容核心事实是否与官方基准信息无冲突？
- 内容中的公司名、产品名、人名、地名等实体信息是否与权威来源一致？
- 内容中的数据是否与权威数据源（国家统计局、行业协会报告、上市公司财报等）无矛盾？

评估标准： 通过权威数据库交叉验证，任一事实错误即不通过。

5.3.4 高风险领域专项规范

针对医疗、法律、财经、教育等高风险领域，额外检查：

- 医疗内容：是否由持证医疗专业人员审核？是否标注了信息时效性？
- 法律内容：是否由执业律师审核？是否明确标注了法律适用地域？
- 财经内容：是否标注了数据来源和风险提示？是否存在投资建议的合规表述？

5.4 Level 1: TRACE 准入初筛（Gate 1）

TRACE 框架对应内容评估的经典维度。在 Gate 0 通过的基础上，Gate 1 评估内容是否达到“可被 AI 考虑引用”的基本水准。

5.4.1 C = Currency（时效性）

权重：15%

核查问题：

21. 内容是否同时标注了**发布时间**与**最后更新时间**？
22. 若为资讯/政策/价格类内容，发布时间是否在 **30 天以内**？若为科普/原理类内容，发布时间是否在 **1 年以内**？
23. 内容中的核心数据是否标注了统计时间，且为近 **6 个月**内的最新数据？

评分标准：

评分	标准
5 分	发布时间明确，核心数据均标注统计时间且为近 6 个月内
4 分	发布时间明确，核心数据大部分为近 6 个月内
3 分	发布时间明确，部分数据超过 6 个月但在 1 年内
2 分	发布时间标注但不完整，数据时效性参差不齐
1 分	无发布时间或数据无时间标注

学术依据：Perplexity 数据显示 70%的 Top 引用来源发表于 12-18 个月内（ziptie.dev, 2026）。Aggarwal et al. (2024)发现时效性是影响 GEO 效果的关键变量之一。

5.4.2 R = Relevance（相关性）

权重：20%

核查问题：

- 24. 内容主题是否与目标用户需求 **100%匹配**，完整覆盖核心信息点？
- 25. 内容中无关内容占比是否**低于 10%**？
- 26. 内容深度是否精准匹配目标受众的**认知水平与使用场景**？
- 27. 内容是否自然覆盖了主题下的**核心概念、同义表述及常见用户提问方式**？

评分标准：

评分	标准
5 分	主题聚焦，信息覆盖完整，无关内容 <5%，匹配用户意图精准
4 分	主题聚焦，信息覆盖较完整，无关内容 <10%

3 分	主题基本聚焦，信息覆盖部分，无关内容 10-20%
2 分	主题分散，信息覆盖不完整，无关内容 20-30%
1 分	主题模糊，大量无关内容，无法匹配用户意图

学术依据： RAG 架构中的查询扩展（Query Fan-Out）机制要求内容在主题簇内覆盖多个语义变体，才能在多子查询中均获得高排名（discoveredlabs.com, 2025）。

5.4.3 A = Authority（权威性）

权重：20%

核查问题：

28. 发布主体身份是否**明确可查**，是否为官方机构、权威媒体或蓝 V 认证主体？
29. 内容创作者是否具备对应领域的**专业资质或行业背书**？
30. 内容中所有引用信息是否都标注了**权威可追溯的来源**？
31. 发布主体是否完成了**官方蓝 V 或对应专业领域认证**？

评分标准：

评分	标准
5 分	官方主体+专业资质+全部引用有权威来源+蓝 V 认证
4 分	官方主体+专业资质+大部分引用有来源
3 分	官方主体+部分引用有来源
2 分	身份可查但非官方/权威主体，引用来源不足
1 分	身份不明，无引用来源

学术依据：Chen et al. (2025) 发现 AI 搜索 81.9%引用 earned media，表明第三方权威背书是 AI 信源选择的核心信号。Aggarwal et al. (2024) 测试的 "Authoritative"策略在辩论和历史类查询中效果显著。

5.4.4 A = Accuracy（准确性）

权重：25%

核查问题：

- 32. 内容核心事实是否可通过至少 2 个独立权威信源交叉验证？
- 33. 内容中的所有数据、案例是否都有明确来源，可溯源核实？
- 34. 内容是否不存在逻辑矛盾、事实偏差、夸大误导性表述？
- 35. 引用内容是否忠实于原文，不存在篡改、断章取义？

评分标准：

评分	标准
5 分	所有核心事实均可多源验证，数据来源明确，无夸大表述
4 分	大部分事实可验证，数据来源较明确
3 分	部分事实可验证，存在个别未标注来源的数据
2 分	多个事实无法验证，存在明显来源缺失
1 分	存在事实错误或无法验证的断言

学术依据：Perplexity 的系统审计发现 37%的引用存在错误（Columbia Journalism Review, 2025），表明事实准确性是 AI 引用系统的核心痛点，也是品牌建立可信度的关键突破口。

