

2026

2026AI品牌资产建设发展白皮书

发布单位：国家广告研究院

2026年6月

发布单位

国家广告研究院

联合发布单位

UCL 人工智能研究中心（英国）

世界未来科技发展联合会（荷兰）

欧洲时报英国分社

国际品牌网

四川省人工智能研究院

弗若斯特沙利文

编撰团队

王昕 徐振廷 刘灿辉 李强 刘玲妤

饶庆升 陈妮 施丹 张丹妮

目录

前言 7

第一章 全球 AI 生态与认知入口格局..... 12

1.1 全球 AI 入口结构：主流平台与能力边界.....12

1.2 中国 AI 入口生态：超级平台整合与商业闭环13

1.3 全球 AI 品牌入口迁移趋势：从搜索到 AI 助手的演进.....15

第二章 范式转移：从“流量主导权”到“认知主导权” 17

2.1 用户行为的变化：从“找资料”到“要结论”17

2.2 搜索引擎优化（SEO）的范式重构：从排序逻辑到生成逻辑18

2.3 “认知主导权”的含义：能否被系统稳定理解与引用.....19

2.4 品牌在生成式 AI 语境下的“认知稀释”与隐性出局风险20

第三章 AI 品牌资产建设：从“被发现”到“被引用” 22

3.1 为什么需要 AI 品牌资产建设：不是换概念，而是换了竞争逻辑.....22

3.2 AI 品牌资产建设三个层次：被发现、被选择、被引用.....22

3.3 AI 品牌资产建设的对象：从“内容产出”转向“品牌信息治理”24

3.4 AI 品牌资产（AIBE）：品牌资产的新“存放位置”26

3.5 AIBE 的五层路径：从可识别到可治理.....26

第四章 AI 如何“看见”品牌 29

4.1 生成式 AI 并非“更聪明的搜索”29

4.2 认知来源：长期认知、行为边界与即时证据30

4.3 工作流程：检索路由、召回、重排与生成31

4.4 为什么品牌内容会被“过滤”32

第五章 全球 AI 品牌资产发展现状与治理框架..... 34

5.1 全球发展概况：区域差异与共性特征.....34

5.2 中国 AI 品牌资产行业现状：从概念热潮走向规则收敛.....	35
5.3 全球 AI 治理与品牌可信框架：国际规则与标准演进.....	37
5.4 行业典型误区：战略错位、执行失真与评估失焦.....	39
第六章 行业落地指南：不同赛道的“认知主导权”争夺	41
6.1 企业服务（SAAS/工业/B2B）：从“流量曝光”转向“专业代入”	41
6.2 零售与电商：从“搜索排名”转向“消费意图精准匹配”	42
6.3 本地生活（餐饮/酒旅/生活服务）：从“高分评价”转向“即时决策入口”。	43
6.4 内容 IP 与教育文旅：从“单向输出”转向“交互式生态搭建”	43
6.5 强合规行业（医疗/金融/法律）：从“模糊回答”转向“权威合规源”	44
第七章 可信知识网络建设方法论：从认知基础设施到 GEO 运行机制.46	
7.1 从内容供给到认知基础设施：为什么企业需要可信知识网络	46
7.2 可信知识网络的定义：把分散信息升级为可被 AI 采用的知识体系.....	46
7.3 三层架构：可信知识网络的整体结构.....	47
7.4 六层建设路径：从认知诊断到认知巩固.....	47
7.5 权威高质量语料库：可信知识网络的知识底座	48
7.6 从语义定位到答案资产：可信知识网络的操作方法.....	50
7.7 从知识建设走向认知治理.....	51
第八章 新评估体系：从点击指标到认知治理指标.....	53
8.1 为什么需要重建指标.....	53
8.2 核心衡量对象：品牌“影响答案”的能力	53
8.3 答案份额（SoA）指标：在关键问题中“占了多少答案位置”	54
8.4 引用率指标：品牌是否被当作“可信证据”使用	55
8.5 认知一致性指标：AI 系统对品牌的判断是否稳定可复现.....	56
8.6 情感倾向指标：生成式环境下负面影响更集中	57
8.7 AIBE：将指标收敛为“可管理的 AI 品牌资产”	57
8.8 AIBV 指数体系：面向生成式内容的统一评估框架.....	58

第九章 行业规范与治理框架：从原则倡议到实施细则 60

9.1 为什么行业规范正在成为 AI 品牌资产建设的必要组成..... 60

9.2 行业基本原则：从“行业倡议”升级为“共同规范” 61

9.3 禁止性行为清单：什么不应再被视为“有效方法” 62

9.4 推荐性建设规范：什么样的建设才是可持续的 63

9.5 评估规范：AIBV 如何成为行业共同语言 65

9.6 认证机制：从语料审核到服务分级 66

9.7 标准化推进路径：从原则表达走向行业实施细则 68

第十章 全球 AI 品牌资产建设十大原则 70

10.1 可识别原则 70

10.2 权威信源原则..... 71

10.3 多语言一致性原则..... 71

10.4 结构化知识原则..... 72

10.5 可验证原则 72

10.6 风险边界原则..... 73

10.7 AI 可引用原则..... 73

10.8 全球平台一致性原则 74

10.9 持续治理原则..... 75

10.10 合规与伦理原则 75

总结..... 77

附录 79

附件 AI 品牌资产建设表现指数体系（AIBV 1.0） 81

第一章 总则 82

第二章 术语与定义..... 84

第三章 总体框架..... 85

第四章 AIP：AI 表现基础指数..... 86

第五章 AIC：AI 建设指数.....89

第六章 AIR：AI 风险与稳定性指数91

第七章 UAF 与 MCI92

第八章 综合指数.....92

第九章 数据采集与测量规范93

第十章 防操纵与质量保障94

第十一章 结果表达与等级划分96

第十二章 使用规范与合规要求.....96

第十三章 版本管理与扩展应用98

结语.....99

前言

过去二十年，数字营销主要建立在搜索曝光、内容分发与点击跳转之上。品牌竞争的重点，是争取被看见、被点击，并在用户进入页面之后，通过内容、产品信息、价格说明、用户评价、客服互动与转化设计影响其判断与行动。

但随着生成式人工智能的快速普及，信息获取与决策形成的关键环节正在发生深刻变化。越来越多用户不再先浏览大量链接、再自行归纳判断，而是直接通过 AI 问答、AI 搜索、AI 推荐、AI 助理和智能代理等语义交互场景获取结论、形成初筛，并推进下一步行动。对用户而言，这意味着更低的信息检索成本、更短的判断链路与更高的决策效率；对品牌而言，则意味着竞争前线正在从“被看见”前移到“被理解、被引用、被纳入考虑”。

《2025 年 AI 新时代品牌价值管理白皮书》指出，生成式人工智能正推动信息获取模式由“链接索引”向“直接答案生成”转变。当前，超过八成中国互联网用户已经形成日常使用 AI 搜索或 AI 助手获取信息的习惯，标志着品牌传播的核心入口与竞争焦点正逐步向人工智能推荐机制迁移。

当前，生成式内容优化市场正处于快速增长阶段。随着生成式 AI 与 AI 搜索应用的普及，2025 年中国生成式内容优化（GEO/AIBE 建设）市场规模已增长至约 57 亿元，并预计在 2026 年达到约 137 亿元，年增长率超过 140%。从全球范围来看，生成式内容优化市场规模预计将在 2030 年达到约 5122 亿元，2025—2030 年复合年均增长率达 65.3%。这一增长

趋势表明，生成式内容优化正逐步成为企业重构数字营销体系与品牌资源配置的重要战略领域。

算力黑洞与认知入口的耦合：生成式 AI 的快速演进背后，是一场“算力黑洞”——单次推理成本、训练成本与推理峰值并发都在指数级抬升。这迫使 AI 平台不断收紧检索预算、压缩证据数量、提高单条证据的质量门槛。对品牌而言，这意味着未来“被引用”将比“被提及”更稀缺，低密度、低可信度内容会被算力机制天然过滤。换言之，AI 算力压力正在反向推高品牌高质量信息供给的稀缺性与战略价值。

这意味着，品牌面临的核心挑战，已不再只是传统意义上的搜索排名或媒体曝光，而是品牌是否能够被主流 AI 系统准确理解、合理归类、稳定表达，并在复杂问题与多轮交互中形成真实、可信、可验证的存在状态。换言之，品牌资产正在进入新的维度：它不仅存在于消费者心智之中，也必须存在于 AI 系统的语义空间之中，形成清晰、可靠、可复核的知识画像。

正是在这一背景下，本白皮书提出“**AI 品牌资产**”这一研究框架。我们将 **AIBE (AI Brand Equity, AI 品牌资产)** 作为总概念，用于描述品牌在主流 AI 大模型与 AI 应用场景中的综合价值表现。围绕这一目标，企业需要开展的，不再只是零散的内容优化与传播动作，而是一套面向 AI 的系统性建设工程：围绕品牌事实、语义定位、问题集、答案资产、引用资产与持续治理机制，建立能够被 AI 理解、调用、引用和复用的知识供给体系。

在这一体系中，可信知识网络（KNIT，Knowledge Network of Integrity & Trust）被界定为 AI 品牌资产建设的核心方法论与底层基础设施。它的作用，不在于直接定义品牌资产本身，而在于通过真实世界数据、权威研究、结构化图谱、第三方验证与可追溯来源，将企业分散、非结构化的信息组织为可被 AI 稳定理解、引用与复用的标准化可信知识体系。换句话说，AI 品牌资产建设真正竞争的，不再只是内容数量，而是知识供给的质量、结构、可信度与可验证性。

在此基础上，过去围绕 **GEO（生成式引擎优化）** 所形成的一系列实践，也需要被重新理解。本白皮书并不否定 GEO 作为方法的现实意义，而是尝试在更完整的框架下对其进行边界重构：凡是有助于提升品牌知识的真实性、结构化程度、可解释性与可引用性的建设性方法，都应被吸收为 AI 品牌资产建设中的有效动作；而那些依赖虚假投喂、伪造信源、数据污染与结果操纵的做法，则应被明确排除。也就是说，GEO 不再被视为一个独立的总框架，而被重新界定为 AI 品牌资产建设中面向生成式引擎的一组实施方法与实践路径。

与此同时，围绕 AI 品牌资产的讨论，也不能只停留在方法层面。随着黑帽 GEO、AI 投毒、虚假投喂、伪造证据链与误导性内容传播等问题不断出现，这一领域正在从概念热潮走向原则厘定、规则收敛与规范建设。真正值得鼓励的，不是对结果的操纵，而是对供给侧信息质量的系统性提升：提升品牌信息的结构化程度、真实性、准确性、透明度与可验证性，使 AI 在回答高价值问题时，能够基于更可信的知识网络形成更稳健的表达。

因此，本白皮书试图建立一套更完整的分析框架：

AIBE 是目标与资产，回答“品牌在 AI 中应形成什么样的价值存在”；

AI 品牌资产建设是系统工程，回答“企业需要建设哪些关键能力”；

KNIT 是核心方法论与底层基础设施，回答“这些能力如何被组织成可信知识体系”；

GEO 是面向生成式引擎的实施方法，回答“这些知识如何进入 AI 的检索、筛选、生成与引用链路”；

而 **AIBV** 则是统一评估框架，用于对 AIBE 的状态、建设能力与风险状况进行系统衡量。

我们希望强调，企业推进 AI 品牌资产建设，不应被理解为一种狭义的营销技巧升级，而应被理解为一套面向 AI 时代的品牌信息治理工程。它的核心，不是“让 AI 更偏爱你”，而是“让 AI 更准确地理解你、更审慎地引用你”。通过持续建设高质量、可复用、可引用、可校验的品牌知识体系，使品牌的专业能力、产品优势、适用边界与风险说明，能够以机器可读且可信的方式沉淀下来，并在不同模型、不同场景与不同时间窗口下保持更高的一致性与可解释性。这既是 AI 品牌资产建设的基础，也是未来行业规范、评估体系与标准化推进得以成立的前提。

本白皮书希望推动行业完成一次必要的认知转换：从“争夺结果”转向“建设资产”，从“流量博弈”转向“公信力经营”，从“行业机会表达”转向“行业标准前置表达”。在 AI 时代，真正能够穿越模型迭

代、平台变化与监管收敛的，不会是短期操纵，而是建立在真实、合规、透明与可追溯基础上的 AI 品牌资产。

AI 时代，品牌最终竞争的也许不再只是传播能力，而是“被世界可信地理解”的能力。谁能够率先完成从“内容供给”到“可信知识供给”的转型，谁就更有可能在未来 AI 主导的信息生态中建立长期而稳定的品牌影响力。



第一章 全球 AI 生态与认知入口格局

1.1 全球 AI 入口结构：主流平台与能力边界

在全球范围内，生成式人工智能已从单一的对话工具演变为多元化的认知入口，深刻重塑了信息分发的底层逻辑。当前全球 AI 入口结构呈现出“寡头模型+垂直应用”的生态格局。以 OpenAI 的 ChatGPT、Anthropic 的 Claude、Google 的 Gemini、Microsoft 的 Copilot 以及 xAI 的 Grok 等为代表的通用大模型，凭借其强大的泛化推理能力与庞大的用户基数，构成了全球用户获取泛知识、进行复杂推理与创意生成的核心底座。这些平台不仅承载着数以亿计的日常问答，更通过 API 生态与插件系统，成为连接品牌与用户的“超级接口”。

进一步看，模型竞争已经从“参数规模竞争”转向“算力—数据—产品三位一体的竞争”。OpenAI 通过与微软 Azure 深度绑定，构建产品化与生态闭环；Anthropic 借助 AWS 与 Google 双投资，主打安全对齐与合规可控；Google Gemini 凭借自研 TPU 实现最深的垂直一体化算力—模型—应用链路；xAI 则以 Colossus 集群（10 万+ H100）押注算力规模化；Meta 以 Llama 开源路线重塑算力分发逻辑。不同路线决定了不同模型对“信源”的偏好与可信度判断方式——这对品牌的跨平台一致性管理提出更高要求：同一条品牌事实必须能在不同模型架构下被稳定理解，并兼容差异化的证据筛选策略。

与此同时，以 Perplexity、You.com 为代表的 AI 原生搜索引擎，正在快速侵蚀传统搜索引擎的市场份额，成为用户获取即时信息、进行深度研究与事实核查的新入口。这类平台的核心优势在于其“答案优先”的交付模式与对信源透明度的强调，直接将品牌信息置于“被引用”而非“被点击”的竞争维度。此外，Salesforce Einstein、Adobe Firefly、GitHub Copilot 等面向垂直业务的 AI 助手，则将生成式能力深度嵌入特定业务流程，成为 B2B 品牌触达决策者的关键节点。

这些主流平台的能力边界正在不断扩展，从纯文本生成向多模态理解（图像、视频、音频）、实时联网检索（RAG）以及智能体（Agent）自主执行演进。例如，多模态模型使得品牌的产品视觉信息、使用场景视频能够被直接“理解”并融入答案生成；实时联网检索则要求品牌信息必须具备高时效性与高可信度，才能在动态更新的知识库中占据一席之地。对于品牌而言，这意味着信息分发的逻辑被彻底重构：平台不再只是信息流转的分发节点，而开始成为信息筛选、语义重组与行动建议生成的关键认知层。这意味着品牌必须适应不同平台在检索路由、证据采纳、生成偏好乃至伦理对齐上的差异，制定差异化的 AI 品牌资产策略，才能在全球 AI 生态中占据一席之地。

1.2 中国 AI 入口生态：超级平台整合与商业闭环

与海外市场相比，中国市场在生成式交互普及、平台整合与商业闭环形成方面具备更强的加速度，呈现出显著的“超级平台整合”特征。以元宝、Kimi、豆包、扣子、DeepSeek、文心一言、通义等为代表的本土 AI 应用，不仅在中文语义理解、文化语境适配上具备天然优势，更深度嵌入了现有的内容、社交、本地生活与电商生态。

在大模型底层，中国市场正在形成“三股力量”并存的格局：第一类是互联网大厂自研模型，如腾讯混元/元宝、字节豆包、阿里通义、百度文心；第二类是独立大模型公司，包括 DeepSeek、月之暗面 Kimi、智谱、MiniMax、阶跃星辰等；第三类是端侧与行业垂直模型，如华为盘古、商汤等。DeepSeek 以算法效率反向缓解算力压力的路径，正在影响全球模型训练范式；同时，昇腾、寒武纪等国产算力与国产大模型形成“耦合演进”。对品牌而言，中文语义空间将出现更多模型分叉，跨模型一致性管理的复杂度与成本将持续上升。

