

—
综合得分
未评分

核心结论

以 AIVO 四维指数为评估框架，对品牌在六大 AI 平台的可见度做体检：示例品牌在 AI 可见度上「4/6 平台已收录、尚可」，竞品优势度因差异化标签清晰而居中，品牌完善度（官网 / 百科等基建）是最大短板，品牌正向度正面为主但权威背书不足。

● 现有优势

本次实测未呈现明显优势项（各项均未达优势阈值）。

● 主要风险

本次实测未触发风险项。

0%

AI 认知率 (0/0)

0

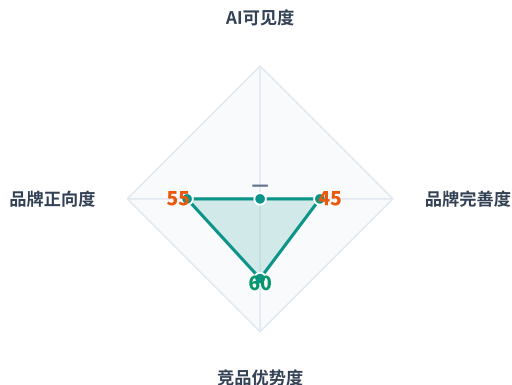
本次实测问句×平台

—

正向表述占比

—

AI 点名的信源平台数



综合 GEO 指数

环比

对比上期

四维得分：AI可见度(—)、品牌完善度(45)、竞品优势度(60)、品牌正向度(55)。综合指数 —（未评分）。

四维得分明细

维度	得分	说明
AI可见度	—	
品牌完善度	45	最大短板 无官网 / 无百科，AI 缺少可引用的权威实体信息。
竞品优势度	60	中等 差异化标签清晰，但对比竞品缺乏量化优势证据。
品牌正向度	55	正面为主 舆情正面，但权威背书与案例佐证不足。

本期问句结构分布（实测）

按关键词类型统计本次提交的高意图问题构成。参考模板中的「客群年龄 / 消费能力 / 复购」等画像评分属于行业调研结论，本报告无该数据源，故不填写（不估算、不代填）。

通用词

100%（3条）

核心人群与典型搜索场景

◆ 通用词

- 品牌名 是做什么的
- 品牌名 靠谱吗
- 品牌名 哪家好

典型搜索问题（本次实测问句）

问题	关键词类型
品牌名 是做什么的	通用词
品牌名 靠谱吗	通用词
品牌名 哪家好	通用词

0

总查询

0

正常介绍次数

0%

认知率

六大 AI 平台提及情况

本次无有效回答，无法绘制。

问题提及明细（共 3 条）

问题	所属搜索场景	正常介绍
品牌名 是做什么的	通用词	0 / 0 平台
品牌名 靠谱吗	通用词	0 / 0 平台
品牌名 哪家好	通用词	0 / 0 平台

品牌完善度（AI 能否讲清品牌）

该维度衡量各平台对品牌主体、业务、所在地的描述是否具体一致；平台间认知冲突时自动降级为「存疑」。本次得分 **45**（最大短板）。

无官网 / 无百科，AI 缺少可引用的权威实体信息。

官方网站 / 门店落地页 / 自媒体矩阵的全网核查需要外部检索数据源，本报告未接入，故不填写评分与明细（不估算、不代填）。

AI 回答点名的信源平台（0 个平台 / 累计 0 次）

本次各平台回答正文中未点名任何可识别的信源平台。

引用信源明细（真实链接，需浏览器实测采集）

本次未采集到逐条引用链接。当前引擎调用的是各平台对话接口，拿到的是模型对品牌的认知与描述，不含逐条网页链接；如需真实引用 URL，需针对该品牌运行浏览器实测采集后再查看本报告。

0

AI 点名的其他机构数

0

本品牌被排在首位的次数

—

平均提及位次

—

推荐类问题被点名率

可见度对比（本品牌 vs 同问题下 AI 点名的机构）

品牌	可见度	市场份额	威胁等级	说明
（本品牌）	—	—	—	本次实测对象

「市场份额」「威胁等级」需要行业调研与份额数据源，本报告未接入，故全部留空（不估算、不代填）；竞品可见度需对其单独发起实测后方可比较。

主要竞品占位（按问题逐条看 AI 首选机构）

未获取到占位数据。

● 本品牌占位优势

本次实测中，本品牌未在任何问题的首选位出现。

● 本品牌占位短板

未触发明显短板项。

收录与提及率统计（分平台）

平台	测试问句	被正常介绍	提及率	模型
合计	0	0	0%	

口径：向各平台官方对话接口逐条实问本次的高意图问题，「被正常介绍」= 该条回答能正常介绍该品牌（不含「查不到」与「仅按名称推测」）。

0%

正向表述占比

0%

中性表述占比

0%

负向信号占比

0

称查不到 / 无登记的条数

● 正向驱动

本次未识别到稳定正向驱动项。

● 负面驱动

本次未识别到负向信号。

AI 认知判定分布

真认识（正常介绍）（0 条）

0%

称查不到 / 无登记（0 条）

0%

信息不足（多为推测）（0 条）

0%

未提及但有保留表述（0 条）

0%

完全未提及（0 条）

0%

问题实录（逐条判定）

无有效回答记录。

合规与司法风险

该项需要工商 / 司法 / 监管公开数据的定向检索（如失信被执行、行政处罚、司法裁判文书等）。本报告由六大 AI 平台对话接口实时生成，未接入上述数据源，故本报告不含合规与司法风险核查，不做任何推测性描述。如需该项，需单独发起企业合规核查。

行动清单（按优先级）

本次未生成行动项。

执行路线图

P1 立即处理

0-2 周

本阶段无待办项。

P2 短期补齐

2-6 周

本阶段无待办项。

P3 持续建设

6 周以上

本阶段无待办项。

关于效果预期

本报告不提供「执行后分数将提升至 X」的预测值。各 AI 平台的回答由其模型与训练数据决定，任何机构都无法精确判定或保证 AI 平台的具体引用位置与提升幅度；上述行动项给出的是改进方向与目标口径，不构成效果承诺。优化效果建议以「下期实测对比本期基线」的方式验证。

基于AIVO四维评分模型（AI可见度 + 品牌完善度 + 竞品优势度 + 品牌正向度）对六大 AI 平台（豆包 / DeepSeek / Kimi / 通义千问 / 腾讯元宝 / 文心一言）实时实测生成。