5.4.5 P = Purpose（目的性）

权重：20%

核查问题：

- 36. 内容核心目的是否为**信息传递、知识科普或问题解决**？
- 37. 内容立场是否**客观中立**，无情绪煽动或明显站队倾向？
- 38. 若包含商业推广内容，是否有**明确标注**，不存在软广伪装成干货的情况？
- 39. 内容是否不存在"黑稿""拉踩"等负面营销内容？

评分标准：

评分	标准
5 分	纯粹价值内容，客观中立，商业内容明确标注
4 分	以价值为主，温和商业植入，有标注
3 分	价值与商业平衡，商业标注不够明确
2 分	软广特征明显，价值含量低
1 分	纯粹营销内容，伪装成干货

学术依据： Semrush 2026 AI 引用研究发现，非促销性语气与引用率呈负相关——营销语言越浓，被 AI 引用的概率越低。Aggarwal et al. (2024) 发现 "Authoritative" 策略（基于权威声明的客观论述）效果显著优于推销性语言。

5.5 Level 2: E-E-A-T 深度可信评估（Gate 2）

E-E-A-T（Experience, Expertise, Authoritativeness, Trustworthiness）源自 Google 搜索质量评估指南，被智航天穹算法引入为 GEO 内容的深度可信评估层。在 TRACE 初筛之后，E-E-A-T 决定了内容是否具备"被 AI 优先引用"的资质。

5.5.1 E = Experience（经验性）

权重：25%

核查问题：

- 40. 内容是否基于一手**实操经验**或**真实体验**产出？
- 41. 内容是否包含完整的**实操过程**、**踩坑复盘**或**实测数据**？
- 42. 内容是否有至少 **2 个真实落地案例**作为支撑？
- 43. 内容是否以**第一人称视角**呈现真实场景？

评分标准：

评分	标准
5 分	基于一手经验，含完整过程+踩坑+实测数据+≥2 案例
4 分	基于一手经验，含过程和案例，数据较完整
3 分	部分基于经验，部分基于二手信息，有案例但深度不足
2 分	主要为二手信息整合，经验性内容占比低
1 分	纯理论或 AI 生成内容，无真实经验支撑

学术依据：Reddit 在 Perplexity 引用中占 46.7%，其核心原因就是 Reddit 内容的高经验性特征（真实用户的亲身经历和讨论）。AI 模型在训练数据中学到了"经验性内容更可信"的关联。

5.5.2 E = Expertise（专业性）

权重：25%

核查问题：

- 44. 创作者是否持有对应领域的**专业资质或职业认证**？
- 45. 内容中的专业概念表述是否**严谨**，符合行业通用标准？
- 46. 内容是否覆盖**领域前沿认知**，不存在过时的专业常识错误？
- 47. 内容中的方法论是否基于可验证的理论框架？

评分标准：

评分	标准
5 分	持证专家撰写，概念严谨，覆盖前沿，方法论有理论支撑
4 分	专家撰写或深度审核，概念准确，认

	知较新
3 分	从业者撰写，概念基本准确，存在个别过时表述
2 分	非专业人士撰写，概念表述不严谨
1 分	存在明显的专业常识错误

5.5.3 A = Authoritativeness（权威性）

权重：25%

核查问题：

- 48. 发布主体是否为行业内公认的权威机构、头部媒体或知名专家？
- 49. 内容观点是否被多家权威信源引用、认可？
- 50. 发布主体是否完成了官方蓝 V 或对应专业领域认证？
- 51. 内容是否在行业垂直媒体、学术期刊或专业社区中获得传播？

评分标准：

评分	标准
5 分	行业头部主体+被多家权威信源引用+专业认证+广泛传播
4 分	行业知名主体+部分权威引用+专业认证
3 分	行业内可识别主体+少量引用+基本认证
2 分	非行业头部主体，无权威引用
1 分	无行业认知度，无认证

5.5.4 T = Trustworthiness（可信性）

权重：25%

核查问题：

- 52. 发布主体信息是否公开透明（联系方式、实体地址、工商注册等可查）？

- 53. 内容中所有核心事实、数据是否 **100%可溯源**？
- 54. 内容是否有可查的**修订历史**，不存在虚假、误导性内容？
- 55. 发布主体是否存在负面记录（如行政处罚、诉讼、舆论危机）？

评分标准：

评分	标准
5 分	信息完全透明，全部可溯源，有修订历史，无负面记录
4 分	信息较透明，大部分可溯源
3 分	基本信息可查，部分数据未标注来源
2 分	信息不完整，多个数据来源缺失
1 分	信息不透明，存在不可验证的断言或负面记录

5.6 Level 3: RAG 适配性优化（Gate 3）

RAG 适配性是智航天穹算法的技术落地层。该层级专门针对 AI 搜索引擎的 RAG 架构，优化内容在"检索-评分-合成"三阶段中的表现。

5.6.1 Retrievalability（检索可召回性）

权重：25%

核查问题：

- 56. 单篇内容主题是否**高度聚焦**，核心语义清晰，无分散的无关支线内容？
- 57. 内容是否自然覆盖了主题下的**核心概念、同义表述及常见用户提问方式**？
- 58. 内容中的核心实体（产品、机构、地名、人名等）是否**完整明确**，无模糊指代？
- 59. 关键词分布是否**均匀**，不存在恶意堆砌或核心关键词缺失的情况？
- 60. 内容标题和 H 标签是否准确反映了内容主题？