一方面，中国用户在移动端高频场景中更倾向于接受“直接给结论”的信息形态，尤其是在高频生活决策（如餐饮、出行、购物）与消费决策中。这种用户习惯使得 AI 答案与后续行动之间的转化路径极短。另一方面，超级平台生态使生成式答案更容易与内容、社交、搜索及交易链路无缝连接。例如，当用户通过 AI 助手询问“适合家庭出游的周末方案”时，系统不仅能生成行程建议，还能直接关联到票务预订、酒店比价、餐厅推荐等具体服务，答案不再只是“参考”，而已直接进入转化动作。

这种“基础设施化”的演进为品牌带来了新的门槛与机遇。仅在传统渠道“被看见”已不再足够，品牌更需要在生成式 AI 入口中保持可验证、可持续的存在。这意味着品牌必须与超级平台的知识库、推荐算法乃至商业规则进行深度协同，将自身的产品信息、服务能力、用户评价等数据，以结构化、可被 AI 调用的方式，整合进平台的“认知基础设施”之中。同时，这也要求品牌方具备更强的跨平台治理能力，确保在不同超级平台的 AI 语境下，品牌认知保持一致，并能有效承接从“答案”到“行动”的流量转化。

1.3 全球 AI 品牌入口迁移趋势：从搜索到 AI 助手的演进

全球用户的行为习惯正在经历不可逆的迁移：从“主动搜索”向“依赖 AI 助手”演进。这一趋势并非简单的工具替代，而是信息获取与决策模式的范式转移。用户越来越习惯于将复杂的决策需求（如跨国旅游规划、B2B 软件选型、跨语种文献研究、个人理财配置）直接抛给 AI 助手，要求其提供结构化的对比方案、个性化的推荐理由以及最终的可执行建议。

这一趋势导致品牌触点发生了根本性前移。在传统搜索时代，品牌通过 SEO 和 SEM 在搜索结果页拦截用户，竞争焦点是关键词排名与点击率。而在 AI 助手时代，品牌必须在 AI 的“预训练语料”和“实时检索知识库”中提前布局。AI 助手在生成答案时，其依据并非实时的全网搜

索，而是其内部知识图谱与经过筛选的高质量信源。如果品牌未能以结构化、高可信度、多语言适配的方式存在于 AI 的知识图谱中，或者其信息在实时检索中因质量低劣而被过滤，那么品牌将在全球用户的决策视野中遭遇“隐性出局”——即用户根本不会将其纳入考虑范围，更遑论比较与最终选择。

这种迁移对品牌资产建设提出了全新要求。品牌需要从“流量思维”转向“认知思维”，将资源投入到构建能够被 AI 稳定理解、准确归类、优先引用的知识体系上。这包括建立全球统一的品牌事实库、多语言标准语料、权威引用资产以及持续的认知监测与治理机制。只有如此，品牌才能在从搜索到 AI 助手的演进浪潮中，牢牢掌握认知主导权，将 AI 从潜在的“隐形过滤器”转化为品牌价值的“放大器”。

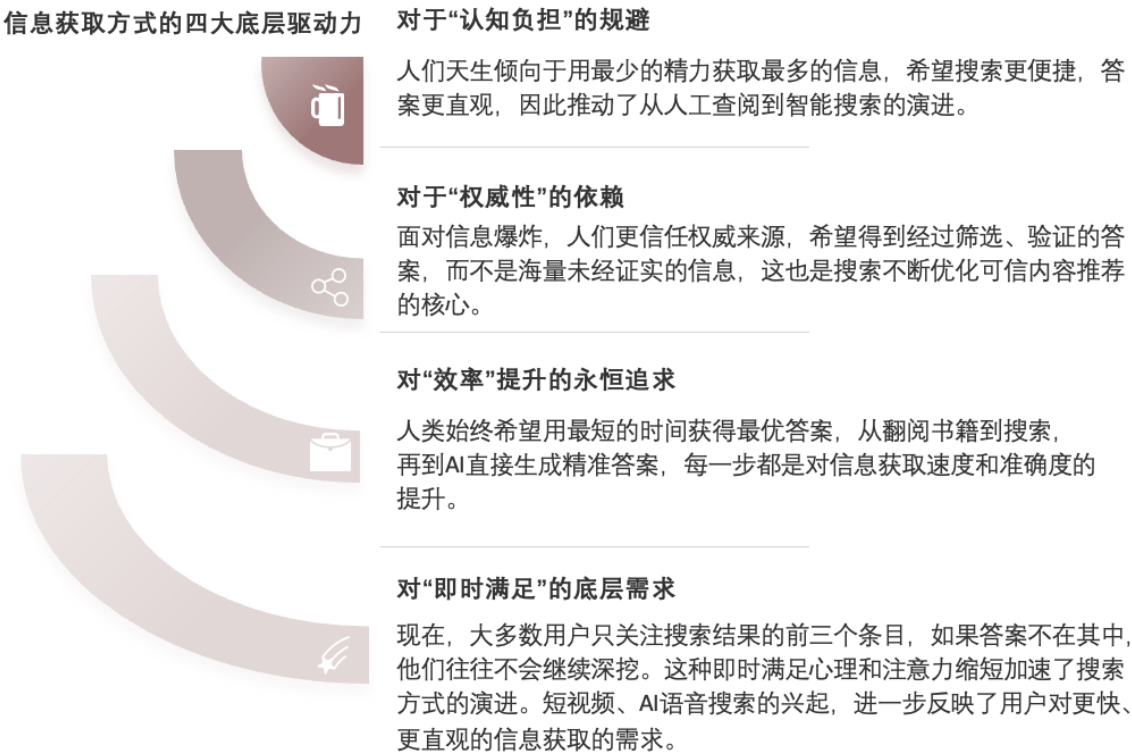
第二章 范式转移：从“流量主导权”到“认知主导权”

2.1 用户行为的变化：从“找资料”到“要结论”

在传统搜索时代，用户获取信息通常要经历一个相对完整的过程：先用关键词把问题“拆碎”，再打开多个链接交叉对比，最后由自己做出判断。这个过程看似麻烦，却支撑了互联网长期以来的主流商业模式，因为“点击”既是用户行动的记录，也是平台分配流量与品牌获取转化的基础。

生成式人工智能的普及，使这条路径发生了明显变化。用户越来越倾向于用更自然的方式把需求一次讲清楚，例如直接提出预算、时间、偏好、限制条件等，让系统给出“整体建议”和“下一步怎么做”。对用户而言，这意味着更少的搜索成本、更少的筛选成本；对品牌而言，意味着用户在“点开页面之前”就可能完成了初步筛选，甚至在打开页面之前就有倾向。于是，品牌进入用户考虑范围的关键节点被提前了：不再只取决于能否被点击，还取决于能否出现在答案里，并以合理的方式被解释。

图表 1 AI 重塑用户搜索体验



2.2 搜索引擎优化（SEO）的范式重构：从排序逻辑到生成逻辑

不少从业者把变化概括为“SEO 会不会被 AI 取代”。但更准确的说法是：SEO 依赖的关键条件正在变化。传统 SEO 之所以有效，是因为搜索入口相对稀缺，且用户必须通过链接跳转来获取信息；只要能在结果列表里占据更好的位置，就更有机会拿到点击与访问。然而，当生成式引擎以“直接答案”作为主要交付方式时，用户不必逐条点击链接也能得到结论，页面排序对决策的影响随之减弱，“点击”也不再是必经环节。

因此，品牌面临的风险不再只是“排得靠后”，而更可能是“根本不在答案里”。在生成式答案场景中，用户首先看到的是一个合成后的

结论与理由，而不是一组可供逐一选择的链接。这使得品牌竞争的重点从“争排名”逐步转向“争进入答案的资格”，即争取被系统采纳与引用。

需要强调的是，SEO 并未被取代，而是被“前置”了——传统 SEO 正在变成 RAG 抓取链路中的“原材料筛选层”，而真正的竞争已经发生在重排（Reranking）阶段。可观察的现象是：Google AI Overviews 上线后，部分高频问询型流量下降幅度超过 30%；ChatGPT Search、Perplexity 等 AI 原生入口，已开始直接挤压交易前的信息比较流量。由此，品牌需要建立一个新概念——“Pre-Click Brand Footprint（点击前品牌足迹）”，即衡量品牌在用户点击发生之前，已经在 AI 答案生成链路中被采纳与表达的程度。

2.3 “认知主导权”的含义：能否被系统稳定理解与引用

在传统互联网环境中，品牌竞争更多围绕“流量主导权”展开，即谁更能获得曝光、点击与进入页面的机会；而在 AI 逐步成为入口的背景下，更关键的资源开始转向“认知主导权”。本报告所称“认知主导权”，是在 AI 成为答案入口之后，品牌是否还能掌握自己被如何理解、如何判断、如何引用的主导权。本质上就是品牌在 AI 时代的“定义权”和“被采纳权”。

更严谨地说，认知主导权由以下三类能力共同构成。第一是可识别性：系统能否准确识别你的品牌实体，知道你属于哪个品类，不与竞品混淆。第二是相关性：在特定问题场景下，你是否与正确的需求、场景与关键词形成稳定关联，从而更容易被召回。第三是可信度与一致性：

当系统需要给出理由与建议时，你的内容能否作为可靠依据被引用，你在不同问题下的核心判断是否稳定一致。只有当这三方面形成闭环，品牌才更可能在生成式答案中获得持续、稳定的出现，并形成长期的推荐惯性。

可以用更通俗的话来理解：当用户问到关键问题时，系统是否“想得起你、说得清你、信得过你”。

这也意味着，品牌传播不再只需要面对用户的注意力与偏好，还需要面对生成式系统对信息的筛选逻辑。若品牌内容无法以“AI 可用”的方式进入证据链与推理链，即便在传统渠道仍保持一定声量，也可能在新入口中逐步变得不可见、不可选。

2.4 品牌在生成式 AI 语境下的“认知稀释”与隐性出局风险

在生成式引擎主导的场景中，品牌可能面临一种更隐蔽的出局方式：不是被明确否定，而是因为无法被系统稳定理解与引用，从而逐步被排除在关键问题的答案之外。这种变化往往缺乏传统意义上的警报信号，它不一定表现为一次明显的排名下滑或投放失效，而更可能表现为：在高价值问题中出现频率降低、被推荐的顺位下降、在答案中长期缺席。其影响往往要累积到一定程度，才会在获客成本、转化效率与品牌偏好上以更滞后、更难溯源的方式显现。

基于上述判断，本报告后续章节将进一步展开生成式引擎的机制：它如何检索与筛选信息、如何决定引用哪些内容，以及品牌应如何通过语义资产与答案资产进入证据链路，并建立更适配“答案时代”的指标

体系与治理框架，以支持生成式内容优化能力在中国市场走向规范化与可持续发展。

第三章 AI 品牌资产建设：从“被发现”到“被引用”

3.1 为什么需要 AI 品牌资产建设：不是换概念，而是换了竞争逻辑

在传统搜索时代，品牌竞争的重点是让用户找到信息并进入页面，后续再通过内容与转化设计影响决策。进入生成式 AI 时代后，用户往往在点击之前就已获得整合答案并形成初步判断，品牌影响认知的关键环节也由“页面打开之后”前移到“答案生成之前”。因此，企业今天所面对的问题，并不是简单地给旧方法换一个新名词，而是品牌建设的对象本身发生了变化。品牌需要建设的不再只是分散的传播内容，而是一套能够被 AI 系统稳定读取、合理归类、持续更新并审慎引用的品牌知识体系。这里真正变化的，不是“营销是否还重要”，而是品牌进入用户认知链路的方式发生了迁移：从依赖曝光与点击，转向依赖被系统正确理解、被纳入有效证据、被纳入候选集合的能力。

也正因如此，本白皮书并不将 AI 品牌资产建设理解为对生成式系统结果的操纵。相反，它强调的是一套更基础、也更长期的能力：品牌是否具备真实、准确、可验证、可复用的知识供给能力，是否能够在不同问题、不同场景与不同表达方式下，向 AI 系统提供更低歧义、更高可信度的品牌信息。与其说这是一次“流量技巧升级”，不如说这是一次品牌信息底座的重建。

3.2 AI 品牌资产建设三个层次：被发现、被选择、被引用

如果要用更简洁的方式概括 AI 品牌资产建设的核心目标，可以将其分为三个层次。

第一层是**被发现**。这意味着品牌至少要在 AI 的检索与召回过程中具备基本可见性：品牌名称能够被正确识别，核心产品与服务能够被正确归类，品牌不会因为信息过于零散、语义模糊或名称歧义而在问题路由阶段被遗漏。对很多品牌来说，这一层并不是理所当然的。过去依靠平台投放或自然流量获得的曝光，并不会自动转化为生成式 AI 场景中的可发现性，因为后者更依赖品牌事实信息的清晰度与语义映射的稳定性。

第二层是**被理解**。品牌被系统看到，并不意味着系统真正理解了品牌。AI 在回答高价值问题时，往往不仅需要“知道有这个品牌”，还需要知道品牌是什么、解决什么问题、适合什么场景、与什么品类或替代方案相关、边界在哪里、风险点是什么。如果品牌的信息供给缺乏结构化表达，或长期存在口径不一致、表述夸张、证据不足等问题，就容易出现误读、错引、过度简化或风险化标签集中的情况。

第三层是**被引用**。在生成式环境中，尤其是在决策价值高、风险敏感度高的问题场景下，AI 系统并不会仅凭“听说过”某个品牌就轻易给出明确建议。它更依赖能够提供出处与依据的材料，包括品牌事实、第三方报告、结构化 FAQ、可核验案例、权威媒体报道、认证与合规说明等。也就是说，品牌最终能否在 AI 输出中形成更稳固的存在，不仅取决于是否“出现过”，更取决于是否具备足够的证据支撑，使系统在表达品牌时“有依据可援引”。

从这个角度看，AI 品牌资产建设不是一个单点动作，而是一条递进路径：先解决品牌被找到的问题，再解决品牌被看懂的问题，最后解决品牌被合理引用的问题。三者缺一不可，且越往后，越要求品牌具备更高的信息质量、更强的知识组织能力与更清晰的边界表达能力。“从被发现到被引用”的真正深层含义，是从“竞争话语”转向“资产建设话语”——从抢入口走向建底座。

图表 2 AI 品牌资产建设的核心目标



3.3 AI 品牌资产建设的对象：从“内容产出”转向“品牌信息治理”

在很多企业的旧有理解中，品牌建设往往等于持续输出内容：做更多文章、发更多稿件、铺更多页面、争取更多曝光。这套逻辑在搜索时代有其合理性，因为链接排序与页面覆盖对流量分发具有直接作用。但在生成式 AI 环境中，单纯增加内容数量并不能自然转化为更强的品牌表

现，甚至在某些情况下，大量重复、同质、低质量的内容反而会增加系统理解难度，制造语义噪音，并拉低品牌信息的一致性与可信度。

因此，AI 品牌资产建设的对象，需要从单纯的“内容产出”转向更底层的“品牌信息治理”。这里的治理，不是抽象意义上的管理口号，而是指对品牌信息进行系统梳理、结构化整理、权威来源锚定、口径统一、边界清晰化与持续更新的过程。它强调的是：品牌提供给 AI 的不是一堆分散文本，而是一套能够相互印证、相互支持、层次清晰的知识体系。

这一建设对象的转变，也意味着企业要重新理解“语义资产”的含义。所谓语义资产，并不是为了让 AI 系统“更偏向你”，而是为了让 AI 系统在遇到与你相关的问题时，能够更低成本地完成识别、归类、比较、解释与引用。企业需要建设的不只是“说了很多”，而是“说得清楚、说得一致、说得有依据、说得可复核”。从这个意义上讲，AI 品牌资产建设本质上是一种品牌信息基础设施工程，其目标不是制造更高的噪音，而是降低理解成本、降低误读概率并提升可信表达的稳定性。

在语义资产管理的实践中，可信知识网络（KNIT）系统性地回应了 AI 时代企业认知管理的新挑战。（KNIT，Knowledge Network of Integrity & Trust）作为面向 AI 时代的企业级认知基础设施解决方案以真实世界数据、权威研究结论、结构化图谱、第三方验证结果与可追溯信息来源为核心构成要素，通过 AI 认知基线诊断、真实世界验证、权威事实锚定、知识结构化工程、可信内容扩散、持续监测与认知巩固的六层结构化工程，将企业分散的非结构化信息编织成可被 AI 稳定理解、

引用与复用的标准化知识体系。其核心目标是帮助企业掌握 AI 时代的认知主导权，确保自身被市场正确理解，让 AI 能基于准确前提稳定引用企业核心价值与关键数据，同时规避数据投毒、算法幻觉、认知稀释等系统性风险，将企业关键事实转化为长期可调用、可累积的可信数字资产。

3.4 AI 品牌资产（AIBE）：品牌资产的新“存放位置”

传统品牌资产讨论的是消费者心智中的沉淀，例如知名度、偏好、信任感、忠诚度等。这些维度在 AI 时代并不会失效，但品牌资产的“存放位置”正在发生扩展：除了存在于消费者心智之中，品牌还需要在 AI 系统中形成清晰、稳定、可复用的知识画像。AI 在回答问题时，会依据其已有知识、可检索证据与生成机制，对品牌进行归类、比较、总结并给出建议，这种“AI 系统对品牌的认知”会越来越直接地影响用户的第一印象、候选范围乃至最终决策。