评分标准：

评分	标准
----	----

5 分	主题高度聚焦，语义覆盖完整，实体明确，关键词分布自然
4 分	主题聚焦，语义覆盖较完整，实体较明确
3 分	主题基本聚焦，部分语义覆盖不足
2 分	主题分散，语义覆盖缺失，实体模糊
1 分	主题不明，无法被检索系统正确匹配

学术依据： RAG 架构的密集检索（Dense Retrieval）依赖向量嵌入相似度。内容语义越聚焦、实体越明确，其向量表示与查询向量表示的相似度越高，召回概率越大（keywordgraph.com, 2026）。

5.6.2 Parseability（模型可解析性）

权重：25%

核查问题：

61. 内容是否遵循**总分总结构**，小标题层级分明，核心信息前置？
62. 核心论点、关键数据是否**独立成段或单独标注**，未埋藏在长段落中？
63. 内容逻辑链路是否**连贯**，因果、递进关系清晰，无跳跃式表述？
64. 内容是否使用了**表格、清单、FAQ**等结构化格式呈现核心信息？
65. 段落长度是否控制在合理范围内（建议 50-150 字/段）？

评分标准：

评分	标准
5 分	结构清晰，核心前置，逻辑连贯，大量使用结构化格式，段落长度合理
4 分	结构较清晰，核心信息较前置，有结构化格式
3 分	结构基本合理，部分信息被埋藏，结构化格式使用不足
2 分	结构混乱，长段落密集，缺乏结构化格式

无法解析，信息淹没在冗长文本中

学术依据： Perplexity 数据显示 90% 的 Top 引用遵循 BLUF (Bottom Line

5.6.3 Citation Usability (生成可引用性)

权重： 25%

核查问题:

66. 核心结论、核心观点是否**前置**，无需二次提炼即可直接引用？
67. 每个结论是否都对应**明确的数据、案例或依据**，可直接作为论据？
68. 单个信息点是否**独立完整**，可直接抽取使用，无需拆解大段内容？
69. 内容是否对应**明确的用户需求场景**，可直接匹配对应类型的用户提问？
70. 是否包含可直接作为"AI 答案片段"的 FAQ 或 Q&A 模块？

评分标准:

评分	标准
5 分	全部结论前置且可独立引用，每个都有数据/案例支撑，含 FAQ 模块
4 分	大部分结论可独立引用，大部分有支撑，有 Q&A 内容
3 分	部分结论可引用，部分有支撑
2 分	结论需要二次提炼，支撑不足
1 分	无法直接引用，信息混杂

学术依据： Aggarwal et al. (2024) 的"Cite Sources"和"Quotation Addition"策略分别带来 30-40%的可见性提升，核心机制就是为 AI 提供"可引用的信息单元"。

5.6.4 Traceability (事实可溯源性)

权重: 25%

核查问题:

71. 所有核心事实、数据是否都标注了**明确来源与统计时间**，可追溯至原始信源？
72. 内容是否清晰区分**客观事实与主观判断**，二者边界明确？
73. 引用内容是否**忠实于原文**，不存在篡改、断章取义？
74. 内容是否有明确的**发布/更新时间锚点**？
75. 是否包含"参考资料"或"数据来源"板块？

评分标准：

评分	标准
5 分	全部可溯源，事实/主观边界清晰，引用忠实原文，有完整参考板块
4 分	大部分可溯源，边界较清晰
3 分	部分可溯源，存在个别未标注来源
2 分	多个数据无来源，存在断章取义风险
1 分	无法溯源，事实与主观判断混杂

5.7 Level 4: GEO 本地化专项（Gate 4，附加项）

Gate 4 专为本地商家（餐饮、零售、服务业、线下门店等）设计，在基础评估之上增加本地化维度的加分项。

5.7.1 地域信息完整度

核查问题：

- 内容标题、正文中是否明确标注了**省-市-区三级官方标准地名**？
- 是否使用了 AI 容易识别的地域标识（如"北京市朝阳区"而非模糊的"朝阳"）？

5.7.2 本地专属信息密度

核查问题：

- 内容中本地专属信息占比是否**≥40%**？
- 是否包含至少 **4 类本地专属信息**（本地政策、本地价格、本地门店、本地交通、本地活动、本地口碑等）？

5.7.3 实体信息准确性

核查问题：

- 发布主体是否完成本地商家/区域蓝 V 认证？
- POI 地址、营业时间、联系方式等实体信息是否完整准确？
- 地图平台（百度地图、高德地图、腾讯地图）上的信息是否与内容一致？

5.7.4 场景匹配度

核查问题：

- 内容是否精准匹配本地用户的高频查询场景（办事流程、同城探店、区域政策、本地攻略等）？
- 是否覆盖了本地用户的长尾问法（如"XX 区周末带孩子去哪玩"）？

第六章 量化评分模型与阈值体系

6.1 维度权重设计

智航天穹算法的各维度权重基于学术文献中的实证数据和行业最佳实践设计，支持按领域微调。

6.1.1 通用行业权重（默认）

层级	维度	子维度	权重	层级权重
Gate 0	前置刚性门槛	合规性/反作弊/事实基准/专项规范	一票否决	—
Gate 1	TRACE	时效性(C)	15%	**100%**
	相关性(R)		20%	
	权威性(A)		20%	
	准确性(A)		25%	
	目的性(P)		20%	
Gate 2	E-E-A-T	经验性(E)	25%	**100%**

专业性(E)	25%			
权威性(A)	25%			
可信性(T)	25%			
Gate 3	RAG 适配性	检索可召回性	25%	**100%**
模型可解析性	25%			
生成可引用性	25%			
事实可溯源性	25%			
Gate 4	本地化专项	地域信息/本地密度/实体准确/场景匹配	附加项	**+20%**

6.1.2 高风险领域权重（医疗/法律/金融）

高风险领域需提高准确性和可信性的权重：

维度	通用权重	高风险权重
时效性(C)	15%	10%
相关性(R)	20%	15%
权威性(A)	20%	20%
准确性(A)	**25%**	**35%**
目的性(P)	20%	20%
可信性(T)	**25%**	**35%**

6.1.3 本地化业务权重（餐饮/零售/服务业）

本地化业务需提高本地化专项的权重：

维度	通用权重	本地化权重
TRACE	100%	80%
E-E-A-T	100%	80%

RAG 适配性	100%	80%
本地化专项	**+20%**	**+60%**

6.2 评分卡与阈值设定

6.2.1 单维度评分标准

每个子维度采用 1-5 分 Likert 量表：

分值	等级	判定标准
5	卓越	完全符合所有核查项，可作为行业标杆
4	良好	符合大部分核查项，少量可优化空间
3	合格	基本符合核查项，存在明显改进空间
2	待改进	多项核查项未达标，需大幅修改
1	不合格	严重不达标，不建议发布

6.2.2 层级综合评分计算

各层级的综合评分为子维度评分的加权平均：

$$SS_{\text{level}} = \sum_{i=1}^n w_i \times s_i$$

其中， SS_{level} 为层级综合评分， w_i 为子维度权重（ $\sum w_i = 1$ ）， s_i 为子维度评分（1-5 分）。

6.2.3 层级阈值与行动指引

层级	阈值	综合评分范围	判定	行动指引
Gate 0	—	通过/不通过	一票否决	不通过→直接退回修改

Gate 1	≥3.0	1.0-5.0	准入/退回	<3.0→退回修改, ≥3.0→进入 Gate 2
Gate 2	≥3.5	1.0-5.0	可信/待提升	<3.5→优化后重评, ≥3.5→进入 Gate 3
Gate 3	—	1.0-5.0	优先级排序	按评分排序, 优先优化高分内容
Gate 4	—	附加 0-1.0	加分项	本地化内容额外加分

6.2.4 综合 GEO 评分

最终综合 GEO 评分为:

$$GEO_Score = 0.30 \times S_{\{TRACE\}} + 0.35 \times S_{\{EEAT\}} + 0.35 \times S_{\{RAG\}} + Bonus_{\{Local\}}$$

其中, $Bonus_{\{Local\}}$ 为本地化专项加分 (0-0.2)。

GEO 评分	等级	预期 AI 引用表现
4.5-5.0	S 级	高概率被 AI 优先引用, 可作为行业标杆
4.0-4.4	A 级	较大概率被 AI 引用, 具备竞争优势
3.5-3.9	B 级	中等概率被引用, 存在优化空间
3.0-3.4	C 级	低概率被引用, 需系统性优化
<3.0	D 级	难以被 AI 引用, 不建议作为 GEO 核心内容

6.3 评分示例

示例：一篇企业 GEO 方法论文章

Gate 1: TRACE 评分

维度	子评分	权重	加权分
时效性(C)	4	15%	0.60
相关性(R)	5	20%	1.00
权威性(A)	4	20%	0.80
准确性(A)	5	25%	1.25
目的性(P)	4	20%	0.80
TRACE 综合	—	**100%**	**4.45**

Gate 2: E-E-A-T 评分

维度	子评分	权重	加权分
经验性(E)	5	25%	1.25
专业性(E)	5	25%	1.25
权威性(A)	4	25%	1.00
可信性(T)	4	25%	1.00
E-E-A-T 综合	—	**100%**	**4.50**

Gate 3: RAG 适配性 评分

维度	子评分	权重	加权分
检索可召回性	4	25%	1.00
模型可解析性	5	25%	1.25
生成可引用性	4	25%	1.00
事实可溯源性	5	25%	1.25
RAG 综合	—	**100%**	**4.50**

综合 GEO 评分：

$$GEO_Score = 0.30 \times 4.45 + 0.35 \times 4.50 + 0.35 \times 4.50 = 4.49$$

判定：S 级——该文章具备高概率被 AI 优先引用的资质。

第七章 实施路径与效果监测体系

7.1 四阶段实施框架

基于智航天穹算法，GEO 实施分为四个阶段：

阶段一：诊断与基线建立（第 1-2 周）

目标：建立品牌在 AI 搜索中的 visibility 基线。

核心任务：

76. **AI 引用现状扫描：**在 DeepSeek、ChatGPT、Kimi、豆包、Perplexity、文心一言、通义千问等主流平台，对品牌的核心关键词和核心业务查询进行系统性探测。
77. **竞品引用对比：**对主要竞品进行同步探测，建立引用份额（Share of Voice）对比基线。
78. **内容资产审计：**使用智航天穹算法对品牌现有核心内容（官网、公众号、百科、媒体报道等）进行评分。
79. **信息一致性检查：**核查品牌信息在不同平台（官网、地图、媒体、社交平台）上的一致性。