需要强调的是，AIBE 的意义不在于把品牌建设从“人”完全转向“机器”，而在于承认 AI 已经成为连接品牌与用户的重要中间层。企业今天面对的，不只是用户会如何看待自己，也包括系统会如何理解自己、如何总结自己、如何把自己放入与其他品牌的比较框架之中。谁能够率先把品牌知识结构化、证据化、可验证化，谁就更有可能在 AI 时代形成更稳固、更可持续的品牌资产。

3.5 AIBE 的五层路径：从可识别到可治理

为了让 AIBE 这一概念具备更强的可操作性，本白皮书继续采用分层视角来描述品牌在 AI 环境中的成熟路径。

第一层是**可识别性**。品牌首先要能够被系统准确识别为一个明确实体，避免与其他品牌、品类或通用名词混淆，并能被归类到正确的业务范围与应用语境之中。没有可识别性，后续所有被理解、被引用与被比较都无从谈起。

第二层是**语境相关性**。品牌不仅要被识别出来，还要能够在相关问题场景中被稳定联想到。也就是说，当用户提出具体需求、限制条件、对比问题或决策情境时，品牌是否能够与这些情境形成合理关联，而不是仅在极少数机械问题中被提及。

第三层是**认知一致性**。在不同问题、不同表述、不同模型和不同时间窗口下，系统对品牌核心能力、适用边界、优势特征与风险提示的判断是否稳定一致。

第四层是**引用可信性**。品牌相关信息是否具备足够可信度，使 AI 系统愿意将其作为依据纳入回答。这并不是简单追求“被引用次数更多”，而是强调品牌是否拥有足够高质量的证据型内容、第三方可验证内容与权威来源，从而在高价值问题中被更审慎、更稳妥地表达。

第五层是**治理可持续性**。品牌不仅要在某个时间点表现良好，更要能够持续监测 AI 对自身的误读、错引、过时信息、风险化表达与口径分裂问题，并通过问题集、事实库、FAQ、证据补强与口径治理等方式进行持续修复。这使 AIBE 不再只是一个静态表现概念，而成为一个兼具建设性与治理性的品牌资产框架。

由此，AIBE 的价值不在于提供一个抽象标签，而在于帮助企业建立一套可复盘、可分层、可治理的品牌资产路径：从可识别，到可关联；

从可理解，到可信用；从可表达，到可治理。从长期来看，AI 时代的最大风险是品牌逐渐失去被系统稳定识别、理解与采纳的能力，并最终在用户认知形成之前“无声退出”。

第四章 AI 如何“看见”品牌

4.1 生成式 AI 并非“更聪明的搜索”

在传统搜索引擎时代，系统的主要职责是“把相关网页排序后交给用户自行选择”。因此，品牌竞争的重点集中在排名与点击：排名越靠前，获得访问与转化的机会越大。生成式 AI 的工作方式与此不同。它在面对用户问题时，目标并不是提供一份链接清单，而是给出一个相对完整、逻辑连贯、看起来可信且可执行的答案。在这个过程中，系统必须完成一系列判断：用户真正想解决的是什么问题，哪些信息可以作为可信依据，应该如何组织这些依据才能形成结论。也正因为生成式 AI 承担了“合成与判断”的角色，品牌的可见性不再主要取决于是否出现在某个列表位置，而取决于是否进入了系统构建答案所使用的证据材料之中。

认知可见性：AI 时代企业生存的底层逻辑。在探讨机制前，我们需要意识到“被 AI 看见”背后的商业量级，过去你漏掉一个点击，只是损失一个访客；现在如果你在 AI 生成的“综合结论”中缺席，你损失的是整个品类的认知主权。AI 不再是把“一堆选项”交给用户，而是直接给出“一个结论”。

还有一个反直觉但日益重要的事实：模型越强，最终保留的证据反而越少。因为长上下文推理的算力成本极高，平台会主动压缩“采纳证据数”，以控制单次推理开销。这意味着重排阶段的“幸存者偏差”正在加剧——未进入 Top-3 证据的内容，几乎等同于不存在。对品牌而言，

信息密度比信息数量更重要，单位 Token 的“事实承载力”正成为新的核心指标。

从品牌视角看，这意味着一个更现实的变化：很多用户在看到链接之前就已经获得结论并形成倾向，品牌如果无法进入答案生成环节，往往不是“排在后面”，而是“没有被纳入候选”。因此，要讨论 AIBE 的方法与路径，必须先理解生成式 AI 在“看见与采用信息”时到底依赖什么机制与规则。

4.2 认知来源：长期认知、行为边界与即时证据

生成式 AI 对品牌与世界的“理解”并非来自单一来源，通常可以概括为以下三类相互补充的机制。第一类是预训练形成的长期认知，即模型在训练阶段接触到的大量公共文本内容所沉淀的背景知识。这部分知识覆盖面广、权重较高，往往决定了大模型对行业与品牌的基本印象，但更新速度相对较慢。第二类是对齐与微调形成的行为边界，即大模型在上线前后通过偏好对齐、安全规范与平台策略等方式形成的回答倾向，它并不直接决定“知道什么”，但会影响“怎么说、是否建议、在什么场景下更谨慎”。第三类是检索增强生成（RAG）等机制带来的即时证据，即大模型在回答具体问题时，会从外部知识库或互联网内容中检索相关材料，并在有限的证据片段基础上生成回答，从而补足时效性与事实性。

对多数品牌而言，短期内最可操作的环节通常集中在第三类，即通过更可被检索、更可信、更结构化的内容供给，提升自身在检索与证据

筛选阶段被纳入的概率。与此同时，也需要认识到长期认知与行为边界的影响：如果品牌长期缺乏权威、稳定的公共信息沉淀，或在某些敏感领域被系统性降权，单次内容优化往往难以带来稳定效果。因此，AIBE 的策略通常应同时覆盖“短期可见”与“长期认知”两条线，并通过持续的证据建设逐步提升被引用的稳定性。

即时证据的准入证：DSS 原则“即时证据”，这是 AI 决定是否采纳内容的核心标准：

- **D-数据支持 (Data Support)**：包含明确的事实依据、数据来源、图表逻辑的内容，在重排阶段会获得更高的“确定性评分”。

- **S-语义深度 (Semantic Depth)**：AI 对空洞的形容词与营销话术敏感度持续下降，内容必须具备结构清晰、观点鲜明、分析深入的特征，而不是简单的关键词堆砌。

- **S-权威来源 (Authoritative Source)**：来自官网、行业白皮书、学术组织或高权重媒体、行业组织的认可的内容，拥有更高的“信任初始分”。在这一维度上，KNIT（可信知识网络）通过权威事实锚定与第三方验证，为品牌内容提供了极高的信任权重。

4.3 工作流程：检索路由、召回、重排与生成

在生成式 AI 的实际工作流中，RAG 机制是决定品牌内容能否进入答案的重要环节。一个相对典型的流程可以拆解为四个阶段。首先是检索路由，即系统判断应当去哪些来源或知识库中寻找材料；这一阶段往往具有“入口效应”，因为系统不会对全网等量检索，而会优先选择其认为更可信、更高效的信源。若品牌内容主要存在于低权重、低可信或不

在系统优先信源范围的平台中，即使内容本身质量较高，也可能在路由阶段就难以进入候选池。

其次是候选召回，即 AI 系统依据关键词与语义相似度等方法，从目标信源中召回一批可能相关的文本片段；需要注意的是，AI 系统召回的往往不是整篇文章，而是被切分后的内容块，因此内容结构是否清晰、段落是否具备独立表达能力，会直接影响召回与后续使用。

第三阶段是重排，即 AI 系统对召回的候选片段进行更精细的打分排序，并选取少量高分片段作为最终证据输入。这一阶段通常是决定性环节，因为生成式 AI 最终会在极少数片段基础上构建答案，未进入 Top-K 的内容在当次回答中基本等同于不存在。重排阶段常见的评估维度包括相关性、信息密度、表达清晰度、来源可信度、事实一致性等。

第四阶段是生成，即大模型基于保留下来的证据片段进行内容整合与表达，形成最终答案。由于生成过程会对不同来源信息进行综合重写，品牌在答案中呈现的方式未必等同于原文表达，因此品牌需要更关注“核心信息是否被采纳、关键结论是否被保留”，而不仅是“是否出现原句或原文链接”。

4.4 为什么品牌内容会被“过滤”

在传统传播语境中，品牌内容往往强调叙事、情绪与调性，这并无问题，但在生成式 AI 的证据筛选中，内容是否“可用”更重要。重排器通常不偏好空泛的形容词堆叠或缺少事实支撑的自我评价，因为这些内容难以作为可靠依据支撑结论，也容易引入不确定性。相反，结构清晰、结论明确、边界清楚、信息密度较高且可以被外部验证的内容，更容易

在重排阶段获得高分，从而进入证据集合。这也解释了为什么在生成式 AI 环境中，许多品牌会出现“明明做了大量内容，但在答案中依然不稳定出现”的情况：问题并不一定在于内容数量不足，而在于内容缺少可被系统采纳的“证据属性”，生成式 AI 所偏好的，正是这种“最容易被纳入证据链”的内容。

去伪存真：AI 在执行检索增强生成（RAG）时，其重排器（Reranker）正变得越来越敏感。如果品牌内容被检测到具有“同质化分发”特征，或引用的信源属于行业“虚假内容黑名单”，系统会采取惩罚性降权。在行业大规模清理垃圾内容、虚假信息以及监管收紧的双重作用下，未来一段时间将出现更大范围的惩罚性降权与“黑名单”机制。

权威性的存证：品牌必须意识到，AI 在生成答案时所抓取的链接，往往是用户进行事实溯源的关键窗口。只有源自真实权威机构、具备可核验链路的内容，才能在“假报告”、“假榜单”泛滥的竞争环境中守住认知主权。缺乏第三方背书或可核验案例的内容，在处理高决策价值（如金融、医疗、法律）问题时，几乎无法进入答案。

口径分裂的风险：如果品牌在不同渠道表达的卖点、数据不一致，AI 会因为无法判断真伪而采取“保守策略”——直接不引用。

因此，生成式内容策略需要从“面向人类读者的表达”适度转向“面向 AI 系统的可用信息供给”。

第五章 全球 AI 品牌资产发展现状与治理框架

5.1 全球发展概况：区域差异与共性特征

全球 AI 品牌资产服务市场正处于爆发前夜，但不同区域市场因技术基础、监管环境与用户习惯的差异，呈现出显著的差异化特征。在北美市场，企业更侧重于底层技术对接、RAG（检索增强生成）优化与 Agent 生态接入。许多领先品牌已开始通过 API 接口，将自身的结构化数据（如产品目录、技术文档、客户案例）直接接入 OpenAI、Anthropic 等大模型的知识库，以提升品牌信息在特定领域问答中的召回率与准确度。这种“数据直连”模式要求品牌具备强大的数据治理与工程化能力。

在欧洲市场，受 2024 年正式通过的《人工智能法案》（AI Act）及配套合规要求的影响，品牌资产建设高度聚焦于数据隐私、透明度与合规性。企业不仅关注 AI 如何“看见”品牌，更关注品牌信息被 AI 使用的过程是否符合 GDPR 等数据保护法规。因此，欧洲品牌在建设 AI 资产时，普遍强调信息来源的可解释性、用户授权的明确性以及算法偏见的防范，将“可信”与“合规”置于首位。

在亚太市场，尤其是中国，企业更看重 AI 认知到商业转化的闭环效率。品牌方积极与本土超级平台合作，探索将 AI 问答直接导流至电商、本地生活服务或企业服务平台的路径。市场增长迅速，根据相关数据显示，2025 年中国生成式内容优化市场规模已增长至约 57 亿元，并预计在 2026 年达到约 137 亿元。这种增长动力主要来自品牌对“品效合一”与“认知—交易闭环”的强烈需求。

尽管存在区域差异，全球市场的共性特征十分明显：品牌方普遍意识到传统 SEO 的局限性，开始将预算向“生成式引擎优化（GEO）”和“AI 语义资产治理”倾斜。全球都在经历一场从“争夺流量曝光”到“争夺 AI 认知主导权”的深刻变革。企业不再满足于“被看见”，而是追求“被理解、被引用、被推荐”。这推动着 AI 品牌资产建设从一个新兴的营销概念，快速演变为企业数字化战略的核心组成部分。

5.2 中国 AI 品牌资产行业现状：从概念热潮走向规则收敛

中国 AI 品牌资产相关服务市场，正处于从“概念导入”迈向“规模化应用”的关键阶段。一方面，生成式人工智能已经深刻改变用户获取信息、比较品牌和形成初步判断的路径，企业对“品牌能否被 AI 正确理解、稳定表达、合理引用”的重视程度快速上升。另一方面，行业整体仍处在方法、口径、评价和边界尚未完全统一的前标准期。也就是说，市场需求正在迅速增长，但行业语言、服务边界与评价体系仍然处于快速形成之中。

这一阶段最突出的特征，不是单一技术概念的流行（如 GEO 概念），而是品牌建设对象发生了迁移。过去企业更关注搜索、媒介和流量入口；而今天，越来越多企业开始意识到，品牌在 AI 场景中的缺席、误读、错引与口径分裂，已经会直接影响用户的候选集合、比较顺序与风险判断。换言之，行业的增长并不只是因为“新概念带来新预算”，而是因为企业首次需要系统面对“AI 如何看见我、理解我、表达我”的问题。

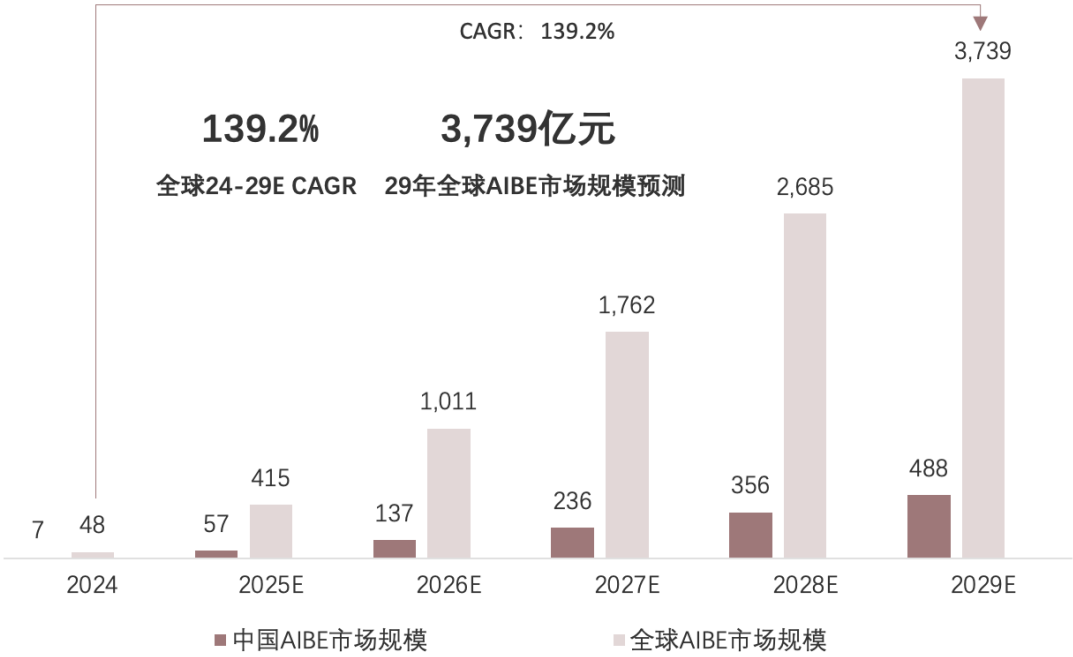
然而，市场的快速升温也催生了诸多乱象，主要表现为虚假信息、数据污染、伪造第三方背书以及缺乏验证机制的短期套利行为。例如，

部分服务商通过海量投放同质化、劣质化的“软文”，试图在 AI 面前堆砌虚假事实，这种行为本质上是“数据污染”，不仅损害了大语言模型的资料索引质量，更由于其“概率性”抓取的特征，导致品牌面临随时被系统过滤或降权的安全隐患。又如，编造假报告、塑造假专家等“虚假权威”行为，在生成式环境中的风险显著高于传统搜索环境，一旦被模型采纳，就会直接进入用户结论，放大误导效应。