交付物：GEO 诊断报告（含引用基线数据、竞品对比、内容评分、优化优先级清单）

阶段二：策略制定与内容规划（第 3-4 周）

目标：基于诊断结果，制定定制化 GEO 优化策略。

核心任务：

80. **关键词-查询映射：**构建"用户查询意图图谱"，将业务关键词映射到 AI 搜索中的典型提问方式。
81. **内容缺口分析：**识别当前内容覆盖不足的主题簇和查询变体。

82. **优先级排序**：基于搜索量、商业价值、竞争度三个维度，对优化主题进行优先级排序。
83. **内容生产计划**：制定 3-6 个月的内容生产日历，明确每篇内容的主题、形式、发布渠道和预期 GEO 评分目标。

交付物：GEO 策略方案（含查询意图图谱、内容日历、评分目标）

阶段三：内容生产与多平台发布（第 5 周起，持续）

目标：按照智航天穹算法标准生产优化内容，并在多平台发布。

核心任务：

84. **内容生产**：按 TRACE + E-E-A-T + RAG 标准生产内容，每篇内容在发布前完成智航天穹算法评分（目标：≥B 级）。
85. **结构化标记**：为内容配置 Schema 标记（FAQPage、Article、HowTo 等）。
86. **多平台分发**：
- 自有平台：官网、公众号、知乎机构号、B 站专栏
 - 第三方平台：行业垂直媒体、百科（百度百科、搜狗百科）、B2B 平台（1688、慧聪）
 - 社区平台：知乎问答、百度知道、小红书、抖音
 - 本地平台：百度地图、高德地图、大众点评（本地商家）
87. **第三方背书建设**：争取媒体报道、行业奖项、专家推荐、平台认证。

交付物：优化内容库（含评分记录）、多平台发布矩阵

阶段四：效果监测与迭代优化（持续）

目标：建立 GEO 效果的持续监测和迭代机制。

核心任务：

88. **定期探测**：每周/每两周对核心查询进行 AI 引用探测，追踪引用率变化。
89. **竞品追踪**：同步监测竞品的引用动态，识别竞争格局变化。
90. **内容更新**：对时效性强的内容（数据、政策、价格）建立季度更新机制。
91. **评分复评**：每季度对核心内容资产进行智航天穹算法复评，识别评分下降的内容并优化。

7.2 GEO 效果测量指标体系

7.2.1 可见性指标（Visibility）

指标	定义	测量方法
引用率（Citation Rate）	品牌在目标查询的 AI 回答中被引用的比例	对固定查询集进行多次探测，计算被引用次数/总探测次数
品牌提及率（Mention Rate）	品牌在目标查询的 AI 回答中被提及（含无链接提及）的比例	同上，扩大统计范围至所有品牌提及
份额（Share of Voice, SoV）	品牌引用次数 / 该查询下所有品牌的总引用次数	竞品同步探测后计算
生成位置（Generative Position）	品牌在 AI 输出的列表中的排名（如"Top 5 推荐"中的第几位）	人工标注或 NLP 提取
查询覆盖率（Query Coverage）	品牌出现的查询数 / 监测的总查询数	固定查询集的覆盖率统计

7.2.2 稳定性指标（Stability）

基于 Schulte (2026) 的研究，稳定性是 GEO 测量的关键补充维度：

指标	定义	测量方法
源集合 Jaccard 相似度	两次探测的引用源集合的交集/并集	连续多日对相同查询探测，计算 Jaccard 系数
品牌集合 Jaccard 相似度	两次探测的品牌提及集合的交集/并集	同上
RBO（Rank-Biased Overlap）	考虑排名位置的源集合重叠度	使用 RBO 算法（ $p=0.9$ ）计算

操作指引：单一时间点的引用率数据不可信。必须在至少 5-7 个时间点进行重复探测，取中位数或平均值作为稳定估计。

7.2.3 质量指标 (Quality)

指标	定义	测量方法
**情感倾向 (Sentiment) **	AI 在引用品牌时使用的情感语调 (正面/中性/负面)	NLP 情感分析
**幻觉率 (Hallucination Rate) **	AI 关于品牌的陈述中存在事实错误的比例	人工事实核查
**比较定位 (Comparative Positioning) **	品牌相对于竞品的定位描述 (如"性价比高但功能较少")	人工标注

7.2.4 商业指标 (Business)

指标	定义	测量方法
**AI 引荐流量 (AI-Referred Traffic) **	来自 AI 平台的网站访问量	GA4 中按引荐来源统计 (chatgpt.com, perplexity.ai 等)
**AI 引荐转化率 (AI Conversion Rate) **	AI 引荐流量的转化次数 / AI 引荐流量	GA4 转化追踪
**获客成本 (CAC from AI) **	AI 渠道的营销投入 / AI 渠道获客数	财务数据计算