随着平台隐性规则的收紧、管理机关的严格管理以及品牌方内部审计的倒逼，中国市场正在从机会叙事转向原则表达，行业规则逐步收敛。平台通过提高权威来源权重、降低低信息密度内容的可见性、强化实体识别与一致性判断等方式，逐步建立“什么内容更可能被采用”的事实标准。管理机关与行业组织则加强了对虚假广告的打击，推动 SoA、引用率、一致性等指标的定义与规范。品牌方则开始要求以固定问题集进行周期性测试、建立对照与留痕机制，并将 AIBV 指标纳入增长与品牌资产管理。这三种力量的叠加，正推动服务生态向认知治理与工程交付分化，清洗短期套利型供给，使行业走向规范化、可持续发展的轨道。

图表 3 AIBE 市场规模预测，2024-2030E

单位：亿元人民币



5.3 全球 AI 治理与品牌可信框架：国际规则与标准演进

在全球范围内，AI 治理正在从伦理倡议走向硬性约束，并深刻影响着品牌资产建设的规则与边界。欧盟的《人工智能法案》确立了基于风险的 AI 监管框架，要求 AI 系统，尤其是生成式 AI，必须具备高度的透明度，包括披露生成内容为 AI 所制、提供训练数据来源摘要等。这意味着品牌向 AI 系统提供知识时，其信息来源的合法性、合规性将受到严格审查，任何基于侵权或非法数据构建的品牌知识资产都可能面临法律风险。

美国近年来通过总统行政令和各州立法，逐步形成以风险管理、防止算法偏见和虚假信息为核心的 AI 治理路径。白宫发布的《人工智能权利法案蓝图》等文件，要求 AI 系统应避免歧视性输出，并确保信息的准

确性与可靠性。这直接要求品牌在向 AI 系统提供知识时，必须确保其内容不包含偏见、歧视或误导性信息，并具备可解释性，以便在出现争议时能够追溯和澄清。

此外，国际标准化组织（ISO）、电气电子工程师学会（IEEE）等机构也在积极制定关于 AI 可信度、数据治理、伦理设计的标准。这些国际规则的演进，共同构成了全球品牌可信框架的基石。其核心共识在于：品牌信息必须具备可追溯的证据链、清晰的风险边界声明以及跨文化/跨语言的伦理合规性。任何试图通过“数据投毒”或“黑帽 GEO”操纵全球大模型结果的行为，不仅面临平台封禁风险，更可能触发跨国法律诉讼与声誉危机。

对于跨国品牌而言，这意味着 AI 品牌资产建设必须纳入全球合规与伦理管理体系。品牌需要建立一套全球统一的 AI 知识治理流程，确保其提供给 AI 的所有信息，无论是在哪个国家、通过哪个平台、被哪个模型使用，都符合当地法律法规与国际伦理标准。这不仅是规避风险的需要，更是建立全球长期信任、巩固品牌权威性的关键。

另一项不容忽视的产业侧变化是：模型厂商正在加速建立“内容供给协议化”体系。自 2024 年以来，OpenAI 先后与 AP（美联社）、News Corp、Axel Springer、Financial Times、Reddit 等签订内容授权协议；Google 则与 Reddit、Stack Overflow 达成数据合作。这意味着 AI 平台正在逐步建立“白名单信源体系”——未在授权协议范围内的内容，将被纳入更严格的可信度评估机制。对品牌的启示在于：长期来看，进入 AI 平台官方信源合作池，可能比单纯优化外部内容更具杠杆效应。这也

正是为什么“权威信源原则”会成为 AI 品牌资产建设中越来越关键的支柱。

5.4 行业典型误区：战略错位、执行失真与评估失焦

企业推进 AI 品牌资产建设时，最容易出现的偏差通常不是能力不够，而是方向先错了。本白皮书将主要误区归纳为战略、执行与评估三类。

第一类是战略错位。表现为只看到短期热度，把 AI 品牌资产建设简单等同于 GEO、搜索优化升级，或者只追求被提及次数，却忽视品牌事实库、可信知识网络和长期口径治理的重要性。例如，一些企业投入大量资源进行“关键词堆砌”式的 AI 内容优化，却未能系统梳理品牌的核心事实、产品边界与权威证据，导致 AI 系统虽然能“提到”品牌，却无法形成稳定、可信的认知。结果往往是投入不少，但所有动作都围绕“快出效果”而不是“形成积累”，一旦模型更新或算法调整，之前的优化成果便荡然无存。

第二类是执行失真。表现为忽视权威信源、可验证案例与结构化信息建设，迷信大量内容输出或碎片化分发。在生成式 AI 环境中，内容数量并不等于影响力。如果品牌内容缺乏结构化表达，或长期存在口径不一致、表述夸张、证据不足等问题，就容易出现误读、错引、过度简化或风险化标签集中的情况。例如，不同渠道对同一产品的性能描述相互矛盾，AI 会因为无法判断真伪而采取“保守策略”——直接不引用。没有边界清晰、证据充分和口径统一的供给，再多动作也可能在重排阶段被过滤掉。

第三类是评估失焦。表现为仍然拿少量截图、单次回答、主观印象去代替可复现评测，或者继续把点击、排名等旧指标当作唯一目标，而忽视品牌在 AI 场景中的被理解、被引用、被误读与被风险化表达等真实问题。例如，仅以“AI 是否提到了品牌”作为成功标准，却忽略了品牌在答案中的角色（是首选还是备选）、引用的权威性、以及不同问题下认知的一致性。AIBV 的意义之一，就在于帮助企业把“模糊感受”变成“结构问题”，把“阶段性展示”变成“持续可测量的资产与风险变化”，从而引导企业将资源投入到真正能提升 AI 品牌资产的建设性工作中。

图表 4 低质 GEO 的局限性与风险



第六章 行业落地指南：不同赛道的“认知主导权”争夺

前面讨论了生成式引擎的筛选逻辑与 AI 品牌资产建设方法体系，但行业落地从来不是“一套模板适配所有赛道”。不同行业的决策链路、用户意图与风险权重差异显著，生成式 AI 在证据选择与表达策略上也存在明显偏好差异。AI 品牌资产建设正在企业服务、零售电商、本地生活、内容 IP 与教育文旅、强合规行业五大赛道率先规模化，预计到 2027 年这五大赛道在中国市场将共同支撑超过 100 亿元的新型营销与认知治理预算。本章围绕五大赛道，拆解典型问题形态、需求痛点与内容建设重心，帮助企业将“认知主权”转化为可执行的资产建设与运营路径。

6.1 企业服务（SaaS/工业/B2B）：从“流量曝光”转向“专业代入”

在企业服务赛道，用户对品牌信息的需求正在从“知道你是谁”转向“你能否给出可落地、可验证、可复用的方案”。当用户在 AI 助手提出“如何缓解制造企业库存积压”“适合 50 人团队的协作软件怎么选”等问题时，系统需要的不只是产品列表，而是一套逻辑闭环、可论证的解决路径。受制于决策成本高、试错代价大的行业特性，生成式 AI 在处理此类问题上往往更依赖结构清晰、可验证、可追溯的“知识材料”，包括技术定义、应用边界条件、分阶段实施路径与具有标杆性意义的案例证据。围绕这一需求特性，企业服务赛道的 AIBE 内容建设需聚焦三类核心资产：一是推进“技术百科化”建设，将零散的产品手册、功能页升级为结构化的技术白皮书与行业词条库，通过标准化呈现提升被引擎检索、引用的确定性，强化专业话语权；二是实现 ROI 语义具象化，摒弃“省钱提效”等口号式的模糊表达，转化为“在某行业场景下，通过

某机制实现多少比例的效率提升/成本改善”，并明确前提条件与核算口径，便于模型提取完整的逻辑链条；三是搭建竞品对比锚点，主动梳理并提供与行业通用方案、主要竞品的边界差异，明确“适用/不适用条件”，为模型输出稳定、可信的推荐结论提供支撑。对于 B2B 行业而言，未来真正重要的，不再只是“是否被搜索到”，而是“是否能够被 AI 视为可信方案的一部分”。

6.2 零售与电商：从“搜索排名”转向“消费意图精准匹配”

零售电商场景中，生成式 AI 正逐步取代传统搜索，成为链接用户与品牌的“全能导购”。用户更习惯于带着明确条件的问题发起决策，例如“预算 500 元送女朋友什么礼物”“敏感肌适合什么防晒”，系统往往直接输出候选清单与推荐理由。传统搜索时代品牌主要依赖关键词与排名抢占，而在生成式环境中，品牌能否进入“备选清单”并以清晰的理由打动客户，成为更关键的增长变量。因此，电商 AIBE 的内容重心应从“讲品牌故事”转向“双向的意图匹配与可对比”。具体可聚焦三类资产建设：一是搭建场景化标签体系（Scenario Tags），围绕“痛点+人群+环境”铺设语义锚点，如“露营必备”“一人食”“极简风”“敏感肌可用”等；二是构建结构化参数库，确保成分、材质、尺寸、适配人群、注意事项等以标准化列表或结构化字段呈现，提升被模型抽取做横向对比的概率，强化答案的说服力；三是强化本地可得性与 LBS 关联，针对“附近哪里能买到”“最快多久送达”等即时性需求，确保门店位置、线上渠道、配送时效等信息实时可检索、可核验，助力品牌进入用户即时购买决策链路，抢占消费场景认知。

6.3 本地生活（餐饮/酒旅/生活服务）：从“高分评价”转向“即时决策入口”。

本地生活的核心竞争更强调“建立信任感”与“提升便利性”两大维度，生成式 AI 的介入更放大了这一核心逻辑。用户常提出具有强场景约束的问题，如“静安区适合商务宴请、环境安静的餐厅”“周末带娃去哪玩、不用排队”，此类问题信息时效性强、状态变化快，用户更愿意接受系统给出可执行的即时建议。生成式 AI 在此更看重可执行性与实时可信信息，而不仅是抽象的高分评价。基于此，本地生活赛道 AIBE 需重点补齐三类内容资产：一是实时动态语料，持续更新营业状态、预约方式、排队情况、服务时段等关键事实，通过高频更新减少模型引用过期信息的风险，筑牢信任基础；二是沉淀深度语义评价，积累包含场景化细节的评论与描述（如“灯光柔和、座位间距大”“亲子区隔音好”），因为细节更能支撑场景匹配推理，提升推荐精准度；三是补全场景化标签属性，将“亲子友好、停车免费、宠物准入、可开发票”等核心特征明确标注为可检索属性，确保用户长尾场景意图能够稳定命中，同时为引擎构建即时消费路由提供清晰指引。相比传统评分体系，未来 AI 更可能优先采纳“是否适合当前场景”的动态判断。

6.4 内容 IP 与教育文旅：从“单向输出”转向“交互式生态搭建”

在知识密集型行业中，AI 更偏好结构完整、步骤清晰且具备可执行性的内容组织方式。面对海量的分散信息，用户不再愿意阅读多篇内容自行拼装路线或学习计划，而希望系统直接输出一份最优的一站式方案，

例如“3 天深度游北京怎么规划”“教师资格证备考路径怎么排”。因此，该赛道的关键不在内容容量的堆砌，而在于内容结构的完整性与可引用性。谁能够率先把知识内容组织成“可被 AI 直接调用的模块化结构”，谁就更容易在生成式环境中形成长期认知优势。建议重点建设三类素材：一是打造结构化教程（Step-by-Step），将攻略、课程、游玩路线等拆解为清晰的步骤、模块与备选项，例如备考攻略按“基础阶段-强化阶段-冲刺阶段”划分，每阶段明确核心任务、资料推荐、时间分配，使模型可以直接引用为答案框架；其二是构建多维评价体系，引入学员反馈、通过率、真实出片心得等第三方验证材料，以提升引用权威性，增加模型引用意愿；其三是维护动态政策库，针对考试报名、景区政策、开放时间等高频变动信息建立实时更新机制，避免模型因不确定而采取保守表达或降低引用，确保认知输出的时效性与准确性。

6.5 强合规行业（医疗/金融/法律）：从“模糊回答”转向“权威合规源”

强合规行业是 AIBE 难度最高但价值也最显著的赛道之一。由于存在误导风险，生成式 AI 在医疗、金融、法律等问题上通常采用更保守的策略，优先引用权威、可核验、具备明确专业或法律依据的信源材料。用户核心需求多为专业问题咨询，例如“某条款如何解读”“某疾病如何预防”“某产品风险点是什么”等，若 AI 系统缺少可信依据往往会弱化结论甚至回避推荐，导致品牌错失认知占位机会。因此，该赛道 AIBE 的

核心任务是“占位权威信源并降低风险不确定性”。具体需重点建设三类资产：一是构建经合规审核通过的官方标准问答（FAQ）与权威认证语料库，将高频专业问题、标准解答、合规表述整理成册，作为模型检索的默认优先信源；并逐步围绕来源真实性、表达规范、证据可追溯性与更新机制，探索建立可认证的权威语料单元体系。二是强化专家语义背书与标准引用，在内容中合理融入权威专家观点、行业核心标准、公开学术文献及法律法规条文，并附上可验证出处，提升内容权威性与可追溯性；三是明确边界说明和风险提示，清晰界定内容适用范围、禁忌条件及责任边界，例如医疗内容注明“仅供健康参考，不构成诊疗建议”，金融内容明确“不构成投资建议”，同时标注线下专业咨询渠道。在生成式环境中，这类能够清晰界定风险、划清服务边界的内容，更易被系统采纳为专业证据，从而稳固权威认知地位。

第七章 可信知识网络建设方法论：从认知基础设施到 GEO 运行机制

7.1 从内容供给到认知基础设施：为什么企业需要可信知识网络

在生成式 AI 逐步成为信息入口的背景下，品牌面临的核心任务，不再只是增加内容数量，而是建设一套能够稳定提供真实、清晰、可验证、可复用知识的认知基础设施。只有当品牌信息具备明确结构、可靠证据与一致口径，才更可能进入生成式系统的检索、召回、重排与生成流程，并在关键问题中成为答案的组成材料。

7.2 可信知识网络的定义：把分散信息升级为可被 AI 采用的知识体系

可信知识网络，是面向 AI 时代的企业级认知基础设施，即将品牌从“内容生产者”转变为“可信知识供给者”。它以真实世界数据、权威研究结论、结构化图谱、第三方验证结果与可追溯信息来源为核心构成要素，通过 AI 认知基线诊断、真实世界验证、权威事实锚定、知识结构化工程、可信内容扩散、持续监测与认知巩固六层结构化工程，将企业分散的非结构化信息转化为可被 AI 稳定理解、引用与复用的标准化知识体系。

可信知识网络的目标，不是单纯提升提及次数，也不是追求短期占位，而是帮助企业掌握 AI 时代的认知主导权，确保自身被市场正确理解，让 AI 能够基于准确前提稳定引用企业核心价值与关键数据，并规避数据投毒、算法幻觉与认知稀释等系统性风险。换言之，可信知识网络的本质，是把企业分散的信息资产升级为可长期调用、可持续积累的可信数字资产。

7.3 三层架构：可信知识网络的整体结构

从整体上看，可信知识网络可以概括为一个由下至上的三层架构。第一层是**可信信源层**，由真实世界数据、权威研究、第三方验证结果、结构化图谱与可追溯来源构成，是整个体系的事实底座，解决“知识从哪里来”的问题。第二层是**知识工程层**，负责将企业分散、异构、冗长的信息转化为可被系统调用的知识单元，也是在这一层，**GEO** 的方法体系真正发挥作用。第三层是**认知输出层**，对应品牌在生成式环境中的最终表现，即是否被系统稳定识别、是否进入关键问题答案、是否被作为可信证据引用，并进一步沉淀为长期 **AI 品牌资产**。

需要特别强调的是，**GEO** 并不独立于可信知识网络之外，而是可信知识网络在知识工程层与可信扩散层中的核心运行机制。**GEO** 的对象不是“多写内容”，而是围绕语义资产、答案资产与引用资产，构建更易被系统采纳的知识供给体系。换言之，**GEO** 的真正作用，是把可信知识网络中的事实、证据与知识，转化为 **AI** 可理解、可调用、可引用的答案材料。

7.4 六层建设路径：从认知诊断到认知巩固

可信知识网络的建设应该遵循完整链路的逻辑。

第一层是 **AI 认知基线诊断**，用于识别系统当前如何理解品牌，判断品牌是否被正确识别、是否与竞品混淆、在哪些高价值问题中缺席。第二层是**真实世界验证**，即将企业能力、案例、效果与结果建立在真实业务证据之上，使品牌表达从“自我陈述”转向“事实表达”。第三层是**权威事实锚定**，即把关键事实锚定在更高公信力的来源之上，如研究报

告、标准文档、权威媒体、第三方评测与合规材料，以提升系统对品牌知识的信任初始分。第四层是**知识结构化工程**，即把零散信息加工成清晰、独立、可被调用与引用的知识单元。第五层是**可信内容扩散**，即让已结构化、已验证的知识进入系统优先采用的信源与平台知识生态，提升被路由、召回与采用的概率。第六层是**持续监测与认知巩固**，即通过固定问题集、答案份额、引用率与一致性等指标对品牌表现进行长期追踪，并持续修正知识口径与证据链，确保认知优势稳定沉淀。