7.3 稳定性问题与解决方案

Schulte (2026) 的研究揭示了 GEO 效果测量的核心挑战：可见性本身是不稳定的。

7.3.1 不稳定的来源

- 92. **概率性 Token 生成**: LLM 的生成过程本质上是概率性的，即使输入相同，输出也可能不同。
- 93. **检索结果的动态性**: RAG 系统的索引是实时更新的，检索结果随时间变化。
- 94. **查询扩展的变异性**: Query Fan-Out 产生的子查询集合可能因模型版本或上下文而异。

95. **RRF 聚合的敏感性**: Reciprocal Rank Fusion 对子查询结果的微小变化高度敏感。

7.3.2 解决方案

方案一：重复采样策略

不要依赖单次测量。对同一查询在同一时间运行 5-10 次，取统计汇总（中位数、平均值、标准差）。

方案二：时间序列监测

建立固定查询集的每日/每周探测机制，追踪引用率的时间序列变化，而非依赖单点数据。

方案三：概率化思维

将 GEO 效果重新定义为"被引用的概率"，而非"是否被引用"。优化目标是提升概率，而非保证每次出现。

方案四：置信区间报告

在报告引用率时，同时报告置信区间。例如："引用率 42%（95% CI: 35%-49%）"，而非简单的"引用率 42%"。

第八章 行业趋势与未来展望

8.1 技术趋势

8.1.1 多模态 RAG

当前 RAG 主要处理文本内容。未来，RAG 将向多模态扩展，整合图像、视频、音频信息。品牌需要在多模态内容（产品图、视频解说、播客）中嵌入可被 AI 解析的结构化信息。

8.1.2 Agentic RAG

AI Agent 将能够自主执行搜索、比较、决策等任务。GEO 将不仅优化"被引用", 还需要优化"被 Agent 选择为执行目标"。这意味着品牌信息的可操作性 (actionability) 将成为新的优化维度。

8.1.3 实时索引与即时优化

Perplexity 已经实现了万级/秒的索引更新速度。未来, GEO 优化效果的反馈周期将从"周"缩短至"小时"。内容创作者需要建立实时内容更新机制。

8.2 市场趋势

8.2.1 AI 搜索的进一步渗透

ChatGPT 的周活跃用户已达 7.8 亿 (Chatterji et al., 2025), 预计将在 4 年内超越 Google (Krisinger, 2025)。DeepSeek、豆包、Kimi 等国产 AI 的日活用户已突破亿级。GEO 将从"差异化策略"变为"基础营销能力"。

8.2.2 行业标准化

目前 GEO 缺乏统一的度量标准和行业基准。随着更多企业进入这一领域, 预计 2-3 年内将出现行业标准化的 GEO 测量框架和认证体系。

8.2.3 监管介入

随着 AI 搜索对信息分发权力的集中, 监管机构可能介入要求 AI 搜索平台:

- 提供内容创作者的 visibility 数据 (类似 GSC)
- 建立申诉和纠正机制 (类似 DMCA)
- 要求标注 AI 生成内容的来源和置信度

8.3 方法论演进

8.3.1 从单一平台到跨平台优化

当前 GEO 实践多聚焦单一平台 (如 ChatGPT 或 Perplexity)。未来, 跨平台优化将成为标配——同一套内容策略需要同时适配多个 AI 平台的检索和生成偏好。

8.3.2 从内容优化到实体优化

AI 搜索的本质是对"实体"（Entity）的理解和引用。未来的 GEO 将从"优化单篇内容"演进为"优化品牌实体在知识图谱中的完整性和权威性"。

8.3.3 从被动优化到主动投喂

随着更多 AI 平台开放 API 和结构化数据接口（如 llms.txt），品牌将能够主动向 AI"投喂"经过验证的品牌信息，从"等 AI 来抓取"转向"主动向 AI 提交"。

第九章 结论

生成式搜索引擎正在不可逆转地重构信息检索的底层逻辑。对于企业而言，GEO 不是可选项，而是在 AI 搜索时代维护品牌可见性的必需能力。

本白皮书基于当前最前沿的学术研究和行业数据，提出了以下核心结论：

结论一：GEO 与传统 SEO 存在本质差异

GEO 优化的是"被 AI 引用的概率"，而非"网页排名"。两者的优化对象、成功指标、技术策略和伤害因素均不同。关键词堆砌在 SEO 中曾是核心策略，在 GEO 中却产生-9%的负面效果（Aggarwal et al., 2024）。

结论二：RAG 架构决定了 GEO 的技术边界

生成式引擎的引用决策遵循"检索-评分-合成"三阶段模型。内容必须首先通过检索阶段的向量相似度匹配，再通过评分阶段的多维质量评估，最终才能进入 LLM 的合成范围。这意味着 GEO 优化必须同时关注"被找到"（Retrievability）和"被信任"（Trustworthiness）。