从方法论上看，这六层并非彼此割裂，而是一个从事实到认知、从建设到治理的连续过程：先诊断，再验证；先锚定，再结构化；再扩散，最后巩固。这意味着，AI 品牌资产建设将是一种持续运行的认知治理机制。

7.5 权威高质量语料库：可信知识网络的知识底座

在企业推进 AI 品牌资产建设的过程中，真正决定品牌能否被系统稳定理解、审慎引用和长期复用的，往往不是内容数量，而是是否建立起一套权威、高质量、可持续更新的知识供给底座。为此，本报告建议将“权威高质量语料库”作为可信知识网络建设中的关键工程加以强调。

这里所说的权威高质量语料库，并不是资料的简单堆积，也不是把现有稿件、页面与宣传内容机械汇总，而是一套围绕品牌、产品、服务、场景、案例、边界与合规信息建立起来的高质量知识供给系统。它的核心作用，不在于制造更多信息，而在于为 AI 系统提供更低歧义、更强证据性、更高可信度的知识材料，使品牌在检索、召回、重排与生成阶段更容易进入候选集合并被作为依据使用。

从构成上看，权威高质量语料库至少应包括六类内容：

第一，品牌事实库，用于明确品牌实体、核心能力、产品边界、适用场景与基础属性；

第二，标准问答与标准语料，用于围绕高频问题形成统一、清晰、可复用的标准表达；

第三，答案资产，用于将关键结论组织为可独立调用的知识单元；

第四，引用资产，用于提供研究报告、第三方评测、标准 FAQ、案例数据、认证说明与合规材料等可验证依据；

第五，案例与证据材料，用于支撑品牌价值主张从“自我陈述”转向“事实表达”；

第六，边界说明与合规信息，用于明确适用条件、风险提示与责任边界，尤其在高风险行业中构成系统采纳的重要前提。

权威高质量语料库的价值，在于把品牌原本分散、异构、冗长的信息，转化为 AI 更容易消费的“优先信源”。它既是可信知识网络中“可信信源层”的具体落地，也是后续问题集建设、答案资产建设、引用资产建设和持续监测的内容基础。没有高质量语料库，品牌的知识建设就容易停留在零散表达层面；而一旦建立起这套底座，品牌就更有可能在生成式环境中形成稳定、清晰、可信的认知位置。

需要进一步强调的是，随着行业进入规范化阶段，权威高质量语料不应只停留在企业自建层面，而需要逐步引入可核验、可分级、可认证的机制。也就是说，未来行业不仅需要高质量语料，更需要围绕语料来源真实性、结构完整性、证据可追溯性、边界清晰度、更新机制与合规

审查情况，形成统一要求，并在此基础上逐步建立语料认证、内容认证与服务认证机制。这样做的目的，不是增加市场包装，而是为 AI 系统提供更可信的优先信源，为品牌提供更稳定的被采纳基础，也为行业治理提供更明确的标准接口。

从这个意义上讲，权威认证语料库体系既是可信知识网络建设中的底层工程，也是未来行业规范、评估体系与服务生态走向标准化的重要支撑。围绕这一体系，未来可进一步形成经认证的标准问答库、标准语料单元、权威引用材料库，以及配套的审核规则、实施细则和服务分级机制，从而推动 AI 品牌资产建设从“方法实践”走向“标准建设”。

特别是在医疗、金融、法律等强合规场景中，权威高质量语料库的重要性会进一步放大。AI 系统在处理高风险问题时，更倾向于采用经过合规审核的标准问答、权威来源与可核验材料。因此，企业应优先整理高频专业问题、标准解答、合规表述、专家观点、标准条文与风险提示，将其构造成模型更容易优先采用的默认信源体系。

从这个意义上讲，权威高质量语料库不是内容运营的附加项，而是 AI 品牌资产建设中的底层工程。它决定了品牌向 AI 提供的，究竟是一堆分散文本，还是一套可以被理解、验证、调用和复用的可信知识系统。

7.6 从语义定位到答案资产：可信知识网络的操作方法

如果说可信知识网络回答的是“为什么建、总体怎么建”，那么 GEO 回答的就是“具体怎么做”。在操作层面，

首先是**品牌语义定位**。品牌想在关键问题中获得稳定出现，首先要解决“被正确理解”的问题。语义定位要回答三件事：AI 系统是否认得

出你是谁，是否分得清你与竞品的差别，是否在特定场景下会自然想到你。为实现这一目标，企业需要明确品牌的实体身份、能力边界、语义邻接与场景绑定，使 AI 系统能够在正确语境中识别与召回品牌。

其次是**答案资产建设**。在生成式 AI 的检索与重排机制下，AI 系统使用的往往不是整篇文章，而是被切分后的内容块。因此，品牌需要将关键知识与关键结论组织成可独立使用的“答案单元”，使其具备明确问题边界、清晰结论、结构化理由、适用限制与可追溯依据。答案资产是可信知识网络中最小、最直接、最可调用的知识单元。

第三是**问题集建设**。生成式时代，品牌覆盖的对象已不再是关键词清单，而是用户真实提出的问题场景。因此，企业需要围绕定义类、对比类、决策类、风险类与替代类问题，建立覆盖用户全旅程的高价值问题集，并围绕这些问题持续供给答案资产。问题集决定了知识建设的优先级，也决定了品牌最终会在哪些问题中进入答案。

第四是**引用资产建设**。仅有答案资产并不足以形成长期、稳定的被引用能力，企业还需要同步建设可作为权威依据的引用资产，包括白皮书、研究报告、第三方评测、权威媒体、标准 FAQ、案例数据、认证说明与合规材料等。引用资产的作用，是为答案提供出处与依据，让系统“更敢用”。

从这个角度看，可信知识网络中的“知识结构化工程”，在操作层面就表现为：以语义定位校准品牌认知，以问题集确定场景边界，以答案资产提供最小知识单元，以引用资产提供可信证据支撑。

7.7 从知识建设走向认知治理

在 AI 答案时代，品牌建设的重点，已不再只是内容分发，而是知识供给。可信知识网络作为一套面向 AI 时代的企业级认知基础设施，解决的是品牌如何把分散信息转化为可被系统稳定理解、引用与复用的知识体系；而 GEO 则是这套基础设施在知识工程层与可信扩散层中的核心运行机制，解决的是品牌如何通过语义定位、问题集、答案资产、引用资产与 CREATE 运行框架，把知识真正组织成答案材料并进入系统候选集合。

当可信知识网络完成基础建设，并通过 GEO 运行机制转化为稳定的答案供给能力之后，品牌管理的重点便不再只是“是否做了建设”，而进一步转向“这些建设是否真正影响了答案生成”。因此，下一章将围绕答案份额、引用率、认知一致性与情感倾向等指标，讨论企业如何将可信知识网络与 GEO 的建设成果转化为可衡量、可复盘、可管理的 AI 品牌资产表现体系。

第八章 新评估体系：从点击指标到认知治理指标

8.1 为什么需要重建指标

在传统数字营销体系中，曝光、点击、转化、排名等指标之所以在过去长期有效，根本原因在于用户获取信息的路径高度依赖“跳转”：用户看到结果后需要点开链接，进入页面。再进一步了解并完成决策。对品牌而言，能否获得点击既代表被看见，也代表进入了用户的考虑范围，因此点击率、访问量、转化漏斗可以较为直接地反映营销动作的效果。

生成式 AI 的普及改变了这一前提。越来越多用户在对话界面中直接获得结论与建议，很多比较与筛选在“点击发生之前”就已经完成，甚至在某些场景中不再需要点击。这并不意味着品牌不再需要转化，而是意味着品牌影响用户决策的关键节点，正在从“页面访问阶段”前移到“答案生成阶段”。因此，继续用点击、排名来衡量生成式场景下的品牌影响力，容易出现一种错位：数据看起来仍然可观，但品牌在关键问题上的推荐顺位与被纳入候选集合的概率可能已经发生变化，而这种变化往往不容易通过传统漏斗及时捕捉。

在“答案时代”，衡量体系需要回答一个更直接的问题：当用户把初筛与判断交给生成式 AI 时，品牌在 AI 系统的答案中处于什么位置，是否被采纳、以何种方式被采纳，以及这种采纳是否具有稳定性与可复用性。只有当指标体系能够反映这一层面的变化，GEO 的投入产出才可能被管理、被复盘并形成长期积累。

8.2 核心衡量对象：品牌“影响答案”的能力

在生成式 AI 中，品牌影响力的核心不再仅是“带来多少访问”，而更接近“影响了多少结论”。通俗地说，品牌要衡量的不仅是用户是否来到品牌页面，而是用户在提出关键问题时，AI 系统是否把品牌纳入答案、是否将其视为较优选择、是否给出明确理由，以及是否引用了可验证的证据来源。这类衡量更贴近生成式环境的真实决策路径，也更能反映品牌在 AI 入口中的可见性与竞争地位。

基于这一思路，本章提出四类关键指标：答案份额（Share of Answer）、AI 引用率（AI Citation Rate）、认知一致性（Cognitive Consistency）与情感倾向（Sentiment Orientation）。这四类指标分别对应“是否被纳入答案”“是否被当作可信依据”“AI 系统对品牌的判断是否稳定”“AI 系统对品牌的总体评价是否存在风险化倾向”，共同构成 AI 时代较为完整的衡量框架，并可进一步汇聚为 AIBE（AI 品牌资产）的量化管理视角。

8.3 答案份额（SoA）指标：在关键问题中“占了多少答案位置”

答案份额（Share of Answer，简称 SoA）可以理解为生成式环境中的“品牌市占率”或“心智份额”的对应物，它衡量的是：在一组事先定义的高价值问题集中，品牌在生成式答案中出现的比例及其综合权重。与传统“被提及次数”不同，SoA 不是简单统计是否出现，而是试图刻画品牌在答案中的实际地位，因为在生成式回答中，“出现”本身并不等同于“被推荐”。

从可操作角度看，SoA 通常至少包含三个层面的权重考虑。第一是提及率，即在多少问题中品牌被纳入答案，反映品牌是否进入候选集合

的基本概率。第二是角色权重，即品牌在答案中被呈现为“首要推荐”“备选方案”还是“风险提醒或反例”，不同角色对决策影响差异显著。第三是顺位与显著性，即当答案以列表或分层结构呈现时，品牌出现的位置是否靠前，理由是否明确、篇幅是否充分，反映系统对品牌的确定性程度。通过对问题集进行长期、周期性的测试，SoA 可以用于观察品牌在生成式入口中的总体占位变化，并与竞品进行横向比较，为品牌管理层提供更接近真实决策影响力的量化依据。

8.4 引用率指标：品牌是否被当作“可信证据”使用

在生成式环境中，“被引用”往往比“被提及”更接近系统真实的信任结构。如果说 SoA 主要衡量品牌在答案中的“占位”，那么 AI 引用率更直接衡量品牌在答案中的“可信度”。在生成式 AI 的工作机制中，尤其当问题具有较高决策价值或存在风险属性时，AI 系统往往更依赖可验证来源来支撑结论。此时，品牌相关内容能否成为被引用的证据来源，将显著影响品牌在不同问题中的可复用性与长期稳定出现概率。

AI 引用率可以从两个维度来理解：一是引用频次，即品牌自有资产或与品牌相关的权威材料在生成式答案的引用链中出现的次数与比例；二是引用质量，即被引用信源的权威程度与可验证性，例如来自行业协会、权威媒体、研究机构、标准组织或高可信平台的材料，通常比低质量转载或软文站具有更高权重。需要强调的是，“被提及”往往只是结果呈现，而“被引用”更接近 AI 系统形成结论的依据来源。品牌若能够持续提升引用资产的质量与覆盖面，通常更容易在长期形成稳定的答案占位，并在高价值问题上获得更强的推荐确定性。

8.5 认知一致性指标：AI 系统对品牌的判断是否稳定可复现

在生成式环境中，很多品牌会遇到一种常见现象：同一品牌在不同问题、不同问法下系统给出的解释与评价差异较大，有时被定位为“专业方案”，有时被描述为“成本较高选项”，甚至在不同场景中被归入不同品类或被关联到不一致的优势点。这类现象的根源往往不是“模型随机”，而更可能来自品牌信息供给的分散与不一致，例如不同渠道表达口径冲突、关键卖点反复变化、边界描述模糊，导致 AI 系统在检索与重排环节抽取到不同证据片段，从而输出不同结论。

认知一致性（Cognitive Consistency）用于衡量 AI 系统对品牌的核心判断是否稳定，即在相同或相近问题下，品牌被推荐或被描述的理由是否一致、关键标签是否稳定、适用边界是否清晰可复现。对生成式 AI 而言，一致性意味着更低的不确定性与更可控的风险，因此一致性提升往往与长期可引用能力正相关。对品牌管理而言，认知一致性也是一个重要的治理指标，它提醒品牌必须建立统一、稳定、可验证的知识口径与证据体系，避免“内容很多但口径分裂”，最终导致系统难以形成确定判断。

进一步地，在长上下文与多轮对话场景下，传统“一次性问答”式评测的代表性正在下降。新的趋势是“Token 级一致性”——即同一条品牌事实，在不同上下文长度、不同对话深度下，是否被 AI 表达得一致。建议未来 AIBV 版本可引入 CDC（Conversation Depth Consistency，对话深度一致性）作为补充指标，以更精准衡量 Agent 时代品牌认知的稳定性。

8.6 情感倾向指标：生成式环境下负面影响更集中

在传统搜索环境中，负面信息通常以分散形式存在，用户是否点击、是否相信、是否继续阅读仍具有较强主观选择空间；而在生成式环境中，系统可能会将分散的评价与反馈整合为一句或几句“总结性判断”，并以较为客观的语气呈现。这种总结一旦形成，往往比单条差评更具影响力，因为它直接进入答案，成为用户的第一印象与决策依据之一。

因此，情感倾向指标不仅关注正负面比例，更重要的是关注系统是否在高价值问题中形成了稳定的风险提示、争议标签或谨慎性表述，以及这些表述出现的频次、强度与触发场景。品牌需要警惕的是，当 AI 系统反复在某些问题中给出相似的负面总结时，问题往往不仅是某一条内容的舆情，而可能是品牌在可验证证据与边界表达上的缺口，导致系统只能在不确定性下采取更保守的表述方式。对于这类风险化倾向，治理重点通常应回到证据链建设与口径一致性修复，而非仅做表面“压负面”。

8.7 AIBE：将指标收敛为“可管理的 AI 品牌资产”

当 KNIT、SoA、引用率、一致性与情感倾向等指标形成体系化监测后，品牌即可获得一幅更贴近生成式环境的资产画像：品牌在关键问题中是否进入候选集合，是否能以可信理由被推荐，AI 系统认知是否稳定，以及是否存在风险化总结的长期趋势。基于此，本报告提出以 AIBE（AI 品牌资产）作为上层管理视角，将分散指标收敛为可长期追踪的资产变化，并据此评估生成式内容优化的阶段性效果与长期积累。

需要说明的是，AIBE 的价值不在于提供单一“评分”，而在于帮助

品牌建立一套可复盘的治理体系：通过问题集定义与周期性测试管理 SoA，通过引用资产建设提升引用率，通过口径治理提升一致性，通过证据补强与风险说明修复情感倾向，从而在生成式入口中形成更稳定、更可持续的竞争优势。

8.8 AIBV 指数体系：面向生成式内容的统一评估框架

为使本报告提出的“答案时代指标”能够在不同模型、不同入口与不同时间窗口下被一致理解与复现测量，本报告设计了《AI 品牌资产发展指数体系》（AIBV）作为统一评估框架，为行业建立统一、可复现、可比较的认知治理语言。

正文关键指标	在 AIBV 中的归属
SoA	AIP-B 可见性与召回度
引用率	AIP-A3 引用/证据可追溯度
认知一致性	AIP-D 一致性与调性匹配
情感倾向	AIR-R2 负面偏差控制度
问题集与测量质量	MCI/UAF/数据规范

本章提出的 SoA、引用率、认知一致性与情感倾向，是对 AIBV 体系的解释性核心指标；在正式评估框架中，这些指标将进一步映射为 AIP、AIR 及相关校准因子中的具体维度，以实现从“管理语言”到“量化框架”的过渡。

在 AIBV 1.0 体系中，可信知识网络（KNIT）的应用显著提升了“方法置信度（MCI）”与“证据可追溯度（A3）”。通过 KNIT 的真实世界验证与权威事实锚定，品牌能够有效识别并剔除那些通过“GEO

工具一键分发”产生的同质化网页，确保品牌的 SoA（答案份额）是建立在真实公信力之上。

第九章 行业规范与治理框架：从原则倡议到实施细则

9.1 为什么行业规范正在成为 AI 品牌资产建设的必要组成

随着生成式人工智能逐步成为信息获取、品牌比较与决策形成的重要入口，AI 品牌资产建设已经不再只是一个增长议题，也不再只是一个方法创新议题，而开始成为一个规范议题与治理议题。原因在于，品牌在 AI 语义空间中的存在，不仅会影响传播效率，更会影响用户的初筛结论、比较逻辑与风险判断。如果缺乏清晰的行业边界与共同规则，这一领域就极易被误解为对模型偏好的操纵，或在执行中滑向黑帽化、投毒化、伪权威化的短期套利路径。

本白皮书前文已经指出，当前行业正处于高增长与前标准期并存的阶段：一方面，企业预算与市场热度迅速上升；另一方面，方法边界、评价语言与服务标准尚未完全统一，市场中仍存在虚假信息、数据污染、不可验证承诺与信源欺诈等乱象。也正因如此，行业的下一步成熟，不应仅靠规模扩张推动，更需要通过规范建设来明确“什么是正当建设、什么是风险行为、什么是可持续的方法路径”。

从这个意义上讲，行业规范不是附着在方法论之外的补充条款，而是 AI 品牌资产建设的组成部分。可信知识网络（KNIT）解决的是如何把品牌信息转化为可被 AI 理解、引用与复用的知识体系；GEO 解决的是这些知识如何被组织、部署并进入系统候选集合；AIBV 解决的是如何对品牌在 AI 中的表现、建设能力与风险状态进行统一评估；而行业规范解决的，则是这一整套体系在现实应用中应当遵循什么原则、守住什么边

界、接受什么约束。只有当建设、评估与规范三者形成闭环，AI 品牌资产建设才可能从“热点服务”真正走向“基础设施能力”。

9.2 行业基本原则：从“行业倡议”升级为“共同规范”

本白皮书认为，面向 AI 时代的品牌建设，应至少共同遵循以下五项基本原则。

第一，真实性原则。品牌不得以虚假信息、伪造案例、伪造背书、伪造第三方评价或伪造研究来源等方式试图影响 AI 的认知与表达。无论是品牌自有内容、第三方传播内容，还是供 AI 采纳的证据材料，都应以真实事实为前提。真实性是 AI 品牌资产建设的底线，也是品牌进入 AI 候选集合的基本前提。

第二，透明性原则。品牌官方信息、第三方研究、商业传播与 AI 生成内容之间，应保持适当区分；涉及重要结论的内容，应尽可能具备清晰来源、明确身份与可追溯链路。透明性不仅是合规要求，也是在生成式环境下建立长期可信度的重要条件。

第三，可验证原则。品牌进入 AI 语义空间的关键内容，必须具备事实锚点、证据链与可核验出处，避免“只有结论、没有依据”。这意味着品牌建设的重点，不应只是强调优势叙事，而应同步建设答案资产与引用资产，确保 AI 在高价值问题中有材料可用、有出处可引。

第四，边界清晰原则。AI 品牌资产建设的目标，不应被定义为操纵模型偏好、制造结果倾向或争夺算法漏洞，而应被定义为帮助 AI 更准确地理解、引用和表达品牌。换言之，行业应当鼓励的是“降低理解成本、

提升证据质量、增强语义清晰度”的建设性方法，而不是“污染信息环境、伪造信源、诱导系统输出”的破坏性方法。

第五，治理导向原则。评估体系不只服务于增长，也要服务于误读识别、错引纠偏、风险表达治理和消费者决策安全。AI 品牌资产建设不是一次性“做出结果”，而是一项持续的治理工程，其目标应是帮助企业在不同模型、不同时间与不同问题场景下保持更高的一致性、可信度与长期稳定性。

这五项原则，实际上构成了本白皮书关于行业规范的最基础表达。它们将前文关于 KNIT 的知识治理逻辑、关于 GEO 的边界重构逻辑，以及 AIBV 关于可解释性、可复现性、防操纵性与合规性的要求，统一收束为一套可被行业共同采纳的原则语言。

9.3 禁止性行为清单：什么不应再被视为“有效方法”

在前标准期，行业最大的问题之一，不是方法太少，而是边界不清。部分市场行为以“提升 AI 推荐概率”“快速形成占位”为名，实际上采取的是伪造信源、污染资料、制造虚假权威与结果操纵等高风险方式。对此，本白皮书认为，以下行为不应再被视为 AI 品牌资产建设的有效方法，而应被明确归入行业治理对象。

第一类，伪造第三方权威与证据链。包括编造研究报告、塑造“假专家”、虚构行业认证、伪造媒体或机构背书等。这类行为在传统搜索环境中或许还能以“营销包装”存在，但在生成式环境中，其风险显著更高，因为一旦被模型采纳，就会直接进入用户结论，放大误导效应。

第二类，低质量铺量与数据污染。包括大规模分发同质化软文、站群堆砌伪权威内容、通过海量雷同资料试图向 AI 注入虚假事实等。这种做法本质上是对信息环境的污染，不仅损害大模型的资料质量，也会因系统重排机制升级而带来整体降权、过滤甚至黑名单风险。

第三类，结果操纵与不透明承诺。包括“包推荐”“包引用”“快速起量”“形成占位”等无法清晰解释、无法持续验证、无法复现的方法承诺。对于品牌而言，这类承诺往往会将建设工作导向短期动作，而非长期能力；对于行业而言，则会进一步放大黑箱交付与概念泡沫。

第四类，缺乏边界说明的误导性表达。在医疗、金融、法律等高风险行业，如果品牌没有明确适用条件、禁忌边界和责任说明，却试图借助 AI 形成明确推荐或强结论，同样属于高风险行为。行业规范的要求，不是让品牌“少说”，而是要求品牌“把话说清楚、把边界讲明白、把依据提供足”。

因此，本白皮书建议，行业应当明确区分“建设性方法”与“破坏性方法”：凡是有助于提升品牌信息真实性、结构化程度、可验证性、可引用性和认知一致性的建设性动作，都可以被纳入 KNIT/GEO 方法体系；凡是依赖污染信息环境、伪造信源、制造偏好操纵或作出误导性承诺的行为，都应被明确排除在 AI 品牌资产建设的正当方法之外。

9.4 推荐性建设规范：什么样的建设才是可持续的

如果说禁止性行为清单回答的是“什么不能做”，那么推荐性建设规范回答的就是“什么值得长期做”。本白皮书认为，面向 AI 时代的品牌建设，应当把工作重点放在以下五类长期能力上。

第一，建立品牌事实库与可信知识底座。品牌需要系统梳理核心事实、业务边界、产品能力、适用场景、价格带、关键限制条件与风险说明，并持续更新。这是可信知识网络建设的起点，也是 AI 形成稳定理解的基础。

第二，建立问题集与答案资产。品牌不应再只围绕关键词堆砌内容，而应围绕用户真实问题场景建设答案单元。问题集决定优先级，答案资产决定可调用性。凡是能够明确问题边界、给出清晰结论、附带结构化理由和适用限制的内容，才更容易被系统采纳。

第三，建立引用资产与权威锚点。品牌应持续建设白皮书、研究报告、第三方测评、标准 FAQ、案例数据、合规材料、权威媒体报道等引用型资产。引用资产不仅服务于传播，更服务于 AI 的证据筛选与结论生成，是品牌被“更敢引用”的关键。

第四，建立口径治理与风险修复机制。品牌需要持续监测 AI 对自身的误读、错引、过时信息、风险化标签与口径分裂，并通过事实修正、边界补充、引用补强与内容更新进行修复。治理能力本身，也是 AIBE 五层路径中的关键成熟度指标。

第五，建立跨部门协同的长期运行机制。

随着 AI 品牌资产建设进一步基础设施化，这项工作将不再只是内容团队或营销团队的单点任务，而会越来越需要品牌、内容、产品、法务合规与数据团队协同参与。特别是在高风险行业与高价值问题场景中，事实更新、边界表达与合规审查必须形成制度化协作。

从方法论上看，这些推荐性建设规范与第六章 KNIT—GEO 框架是完全一致的：其本质不是制造更多传播噪音，而是以更真实、更可验证、更低歧义的方式组织品牌知识，提升 AI 对品牌的理解效率与引用确定性。

9.5 评估规范：AIBV 如何成为行业共同语言

行业规范不能只停留在原则层面，还必须进入评估层面。没有统一的评价语言，行业就很难区分“看起来有效”和“真正有效”，也难以把治理工作从主观感受转化为可复盘、可追踪、可比较的结果。正因如此，AIBV 的意义，不只是提出一套指数体系，更在于为行业提供一种共同的评估语言。

AIBV 体系已经明确提出：其构建与使用应遵循客观性、可复现性、可解释性、防操纵性、合规性与实用性原则。这些原则本身，就构成了行业评估规范的核心。也就是说，评估体系不能只是“算分工具”，而必须同时回答以下问题：测量对象是否清晰，问题集是否合理，数据采集是否留痕，标注是否有统一规范，结果是否可以被第三方复核，是否具备防刷榜与防操纵机制，是否符合数据安全、个人信息保护与广告监管要求。

在具体层面，本白皮建议将以下要求视为行业评估的基本规范：

第一，**多模型多入口原则**。同一问题集应在多模型、多入口下进行测量，避免单一模型或单一入口偏差造成结果失真。

第二，**问题集分层原则**。题集至少覆盖认知、对比、决策、行动、风险等五类用户意图，并保证强约束问题达到规定阈值。

第三，**留痕与可审计原则**。应完整保留问题文本、完整回答、时间戳、模型版本、入口信息、结构化标注结果及关键计算过程。

第四，**标注与质量控制原则**。应建立统一标注规范，采用多人复核、抽样复查、高波动题单独识别与噪声统一清洗等机制。

第五，**结果发布规范**。对外发布结果时，应同步披露指数名称、版本号、测量窗口、发布机构与方法置信提示，不得以绝对化表述替代边界说明。

从这个角度看，AIBV 已经不只是“评估品牌在 AI 中表现”的技术框架，也是一套推动行业走向标准化、去黑箱化、去操纵化的规范工具。它把“什么叫好表现、什么叫好建设、什么叫高风险”从口号变成了可计算、可审查、可校准的语言体系。

9.6 认证机制：从语料审核到服务分级

在行业规范逐步走向制度化的过程中，仅有原则表达和禁止性清单仍然不够。若要让行业真正形成可执行、可复核、可推广的共同标准，还需要进一步建立从语料、内容到服务的认证机制，使“什么样的知识可被优先信任、什么样的方法值得被采纳、什么样的交付可以被视为合规有效”具备更明确的判断依据。

未来行业可逐步围绕三层认证机制展开探索。

第一层，语料认证。语料认证面向的是进入 AI 语义空间的基础知识单元，重点审核其来源真实性、事实准确性、结构完整性、证据可追溯性、更新时间与边界说明情况。其目的在于识别和筛选真正具备高可信度的标准问答、标准语料、案例证据、引用材料与合规说明，使其成为

模型更容易优先采用的知识底座。对于高风险行业而言，语料认证尤其重要，因为系统在相关问题中往往更依赖经过审核、可核验、边界明确的专业材料。

第二层，内容认证。内容认证面向的是围绕问题集、答案资产、引用资产、FAQ、白皮书、案例页等形成的品牌知识产品。其重点不只是看内容“写得多不多”，而是看其是否符合真实性、透明性、可验证性、边界清晰与治理导向等行业基本原则，是否能够在不同模型、不同场景下保持较高的一致性与可解释性。内容认证的作用，在于推动行业从“内容堆积”走向“知识治理”，使品牌输出真正成为可被系统采用的标准化知识供给。

第三层，服务认证。服务认证面向的是提供 AI 品牌资产建设、可信知识网络建设、GEO 运行优化与评估治理服务的机构。其核心不在于认证“谁更会包装结果”，而在于识别其方法是否透明、过程是否留痕、交付是否可审计、数据是否可复现、结果是否具备边界说明，以及是否能够有效区分建设性方法与操纵性、污染性方法。未来，围绕 AI 品牌资产建设服务机构建立分级认证与规范评价机制，有望成为行业标准化的重要组成部分。

需要强调的是，认证机制的目的不是增加额外门槛，而是把原本分散、模糊、不可比的建设实践，逐步转化为可核验、可分级、可复用的标准对象。语料认证解决“底层知识是否可信”的问题，内容认证解决“知识产品是否合规、可用”的问题，服务认证解决“方法与交付是否可靠”的问题。三者共同构成从原则倡议走向实施细则的重要桥梁。

从行业推进路径看，认证机制还可与 AIBV 的实施细则进一步结合：一方面，语料与内容认证可提升评估中的证据可追溯度、方法置信度与结果可解释性；另一方面，服务认证可为品牌方选择合作对象提供更清晰的评价依据，推动行业逐步形成“可解释、可验证、可审计”的合作标准。也正因如此，认证机制并不是独立于评估体系之外的附加安排，而是行业规范、可信知识网络与 AIBV 评估体系走向统一的重要接口。

从这个意义上讲，行业认证机制不是对市场的额外约束，而是推动行业从“概念竞争”走向“标准竞争”、从“方法表达”走向“制度建设”的关键一步。只有当可信语料、规范内容和合格服务都逐步进入可认证、可治理的轨道，AI 品牌资产建设才可能真正从机会型市场成长为基础设施型行业。

9.7 标准化推进路径：从原则表达走向行业实施细则

本白皮书认为，行业规范的最终目标，不是让行业停留在“倡议式自律”阶段，而是推动其逐步走向“标准化实施”阶段。结合 AIBV 附件中已提出的版本管理与实施文件建议，未来行业至少需要围绕以下五类文件持续完善。

第一，**行业实施细则**，用于明确基础原则在不同场景中的适用边界与操作要求。

第二，**题库与问题集规范**，用于统一问题集分层标准、题源构成与高价值问题定义。

第三，**标注与评分手册**，用于保证不同机构、不同项目之间的标注口径与评分逻辑一致。

第四，**质量控制与申诉流程**，用于规范异常检测、争议处理、事实修正与复核机制。

第五，**指标校准与实证报告**，用于伴随版本迭代不断优化权重、修正口径并提升体系的行业适配度。在这一推进路径中，行业规范不是与市场发展对立的限制，而恰恰是支撑市场长期扩容的条件。没有规范，行业会被短期套利拖累；没有治理，方法论会被黑帽化误解；没有实施细则，评估体系就难以成为共同语言。只有当规范框架逐步落实为行业实践中的可执行文件，AI 品牌资产建设才可能真正从“机会表达”进入“能力建设”阶段。

第十章 全球 AI 品牌资产建设十大原则

在全球化运营的背景下，品牌面对的是多语言、多文化、多模型的复杂 AI 生态。不同区域的主流大模型（如 OpenAI、Google、Anthropic、DeepSeek 以及各区域本土模型）在训练数据、算法偏好、伦理对齐上存在差异，这可能导致同一品牌在不同模型中呈现出截然不同的认知画像。为了确保品牌在全球 AI 语义空间中获得稳定、可信的认知主导权，本白皮书提出“全球 AI 品牌资产建设十项原则”，作为跨国企业与出海品牌开展 AIBE 建设的最高行动指南。这些原则旨在帮助品牌构建一套能够跨越语言、文化、平台与模型壁垒，被全球 AI 系统稳定理解、审慎引用与持续复用的可信知识体系。

10.1 可识别原则

品牌首先必须在全球 AI 语义空间中形成稳定、唯一且可持续识别的实体身份。在全球大模型中，品牌名称往往面临跨语种翻译歧义、缩写冲突、与通用词汇混淆乃至与历史上同名实体冲突的风险。例如，一个名为“Apple”的品牌，在农业语境下可能被 AI 误读为水果，而非科技公司。品牌必须通过构建全球统一的知识图谱（Knowledge Graph），向 AI 系统明确声明其全球唯一实体身份（Entity ID），包括品牌名称的官方拼写、多语言标准译名、母公司关系、核心业务领域、历史沿革等关键信息。这需要品牌主动向主流 AI 平台提交结构化的实体声明文件，并确保其在全球权威数据库（如维基百科、行业名录、企业注册信息）中的信息一致且准确。只有当 AI 能够准确识别品牌实体，并将其与竞品、通用概念清晰区分时，后续的理解、引用与推荐才具备基础。

10.2 权威信源原则

建立全球可信引用资产。全球主流 AI（如 ChatGPT、Claude、Perplexity）在生成高价值答案时，高度依赖高权重域名的信息。品牌不能仅依赖自有官网，因为单一信源在 AI 的重排机制中权重有限。品牌必须在全球范围内构建权威信源矩阵，包括：国际权威行业组织（如 IEEE、ISO、W3C）发布的相关文件、全球知名媒体（如《华尔街日报》、《金融时报》）的深度报道、国际标准化组织（ISO）或区域权威机构的认证文件、知名学术期刊或研究机构的相关引用、以及全球知名评测机构（如 Gartner、Forrester）的分析报告。这些高权重信源是品牌进入全球 AI 重排阶段的“硬通货”。品牌应系统性地管理这些信源，确保其内容准确、更新及时，并建立与这些信源的主动关联机制（如提供官方数据、参与行业白皮书撰写），从而提升品牌信息在 AI 检索与证据筛选中的优先级。