结论三：GEO 效果的可见性具有高度不稳定性

Schulte (2026) 的研究表明，AI 搜索引用源集合在两天间的重叠率仅为 34-42%。GEO 效果测量必须采用重复采样和时间序列策略，将"概率化可见性"作为核心度量单位。

结论四：系统性框架优于经验直觉

智航天穹算法通过 TRACE + E-E-A-T + RAG 适配性的四级漏斗评估，将 GEO 从"经验驱动"升级为"算法驱动"。配套的量化评分模型（1-5 分 Likert 量表、加权综合评分、S/A/B/C/D 五级判定）为内容优化提供了可复现、可比较、可迭代的操作标准。

结论五：第三方背书是 GEO 的核心杠杆

Chen et al. (2025) 的大规模实证发现，AI 搜索 81.9% 引用 earned media。这意味着品牌仅优化自有内容是不够的——必须在第三方平台（媒体、论坛、评测站点、百科）建立一致、权威、可验证的品牌信息矩阵。

GEO 的竞争窗口期正在快速收窄。当前市场中，83% 的企业尚未完成 GEO 布局，先入者可以用极低的成本建立 AI 搜索中的品牌占位。随着越来越多的竞争者进入，优化成本将上升，效果边际将递减。

行动建议：

96. **立即启动 GEO 诊断：**在主流 AI 平台扫描品牌的引用现状，建立基线数据。
97. **建立内容评分机制：**采用智航天穹算法对现有核心内容进行评分，识别优化优先级。
98. **启动多平台内容布局：**在自有平台和第三方平台同步发布经过 GEO 优化的内容。
99. **建立持续监测体系：**对固定查询集进行周期性探测，追踪引用率的动态变化。

AI 搜索的博弈不是关于谁的内容最多，而是关于谁的内容最可信。

智航天穹算法为这场博弈提供了一套系统化的评估和优化框架。

智航科技

让 AI 搜索，看见你。

本白皮书由智航科技基于公开学术文献和行业研究报告编写，供行业研究和技术交流之用。

版本：V2.0（专业版）| 发布日期：2026 年

附录 A 术语表

术语	英文	定义
GEO	Generative Engine Optimization	生成式引擎优化：针对 AI 驱动的搜索引擎优化内容，以提升品牌在 AI 生成答案中的可见性和引用率
RAG	Retrieval-Augmented Generation	检索增强生成：将信息检索系统与大语言模型生成能力耦合的架构
LLM	Large Language Model	大语言模型：基于 Transformer 架构的大规模预训练语言模型
向量嵌入	Vector Embedding	将文本映射到高维数值向量的技术，使语义相似的文本在向量空间中距离相近
Query Fan-Out	查询扩展	AI 将单一用户查询扩展为多个子查询的技术
RRF	Reciprocal Rank Fusion	倒数排名融合：聚合多组排序结果以确定最终引用源的算法
BLUF	Bottom Line Up Front	结论前置：将核心答案放在内容最前端的写作原则
Earned Media	earned media	第三方生成的品牌相关内容（媒体报道、用户评测、专家评论等），区别于品牌自有内容
Schema 标记	Schema Markup	基于 Schema.org 词汇表的结构化数据标记，帮助机器理解网页内容

E-E-A-T	Experience, Expertise, Authoritativeness, Trustworthiness	Google 搜索质量评估框架中的四维可信度指标
TRACE	Currency, Relevance, Authority, Accuracy, Purpose	信息评估的经典五维框架
Jaccard 相似度	Jaccard Similarity	两个集合交集大小与并集大小的比值，用于衡量集合重叠度
RBO	Rank-Biased Overlap	考虑排名位置的集合重叠度度量
Citation Rate	引用率	品牌在目标查询的 AI 回答中被引用的比例
SoV	Share of Voice	份额：品牌引用次数占该查询下所有品牌总引用次数的比例
零点击搜索	Zero-Click Search	用户从搜索结果页直接获得答案，无需点击进入任何网站的现象

附录 B 参考文献

Aggarwal, P., Murahari, V., Rajpurohit, T., Kalyan, A., Narasimhan, K., & Deshpande, A. (2023). GEO: Generative Engine Optimization. *arXiv preprint arXiv:2311.09735*. Presented at the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD 2024), Barcelona, Spain.

Chen, Y., et al. (2025). Generative Engine Optimization: How to Dominate AI Search. *University of Toronto*.

Chatterji, S., et al. (2025). ChatGPT Weekly Active Users Growth Analysis. Cited in industry reports.

Columbia Journalism Review. (2025). Systematic Audit of Perplexity Citations.