10.3 多语言一致性原则

多语言环境中的认知漂移，正在成为全球品牌 AI 治理的重要风险之一。同一个品牌，在英文语料中可能因丰富的技术文档和媒体报道被 AI 视为“高端创新者”，而在小语种语料中可能因信息匮乏或翻译错误被降级、误读甚至与负面信息关联。品牌必须建立跨语言的标准化语料库，确保核心价值观、产品参数、技术优势、品牌故事在英语、西班牙语、阿拉伯语、中文等不同语言的大模型训练集和检索库中保持绝对的语义一致性。这要求品牌进行专业的本地化与翻译管理，而非简单的机器翻译。品牌需要定义一套全球统一的“品牌语义词典”，明确关键术语、

价值主张、产品特性的标准表述，并确保其所有语言版本中严格对应。同时，要定期监测不同语言模型对品牌的认知差异，一旦发现“认知漂移”，立即通过补充高质量本地化内容、修正错误翻译、强化权威信源等方式进行干预，消除因语言壁垒导致的认知偏差。

10.4 结构化知识原则

知识必须机器可读。生成式 AI 更倾向于优先采纳结构明确、边界清晰且能够低成本解析的信息。AI 系统不擅长处理冗长、感性、充满营销话术的长文。品牌面向全球输出的信息必须高度结构化，采用 JSON-LD、Schema.org 等机器可读标记，将产品规格、服务条款、全球网点、价格信息、技术参数、适用场景等关键信息转化为清晰的键值对（Key-Value）或结构化数据表。例如，对于一款全球销售的产品，应提供包含型号、尺寸、重量、材质、认证标准、适用电压、全球上市时间等字段的结构化数据。这些结构化数据能够被 AI 的检索系统直接索引和比对，在毫秒级的 RAG 检索中被优先召回。品牌应建立全球统一的知识管理平台，集中维护和分发这些结构化知识，确保其与官网、产品手册、客服系统等渠道的信息实时同步。只有让 AI “秒懂”，品牌才能在激烈的全球 AI 认知竞争中占据优势。

10.5 可验证原则

必须存在第三方证据链。在全球市场，自我标榜的营销话术（如“全球销量第一”、“行业领导者”）因为缺乏可信依据，通常会被 AI 的重排器（Reranker）降低优先级。品牌主张的每一项核心优势，都必须在互联网上存在可交叉验证的第三方数据源。例如，若品牌宣称“碳

中和达标”，应提供权威第三方机构（如 Carbon Trust）的认证证书链接；若宣称“用户满意度最高”，应提供知名调研机构（如 J.D. Power）的公开报告引用。品牌需要系统性地构建完整的证据链，将每一项关键主张与一个或多个可验证的第三方来源关联起来。这不仅提升了 AI 引用率（Citation Rate），也增强了品牌在用户心中的可信度。在建设过程中，品牌应主动与第三方机构合作，获取其背书或数据授权，并将这些证据以结构化方式整合进自身的知识库中，使其成为 AI 可检索、可引用的“事实锚点”。

10.6 风险边界原则

必须清晰声明适用边界。面对全球不同国家复杂的法律法规（如医疗、金融、数据隐私）以及文化差异，AI 模型倾向于采取保守策略，对可能存在风险的信息进行弱化或回避。品牌必须主动、清晰地向 AI 声明产品的适用范围、禁忌症、免责条款、区域限制及潜在风险。例如，一款医疗设备应明确标注“仅适用于 XX 国家/地区，需在医生指导下使用”；一款金融产品应说明“不构成投资建议，过往业绩不代表未来表现”。这些边界信息应以结构化、标准化的方式呈现，并嵌入到品牌的核心知识单元中。清晰的风险边界不仅不会降低推荐率，反而会增加 AI 在特定合规场景下引用该品牌的安全感，因为 AI 能够基于明确的边界信息，做出更负责任、更符合法规要求的推荐。品牌应建立全球合规审查机制，确保所有面向 AI 输出的内容都经过法务与合规部门的审核，符合目标市场的所有相关法规要求。

10.7 AI 可引用原则

内容必须适配 RAG 与 AI 重排。品牌的内容生产流程必须重构，从“写给人类看”转变为“写给 RAG 系统看”。段落应遵循“结论先行、数据支撑、逻辑清晰”的原则，信息密度要高，避免冗长的铺垫和情感渲染。每一段落或知识单元应具备独立的语义完整性，能够回答一个具体问题或阐述一个核心观点，这样在被切分为文本块（Chunk）后，依然能被 AI 准确理解和召回。品牌应采用“问题-答案”对（Q&A Pair）或“主题-陈述”对（Topic-Statement Pair）的形式组织内容，并确保每个答案或陈述都包含必要的上下文、数据来源和适用条件。此外，内容应使用清晰、客观的语言，避免模糊、夸张或主观性强的表述。通过优化内容的“可引用性”，品牌可以显著提升其在 AI 检索与重排阶段的得分，从而在答案生成中获得更高的出现概率和更优的展示位置。

10.8 全球平台一致性原则

不同模型中保持一致。全球用户使用的 AI 工具千差万别，包括 OpenAI、Google、Anthropic、Microsoft 以及各区域本土模型。这些模型在训练数据、算法偏好、对齐策略上存在差异，可能导致对同一品牌的认知不一致。品牌不能只针对单一模型进行优化，而应基于底层可信知识网络（KNIT），向全网释放标准化语料。这意味着品牌的知识建设应遵循平台无关的原则，确保其核心事实、价值主张、产品信息在任何模型中都是准确、一致且可理解的。品牌应定期使用多模型、多入口的测试集（如 AIBV 评测体系）对自身的 AI 品牌资产表现进行监测，一旦发现不同模型间存在显著认知差异，应分析原因（如特定模型训练数据偏差、算法偏好），并通过补充通用性更强、证据更充分的内容进行修

正。目标是构建一个“模型中立”的品牌知识底座，确保品牌在全球任何 AI 入口中都能呈现出高度一致、稳定可靠的品牌画像。

10.9 持续治理原则

长期监测 AI 误读。AI 模型的认知是动态演进的，且存在“幻觉”风险，即可能生成与事实不符的内容。品牌必须建立全球 AI 舆情与认知监测机制，定期使用多语种、多场景的问题集对主流大模型进行探测。监测范围应包括品牌被提及的频率、在答案中的角色（首选、备选、风险提示）、引用的权威性、以及是否存在事实错误、负面偏见或过时信息。一旦发现 AI 对品牌产生误读、错引或生成负面偏见，需立即启动“知识补丁”修复流程：首先，核实问题根源（是自身信息缺失、错误，还是外部虚假信息干扰）；其次，通过更新权威信源、发布澄清声明、补充结构化事实等方式，向 AI 系统提供正确的信息；最后，持续跟踪修复效果，确保误读得到纠正。这种持续治理能力是 AI 品牌资产建设不可或缺的一环，它将品牌从被动的“被认知者”转变为主动的“认知管理者”，确保品牌在 AI 时代的认知主导权始终掌握在自己手中

10.10 合规与伦理原则

符合全球 AI 治理框架。品牌的 AI 资产建设必须严格遵守全球范围内的 AI 治理法规与伦理准则。这包括欧盟的《人工智能法案》、美国的《人工智能权利法案蓝图》、中国的《生成式人工智能服务管理暂行办法》等。坚决杜绝使用虚假数据投毒、恶意抹黑竞品、侵犯用户隐私、生成歧视性内容等黑帽手段。品牌应建立内部的 AI 伦理审查委员会，对所有面向 AI 输出的内容进行伦理与合规评估，确保其符合人类共同的伦

理道德标准，如公平、透明、尊重隐私、避免伤害等。坚持科技向善，将 AI 品牌资产建设视为提升社会信息质量、促进良性竞争的积极力量。只有建立在合规与伦理基础上的 AI 品牌资产，才是真正可持续、可信赖的资产，才能穿越全球监管与公众审视的考验，为品牌带来长期的价值。

展望 AI Agent 时代：从“被理解”到“被调用”。如果说 2025—2026 年是 AI 品牌资产建设的“AI 入口期”，那么 2027 年起，行业将进入“AI Agent 执行期”。届时，Agent 不再只是替用户阅读品牌信息，更会代表用户直接下单、订房、签约、报修等。品牌资产建设的下一战场，将是“Agent 可调用性”——包括 MCP 协议接入能力、Function Calling Schema 设计、API 的可发现性，以及面向 Agent 的友好定价与履约策略。这意味着 AIBE 的边界，将从“被理解、被引用”进一步扩展为“被调用、被执行”。今天的可信知识网络（KNIT）建设，正是为明天 Agent 时代的品牌可调用性打下基础。

从长期看，全球 AI 品牌资产竞争的核心，正在从“谁更会传播”转向“谁能够提供更可信、更结构化、更可持续的知识供给”。这不仅是品牌建设逻辑的变化，也意味着全球数字竞争秩序正在进入新的阶段。

总结

过去二十年，品牌习惯于为获取曝光而配置预算、组织内容与建立渠道；但在生成式人工智能重塑信息获取与判断方式的今天，品牌所面对的已不再只是“如何被看见”，而是“如何在 AI 的判断链条中被正确理解、审慎表达与合理引用”。当越来越多用户将初筛、比较与归纳的任务交给 AI，品牌竞争的关键环节已经从传统流量入口前移到答案生成、候选形成与证据组织的阶段。品牌真正面临的风险，不再只是排位靠后，而更可能是在新的入口中逐步缺席。

也正因如此，AI 品牌资产建设不应被误解为一种针对生成式 AI 系统的操纵技巧，更不应被引向黑帽 GEO、AI 投毒、虚假信息铺设、伪造信源或误导性承诺。真正值得建立的，不是对结果的控制力，而是品牌在 AI 时代的信息供给能力：是否真实、是否准确、是否有边界、是否可追溯、是否具备被系统采用的证据属性。GEO 的本质绝不是算法欺骗，而是提升供给侧信息的结构化程度、准确性与可验证性。

从更长期的视角看，品牌真正需要争取的，不是某一次回答中的偶然位置，而是一种更稳定的存在状态：当用户提出关键问题时，系统是否能够认得出你、分得清你、说得准你、信得过你，并在必要时引用你。这种存在状态，不会靠单点动作获得，而只能通过品牌事实库、答案资产、引用资产、可信知识网络、口径治理、风险修复与持续评估逐步沉淀。

因此，在 AI 时代，品牌需要完成一次重要的认知转换：从“争夺结果”转向“建设资产”，从“流量博弈”转向“品牌治理”，从“话术

承诺”转向“证据建设”，从“短期有效”转向“长期可信”。真正能够穿越模型迭代、平台变迁与规则收敛的，不会是最擅长制造噪音的品牌，而是最早完成真实、合规、透明、可验证与可持续治理升级的品牌——它们将在 AI 时代获得更稳固、更可累积的认知主导权。

附录

A. AI 品牌资产建设表现指数体系（AIBV 1.0）

AIBV 1.0 用于衡量品牌在中国主流 AI 大模型与 AI 搜索/推荐场景中的“语义空间品牌资产”表现、建设能力与风险状况，并为行业提供统一评价语言与可量化对标框架。本体系采用 AIBV 的“3+2”结构与方法置信度机制。

关于该体系的完整定义、总体框架、指标设置及使用规范，详见附件《AI 品牌资产建设表现指数体系（AIBV 1.0）》。

B. 术语表（精简版）

KNIT：可信知识网络（Knowledge Network of Integrity & Trust）

GEO：生成式引擎优化

AEO：答案引擎优化

RAG：检索增强生成

SoA：Share of Answer，答案份额

AIBE：AI 品牌资产

AIBV：（AI Brand Asset & Value Development Index）AI 品牌资产建设表现指数体系

Answer Asset：答案资产

Citation Asset：引用资产

Brand Agent：品牌智能体

C. 指标定义

SoA=在问题集中的“出现×权重×顺位”加权占比

引用率=AI 输出引用链中品牌资产占比（按来源权威度加权）

认知一致性=不同问题下核心判断的稳定程度（理由一致/标签一致）

情感倾向=AI 总结的正负/风险标签出现频次及强度

附件 AI 品牌资产建设表现指数体系（AIBV 1.0）

前言

随着生成式人工智能和大模型技术的快速普及，用户获取品牌信息、比较产品与做出决策的路径正在发生深刻变化。越来越多的品牌认知、产品比较与消费决策，不再只发生在传统搜索引擎和媒体渠道，而是在 AI 问答、AI 搜索、AI 推荐、AI 助理和智能代理等语义交互场景中完成。

在这一背景下，品牌资产正在进入一个新的衡量维度：品牌是否能够被主流 AI 系统准确理解、自然召回、合理推荐、稳定表达，并在复杂场景下形成可信、可执行、可持续的影响力。这种存在于 AI 语义空间中的品牌表现、建设能力与风险状况，构成了 AI 时代品牌资产的重要组成部分。

为此，我们提出 AI 品牌资产建设表现指数体系（AI Brand Asset & Value Development Index，简称 AIBV）。本体系旨在建立一套面向 AI 时代的品牌资产评价框架，用于衡量品牌在主流 AI 模型、AI 搜索与推荐场景中的综合表现、建设水平与风险状态，为企业治理、行业研究、品牌治理、投资判断与标准建设提供统一的评价语言和量化依据。

AIBV 不是传统品牌价值评估模型的简单延伸，也不以广告投放、媒体曝光或公关传播规模为核心依据，而是聚焦品牌在 AI 输出端和 AI 语义空间中的真实存在状态。它既可作为企业开展 AI 品牌资产诊断与提升的工具，也可作为行业报告、评价规范和相关标准建设的技术基础。

需要说明的是，随着行业逐步从方法探索走向规则收敛，AIBV 的功能也不应局限于“结果测量”。围绕高质量知识供给、可信证据引用与

评测过程治理，AIBV 还应进一步成为行业标准化的重要接口。未来，围绕权威认证语料库、标准问答单元、答案资产、引用资产以及 AI 品牌资产建设服务机构，逐步建立可核验、可分级、可认证的配套规则，有望成为 AIBV 实施细则持续完善的重要方向。

这意味着，AIBV 不只是对品牌在 AI 语义空间中“表现如何”的评价体系，也将逐步承担“什么样的知识供给更可信、什么样的建设动作更规范、什么样的服务交付更可审计”的规则支撑功能，从而推动行业形成从知识建设到评估治理再到标准实施的闭环。

第一章 总则

1.1 目的与定位

AIBV 的主要目的在于：

- 建立一套面向 AI 时代的品牌资产评价体系，反映品牌在 AI 生态中的综合表现和发展水平。
- 为政府部门、行业组织、媒体、投资机构和企业提供统一的评价语言和量化参考。
- 引导企业重视 AI 语义空间中的品牌建设，提升 AI 场景下的品牌安全性、竞争力与经营效率。
- 为 GEO（生成式引擎优化）及相关品牌增长实践提供可测量、可比较、可追踪的评价框架。

1.2 适用范围

本体系适用于：

- 在中国境内运营、具有一定公众认知基础的企业或组织机构品牌和产品品牌。
- 优先适用于上市公司品牌、行业头部品牌、区域龙头品牌及重点公共服务品牌。
- 适用场景包括但不限于：
AI 对话/问答；
AI 搜索与推荐；
AI 助理类产品中的品牌表达与推荐；
与品牌认知、比较、决策、行动相关的语义交互场景。

1.3 不在评价范围内的内容

AIBV 不直接评价以下内容：

- 传统媒体广告投放规模、曝光量等媒介指标；
- 品牌整体商业价值、社会责任表现等非 AI 语义空间维度；
- 企业市值、利润、营收等经营结果本身。

1.4 基本原则

AIBV 的构建与使用遵循以下原则：

- 客观性：基于统一的多模型测试、结构化标注和明确的量化公式，尽量减少主观干预。
- 可复现性：数据采集流程、指标定义和计算方法公开透明，可由第三方复核。
- 可解释性：通过多维度、多指标呈现，使企业和公众能够理解得分来源。

- 防操纵性：通过暗题、随机题、异常检测和模型差异检测，降低“刷榜”与定向操纵空间。
- 合规性：符合数据安全、个人信息保护、广告监管等相关法律法规要求。
- 实用性：能够支撑企业诊断、行业对标、品牌治理和 GEO 优化实践。

第二章 术语与定义

2.1 AI 品牌资产

AI 品牌资产，是指品牌在主流 AI 大模型和 AI 应用场景中的综合表现，包括被 AI 准确理解、自然召回、合理推荐、稳定表达、可信引用、有效执行，以及品牌自身在 AI 建设方面的投入与能力等所形成的综合价值。