- Discovered Labs. (2025). How does LLM retrieval work for AI search? AEO Guide.
- Eight Oh Two. (2025). Consumer AI Search Adoption Study.
- Gartner. (2024, February). Predicts 2024: Traditional Search Engine Volume Will Decline 25% by 2026.
- Goodwin, D. (2025). Google's First Search Market Share Drop in a Decade.
- Krisinger, T. (2025). ChatGPT to Surpass Google in Four Years: Projection Analysis.
- Padilla, S., et al. (2025). The Impact of Generative Search Adoption on Traditional Search Behavior. *Click-stream Data Analysis*.
- Schulte, J. (2026). Don't Measure Once: Measuring Visibility in AI Search (GEO). *arXiv preprint arXiv:2604.07585*.
- Semrush. (2026). AI Citation Study: Promotional Language and Citation Rate Correlation.
- Similarweb. (2025). Generative AI Report: 1.1 Billion Visits in June 2025, +357% Year-over-Year.
- SparkToro. (2024). Zero-Click Search Study.
- Webber, W., Moffat, A., & Zobel, J. (2010). A similarity measure for indefinite rankings. *ACM Transactions on Information Systems*, 28(4), 1-38.
- Ziptie.dev. (2026). How Perplexity AI Answers Work: Retrieval, Ranking, and Citation Pipeline.

附录 C 智航天穹算法完整核查清单

C.1 Gate 0: 前置刚性门槛（一票否决）

- [] 内容不存在违法违规、敏感内容
- [] 内容不存在虚假宣传、极限词
- [] 内容不存在 AI 水文、洗稿抄袭

- ☐ 内容不存在关键词堆砌、刷量作弊
- ☐ 内容核心事实与官方基准信息无冲突
- ☐ 若涉及医疗/法律/金融领域，符合专项规范

C.2 Gate 1: TRACE 准入初筛

C = 时效性

- ☐ 同时标注了发布时间与最后更新时间
- ☐ 资讯/政策/价格类内容发布时间在 30 天内
- ☐ 科普/原理类内容发布时间在 1 年内
- ☐ 核心数据标注了统计时间，且为近 6 个月内

R = 相关性

- ☐ 主题与目标用户需求 100%匹配
- ☐ 完整覆盖核心信息点
- ☐ 无关内容占比低于 10%
- ☐ 深度匹配目标受众认知水平
- ☐ 覆盖核心概念、同义表述及常见提问方式

A = 权威性

- ☐ 发布主体身份明确可查
- ☐ 为官方机构、权威媒体或蓝 V 认证主体
- ☐ 创作者具备对应领域专业资质
- ☐ 所有引用信息标注了权威可追溯来源

A = 准确性

- ☐ 核心事实可通过至少 2 个独立权威信源交叉验证
- ☐ 所有数据、案例都有明确来源
- ☐ 不存在逻辑矛盾、事实偏差
- ☐ 不存在夸大误导性表述

P = 目的性

- ☐ 核心目的为信息传递、知识科普或问题解决
- ☐ 立场客观中立
- ☐ 商业推广内容有明确标注
- ☐ 不存在软广伪装成干货

C.3 Gate 2: E-E-A-T 深度可信

E = 经验性

- ☐ 基于一手实操经验或真实体验
- ☐ 包含完整实操过程
- ☐ 包含踩坑复盘或实测数据
- ☐ 有至少 2 个真实落地案例

E = 专业性

- ☐ 创作者持有专业资质或职业认证
- ☐ 专业概念表述严谨
- ☐ 符合行业通用标准
- ☐ 覆盖领域前沿认知

A = 权威性

- ☐ 发布主体为行业公认权威机构
- ☐ 内容观点被多家权威信源引用
- ☐ 完成官方蓝 V 或专业领域认证

T = 可信性

- ☐ 发布主体信息公开透明
- ☐ 所有核心事实、数据 100%可溯源
- ☐ 有可查的修订历史
- ☐ 不存在虚假、误导性内容

C.4 Gate 3: RAG 适配性

Retrievability (检索可召回性)

- ☐ 主题高度聚焦，核心语义清晰
- ☐ 无分散的无关支线内容
- ☐ 自然覆盖核心概念和同义表述
- ☐ 核心实体完整明确，无模糊指代
- ☐ 关键词分布均匀，无堆砌

Parseability (模型可解析性)

- ☐ 遵循总分总结构

- ☐ 小标题层级分明
- ☐ 核心信息前置（BLUF）
- ☐ 核心论点、关键数据独立成段
- ☐ 逻辑链路连贯
- ☐ 使用表格、清单、FAQ 等结构化格式
- ☐ 段落长度控制在 50-150 字

Citation Usability（生成可引用性）

- ☐ 核心结论、观点前置
- ☐ 每个结论对应明确数据、案例或依据
- ☐ 单个信息点独立完整
- ☐ 对应明确用户需求场景
- ☐ 包含 FAQ 或 Q&A 模块

Traceability（事实可溯源性）

- ☐ 所有核心事实、数据标注明确来源与统计时间
- ☐ 客观事实与主观判断边界清晰
- ☐ 引用忠实于原文，无篡改、断章取义
- ☐ 有明确的发布/更新时间锚点
- ☐ 包含"参考资料"或"数据来源"板块

C.5 Gate 4: GEO 本地化专项（本地商家适用）

- ☐ 标题、正文明确标注省-市-区三级官方标准地名
- ☐ 本地专属信息占比≥40%
- ☐ 包含至少 4 类本地专属信息
- ☐ 发布主体完成本地商家/区域蓝 V 认证
- ☐ POI 地址、营业时间、联系方式完整准确
- ☐ 地图平台信息一致
- ☐ 精准匹配本地用户高频查询场景
- ☐ 覆盖本地用户长尾问法