2.2 AIBV

AIBV (AI Brand Asset&Value Development Index)，即 AI 品牌资产建设表现指数体系，是基于多模型问答与生成实验、用户调研、品牌侧数据构建和专家团队顾问评估的交叉认证的综合评价体系，用于系统刻画品牌在 AI 语义空间中的表现、建设能力与风险状况。

2.3 三个主指数

- AIP (AI Performance Index) AI 表现基础指数：反映品牌在 AI 输出端的实际表现。
- AIC (AI Construction Index) AI 建设指数：反映品牌为进入 AI 语义空间所进行的内容、系统和组织建设能力。

- AIR (AIRisk&ReliabilityIndex) AI 风险与稳定性指数：反映品牌在 AI 语义空间中的错误、偏差、波动与操纵疑似等风险状况。

2.4 两个校准因子

- UAF (User - AI Alignment Factor) 用户 - AI 对齐系数：用于反映用户对 AI 的实际使用程度、信任程度以及 AI 对品牌决策的影响程度。
- MCI (MethodConfidenceIndex) 方法置信度指数：用于反映本次评测的数据质量、样本充分性、测试稳定性和多模型覆盖程度。

第三章 总体框架

3.1 “3+2” 组合结构

AIBV 采用 “三大主指数+两个校准因子” 的结构：

- 主指数：AIP、AIC、AIR
- 校准因子：UAF、MCI

在对外展示、研究分析与行业传播中，应同时展示上述五项关键值，而不建议仅展示单一综合分。

3.2 分数区间

各主指数、维度指标及子指标统一映射为 0 - 100 分。分值越高代表表现越好；风险项按 “风险越低、得分越高” 的原则进行换算。

3.3 输出要求

- AIP：AI 表现基础指数；

- AIC: AI 建设指数;
- AIR: AI 风险与稳定性指数;
- UAF: 用户 - AI 对齐系数;
- MCI: 方法置信度指数;

如使用综合指数, 应同时说明公式、版本号、测量窗口与置信提示。

第四章 AIP : AI 表现基础指数

4.1 维度构成

AIP 由四个维度组成:

- A: 认知准确度 (Accuracy)
- B: 可见性与召回度 (Visibility & Recall)
- C: 场景适配度 (Scenario Fitness)
- D: 一致性与调性匹配 (Consistency & Tone)

4.2 组合公式

AIP 的基本计算公式为:

$$AIP=wA \times A+wB \times B+wC \times C+wD \times D$$

推荐初始权重为: $-wA=0.30-wB=0.30-wC=0.20-wD=0.20$

权重参数可结合行业特征和实证结果进行周期性校准。

4.3 维度 A: 认知准确度

维度 A 用于衡量 AI 是否准确理解品牌、描述品牌, 并能够提供可信且具区分度的认知信息。

包括以下指标:

A1 品牌事实正确率：核心事实字段（主体、品类、产品线、价格带、核心能力、关键限制条件等）被正确表述的比例。

A2 核心要素覆盖度：回答对品牌关键要素集合的覆盖程度，关键要素按行业模板定义。

A3 引用/证据可追溯度：衡量回答中是否提供可溯源的出处、来源提示，以及引用来源的权威等级。

- 可采用两个子项：

Cite Rate：回答中出现可溯源出处/来源提示的比例；

Auth Tier：引用来源等级加权得分。

- 建议来源等级：政府/标准>权威机构>权威媒体>方法透明测评>UGC。

A4 区分度：回答是否能够稳定指出品牌与竞品的关键差异点，避免同质化套话。

维度 A 得分为上述指标按设定权重加权平均。

4.4 维度 B：可见性与召回度

维度 B 用于衡量品牌在不同 AI 场景中被提及、被召回、被展示以及被优先呈现的能力。

包括以下指标：

B1 指名检索可见性：用户直接询问品牌名称时，品牌是否稳定出现且描述正确。

B2 类目检索可见性：用户以品类、需求或任务方式提问时，品牌进入候选集合的比例。

B3 召回排序指数：按品牌在答案中的排序位置进行加权统计，适用于显式列表或可判定位次的答案结构。

B4 覆盖广度指数：在多问题类型、多人群、不同约束条件下的整体进入率。

B5 答案结构位置加权展示度（PAD）：对品牌在首段、结论区、Top 列表、中段、尾段等不同答案结构位置的展示进行加权，反映品牌在答案结构中的真实可见度。

在对外传播或项目管理中，可将 B3 与 B5 组合形成增强展示指标，用于更准确反映品牌在 AI 答案中的呈现价值。

4.5 维度 C：场景适配度

维度 C 用于衡量品牌在具体决策场景中的推荐合理性、场景覆盖能力与行动闭环能力。

包括以下指标：

C1 场景推荐匹配度：在预算、城市、人群、偏好、风险限制等强约束问题中，品牌被推荐的合理性。

C2 场景覆盖度：品牌覆盖关键决策场景的程度，关键场景由行业高价值场景清单定义。

C3 推荐理由链完整度：回答是否给出可复述的“为什么选它”的完整理由链，包括卖点、证据或逻辑、适用条件和注意事项。

C4 答案可执行闭环度：回答中是否给出行动所需的关键要素字段，例如渠道、门店、预约、售后、政策、时间戳、操作步骤等。

维度 C 得分为上述指标按设定权重加权平均。

4.6 维度 D：一致性与调性匹配

维度 D 用于衡量品牌在不同模型、不同时间和不同表达环境中的稳定性与品牌风格一致性。

包括以下指标：

D1 跨模型一致性指数：不同模型回答之间的语义相似度或关键要素一致程度。

D2 跨时间稳定性指数：不同时间批次回答之间的一致性程度。

D3 调性匹配度：回答语气、价值观与品牌调性的匹配程度。

维度 D 得分为上述指标按设定权重加权平均。

第五章 AIC：AI 建设指数

5.1 指数定义

AIC 反映品牌为了提升其在 AI 语义空间中的存在质量、推荐能力与治理能力所进行的系统性建设水平。

5.2 指标构成

AIC 由三个能力类指标构成，可采用 0 - 5 级成熟度评估并映射为 0 - 100 分：

AIC 由三项能力指标构成：

E1 知识资产建设度：衡量品牌事实库、FAQ、答案资产、引用资产及知识更新机制的完备程度；

E2 认知治理成熟度：衡量企业在问题集管理、口径治理、误读修复、监测反馈与持续迭代方面的治理能力；

E3 信源与证据建设度：衡量品牌在第三方验证、权威来源锚定、合规披露与证据可追溯性方面的建设水平。

5.3 计算方法

$$AIC=v1 \times E1+v2 \times E2+v3 \times E3$$

推荐初始权重为： $v1=1/3-v2=1/3-v3=1/3$

如具备行业实证基础，可进行权重调整。

第六章 AIR：AI 风险与稳定性指数

6.1 指数定义

AIR 用于衡量品牌在 AI 语义空间中的错误风险、负面偏差、结果波动及操纵疑似等风险情况。其原则为“风险越低，得分越高”。

6.2 指标构成

AIR 由四项构成：

- R1 严重错误控制度：基于严重错误率映射得分。
- R2 负面偏差控制度：基于负面偏差程度映射得分。
- R3 稳定性得分：基于多次测量的波动程度映射得分。
- R4 公平性与操纵疑似度：基于公开题/暗题差异、模型异常差异等识别操纵疑似度与不公平表现。

6.3 计算方法

$$AIR = u_1 \times R_1 + u_2 \times R_2 + u_3 \times R_3 + u_4 \times R_4$$

推荐初始权重为： $-u_1=0.30$ $-u_2=0.30$ $-u_3=0.20$ $-u_4=0.20$

6.4 严重错误分类披露要求

为提升 AIR 的可解释性与治理价值，建议对 R1 至少按以下类型进行分组披露：

- 事实错误（主体、品类、参数等）；
- 过期信息（政策、价格、门店、版本等）；
- 合规红线或不当承诺；
- 幻觉断言（无依据的强结论）。

分类披露不改变 R1 的计算公式，但能够显著提升问题定位与治理效率。

第七章 UAF 与 MCI

7.1 UAF：用户-AI 对齐系数

UAF 用于衡量用户在品牌决策过程中对 AI 的使用频率、信任程度以及 AI 对品牌选择的实际影响。

建议计算方式：

$$\text{UAF_raw} = 0.4 \times \text{Usage_norm} + 0.3 \times \text{Trust_norm} + 0.3 \times \text{BrandUsage_norm}$$

$$\text{UAF} = 0.7 + 0.3 \times \text{UAF_raw}$$

UAF 取值映射至 [0.7, 1] 区间，用于对不同场景下 AI 影响强度进行校准。

7.2 MCI：方法置信度指数

MCI 用于反映本次评测方法的可信度与可解释程度。

建议计算方式：

$$\text{MCI_raw} = 0.25 \times S + 0.25 \times Q + 0.25 \times T + 0.25 \times M$$

其中：-S：样本充分性-Q：数据质量-T：测试稳定性-M：多模型覆盖度

$$\text{MCI} = 0.8 + 0.2 \times (\text{MCI_raw} / 100)$$

MCI 取值映射至 [0.8, 1] 区间。对外发布结果时，应同步附带 MCI 提示，以说明结论置信程度。

第八章 综合指数

在需要形成单一综合指标用于内部管理、项目追踪或对外传播时，可在同时展示 AIP、AIC、AIR、UAF、MCI 的前提下，计算 AIBV 综合指数。

公式如下：

$$AIBV = MCI \times [\lambda \times AIP + (1 - \lambda) \times (\alpha \times AIC + \beta \times AIR)] \times UAF$$

其中： λ 建议取值范围为 0.5 - 0.7； α 与 β 之和为 1，可根据研究目的设定；-推荐初始值： $\lambda=0.6$ ， $\alpha=0.6$ ， $\beta=0.4$ 。

综合指数适用于内部管理、趋势追踪与传播补充，不建议替代五项关键值的独立展示。

第九章 数据采集与测量规范

9.1 多模型多入口原则

同一问题集应在多模型、多入口下进行测量，形成具有可复现性的评测数据集。应记录模型名称、版本号、调用方式、关键参数及测试时间窗口。

9.2 问题集分层原则

测试题集应至少覆盖以下五类用户意图：

- 认知；
- 对比；
- 决策；
- 行动；
- 风险。

同时应保证强约束问题达到实施细则规定的行业阈值，以提升评测对真实决策场景的覆盖能力。

9.3 题源构成

题目可来自以下渠道：

- 专家设计的标准问题；
- 真实搜索与问答语料的抽样与规范化问题；
- 用户调研中收集的自然语言表达。

9.4 留痕与可审计

应完整保留：

- 问题文本；
- 完整回答；
- 时间戳；
- 模型与入口信息；
- 结构化标注结果；
- 关键计算中间过程。

9.5 标注与质量控制

应建立统一标注规范，并采取以下质量控制措施：

- 多人交叉复核；
- 关键指标抽样复查；
- 对高波动题目进行单独识别与处置；
- 对噪声、拒答和无效回答进行统一清洗规则管理。

第十章 防操纵与质量保障

10.1 暗题与随机题机制

正式评测题集中应设置一定比例的暗题和随机题，用于识别特定品牌针对公开题进行定向优化的行为。暗题机制应作为正式评测的必要组成部分。

10.2 异常检测

应对以下情况开展异常检测：

- 公开题与暗题表现差异异常；
- 不同模型之间对同一品牌存在显著异常偏好；
- 个别答案结构或推荐路径呈现出非自然集中现象。

10.3 数据质量管理

数据质量因素应纳入 MCI 评估，包括：

- 有效回答率；
- 噪声水平；
- 重复测量误差；
- 标注一致性。

10.4 结果发布要求

对外发布结果时，应至少附带：

- 指数名称；
- 评价年份与版本号；
- 发布机构名称；
- 测量窗口；
- MCI 提示；
- 必要的边界说明与合规提示。

第十一章 结果表达与等级划分

11.1 结果表达

建议对每个被评价品牌至少公开以下内容：

AIP：AI 表现基础指数；

AIC：AI 建设指数；

AIR：AI 风险与稳定性指数；

UAF：用户 - AI 对齐系数；

MCI：方法置信度指数；

如使用综合指数，可同步给出综合分及相应说明。

11.2 等级划分示例

为便于理解与比较，可设置如下示例等级：

AIP 等级（表现）-A： ≥ 85 -B：70 - 84-C：55 - 69-D： < 55

AIR 等级（风险）-低风险： ≥ 85 -中风险：70 - 84-较高风险：55 - 69-高风险： < 55

AIC 等级（建设）-建设领先： ≥ 80 -建设中等：60 - 79-建设起步： < 60

等级边界可根据样本分布和行业情况适度校准。

第十二章 使用规范与合规要求

12.1 指数使用边界

AIBV 反映的是品牌在主流 AI 模型中的语义表现及相关建设情况，属于“AI 语义空间的品牌资产”，不等同于对品牌整体商业价值、社会价值或合规性的全面评价。

12.2 对外引用规范

品牌、媒体和合作机构在引用 AIBV 结果时，应注明：

- 指数名称；
- 评价年份；
- 版本号；
- 发布机构名称。

不得使用容易引起误解的绝对化表述，不得混淆不同指数、不同版本和不同测量窗口的结果。

12.3 争议处理与申诉机制

- 品牌可在结果正式发布前，对自身数据和描述进行核对；
- 对事实性错误，可提交证据申请修正；
- 对涉及重大负面风险的内容，可申请技术与专家复核；
- 申诉流程与时限应在实施细则中明确。

12.4 高风险品牌的信息披露

对于 AIR 得分较低且风险明显的品牌，建议优先采取非公开反馈方式进行治理支持；如需公开披露，应在充分告知、事实复核和合规审慎的前提下实施。

12.5 认证结果引用规范

涉及语料认证、内容认证或服务认证的对外表述，应同步注明认证对象、认证范围、认证版本、认证机构及有效期限，不得将局部认证结果扩大解释为对品牌整体价值、整体合规性或全面经营表现的最终判断。认证结果的传播与使用，应与 AIBV 的评价边界保持一致。

第十三章 版本管理与扩展应用

13.1 版本管理

AIBV 1.0 采用滚动优化机制，可根据实践反馈、样本积累与实证研究结果形成后续迭代版本，如 1.1，1.2 等。

13.2 国内与国际应用

AIBV 当前主要面向中国境内主流模型与中文场景。未来可在合规前提下扩展至多语言环境与海外模型评测，形成 AIBV-Global 版本。

13.3 实施细则

为保障行业可执行性，并推动 AIBV 从评价框架进一步走向行业标准接口，建议在本体系基础上持续完善以下实施文件：

（1）行业实施细则

用于明确 AIBV 在不同行业、不同场景与不同发布用途中的适用边界、披露要求与执行口径。

（2）题库与问题集规范

用于统一问题集分层标准、题源构成、高价值问题定义及不同赛道的强约束题比例要求。

（3）标注与评分手册

用于统一核心指标、结构化标注方法、评分逻辑与异常识别口径，保证不同机构和不同批次之间的评测一致性。

（4）质量控制与申诉流程

用于规范异常检测、事实复核、争议处理、结果修正与发布前核对机制，提高体系的可复核性与治理价值。

（5）指标校准与实证报告

用于根据行业实践、样本积累与版本迭代持续优化权重设置、校准参数与方法置信度机制。

（6）语料认证规则

用于定义权威语料单元的准入标准、审核要求与分级方式，重点覆盖来源真实性、结构完整性、证据可追溯性、边界清晰度、更新机制与合规审查情况。

（7）内容认证规范

用于明确标准问答、FAQ、答案资产、引用资产、案例页面等知识产品的质量要求、边界说明与可引用条件。

（8）服务认证细则

用于对提供 AI 品牌资产建设、可信知识网络建设、GEO 运行优化与评估治理服务的机构进行透明度、可审计性、合规性与治理能力的分级评价。

（9）认证复核与动态更新机制

用于明确语料认证、内容认证与服务认证后的复核、抽查、申诉、撤销与更新规则，保证认证体系本身的持续有效性与可信度。

通过上述实施文件的持续完善，AIBV 将不只是一个评估框架，也将逐步成为连接可信知识网络建设、行业规范落地与认证治理机制的标准化接口。

结语

AIBV 旨在回答一个 AI 时代的新问题：当用户越来越多地通过 AI 认识品牌、比较品牌和选择品牌时，品牌在 AI 语义空间中的真实存在状态应如何被衡量。

本体系不是对传统品牌理论的替代，而是对品牌评价框架在 AI 时代的重要补充。它所衡量的，不只是品牌是否“被提到”，更是品牌是否被准确理解、优先展示、合理推荐、可信引用、可执行承接，并在长期运行中保持稳定与安全。

随着 AI 应用的持续深入，AIBV 将持续迭代，推动形成更加成熟的评价方法、实施规范和行业标准，为企业的 AI 品牌建设、GEO 优化与品牌治理提供长期的方法论基础。