

RSI，大模型的下一站？

——AI研发自动化加速下的模型、算力与应用机会

行业研究 · 专题报告

传媒 · 计算机

投资评级：优于大市（维持）

证券分析师：张衡

021-60875160

zhangheng2@guosen.com.cn

S0980517060002

证券分析师：熊莉

021-61761067

xiongli@guosen.com.cn

S0980519030002

证券分析师：陈瑶蓉

021-61761058

chenyaorong@guosen.com.cn

S0980523100001

核心要点：RSI 重构模型生态，关注头部模型厂商、算力硬件及CSP与高价值应用场景



- **自动化AI研发加速，RSI正由理论进入落地前夕的阶段。** 1) 虽然完整RSI尚未实现，但AI已经进入代码编写、研究调试、Kernel优化、自动实验和部分后训练环节，局部研发反馈闭环正在形成。2) 从当前前沿模型厂商落地来看，模型、Agent、Evaluator、训练平台和算力基础设施开始被整合进同一套研发系统和工作流，AI逐步由研发辅助工具转向研发流程的重要参与者和工具。3) 从技术和产业落地来看，当前虽然暂时无法判断RSI实现的确切时点，但是从渐进落地的角度来看，RSI正处于持续演进渗透的阶段。
- **前沿模型厂商：模型竞争由单代能力比拼转向AI Research Factory竞争，研发正反馈与平台化壁垒有望强化。** 1) 从模型竞争来看，能力更强的模型既是争取客户的重要因素，同时也正逐渐成为研发下一代模型的生产工具；远期来看，模型能力如果能够持续提高Research Agent和研究人员产出，领先厂商可能形成“更强模型—更高研发效率—更快推出继任模型”的正反馈闭环。2) 模型厂商竞争要素也将由人才、数据和单次训练，扩展至模型、Agent、Evaluator、数据闭环、训练平台和算力的系统能力。3) 综合各方竞争要素，具备基础模型、云基础设施、Agent平台、企业客户及真实任务反馈的厂商，更有机会将技术领先转化为API调用、订阅收入、商业化和平台生态。
- **算力基础设施：训练与Research Inference共同扩张，AI研发自动化有望进一步打开算力产业链需求空间。** 1) 训练端来看，算力需求不再仅取决于单次前沿预训练规模，RSI落地渗透之下，参数规模、训练Token、模型代际频率、候选实验、Post-training、RL Rollout、验证及研发推理均可能同步增加，更快迭代和更高实验并行度可能形成对训练需求的乘数效应，带动算力需求的天花板持续提升。2) 重点关注计算芯片及ASIC、HBM与存储、先进封装、高速互联与光通信、交换设备、液冷、电力及数据中心等供给壁垒较高的基础设施环节；3) 同时看好具备大规模算力基础设施和持续资本开支能力的头部CSP。随着训练、Research Inference及Agent推理负载同步增长，算力需求将由单次模型训练向持续研发与推理扩散；具备GPU集群、自研ASIC、云平台 and 模型服务能力的CSP，有望同时受益于AI资本开支扩张和云端AI调用量提升。
- **高价值应用：模型能力加速通过MaaS、API与Agent平台向应用端外溢，优先关注可验证、可闭环的高价值场景。** 1) 模型能力持续提升、推理成本下降以及多模型接入门槛降低，将持续降低AI应用开发、试错和规模部署成本，并通过需求弹性提升调用频次、上下文长度和任务复杂度从而促进应用场景的扩展与渗透率提升。2) 软件工程、研发工具、企业知识管理、客户服务和数字内容等领域数字化程度高、结果可衡量、反馈周期短，有望率先形成Agent规模化渗透；药物、材料和芯片设计等研发场景亦具备较大的效率提升潜力。3) 应用公司的核心壁垒将由“能否接入模型”转向专有数据、核心 workflow、可靠 Evaluator 以及持续提高客户产出的能力。
- **风险提示：**完整RSI可能长期无法形成；商业化落地低于预期风险；竞争加剧风险；监管政策风险；技术进步低于预期风险等；

- [01] RSI：AI开始参与生产下一代AI
- [02] 从理论推演到工程实践：RSI的60年演进与当前拐点
- [03] 从自动化AI研发走向RSI：现实进展与远期路径
- [04] 前沿模型厂商的RSI路径：从研发提效到继任模型闭环
- [05] 产业影响与投资观点：AI研发加速下的模型、算力与应用机会

1. RSI: AI开始参与生产下一代AI

1.1 RSI的核心理念：持续提升“开发AI的能力”

- RSI（Recursive Self-Improvement，递归自我改进）是指 AI 系统参与设计、开发和训练更强的继任系统，而能力更强的继任系统又进一步提高 AI 研发能力，形成跨代持续的正反馈：
 - 更强的AI → 更强的AI研发能力 → 更快开发出更强的继任AI → AI研发能力进一步增强
 - RSI与普通模型升级的根本区别，不在于AI是否参与了代码编写、实验执行或训练优化，而在于“开发AI的能力”本身是否成为持续改进的对象。如果当前模型只是提高了单次任务的开发效率，属于AI辅助或自动化AI研发；只有当其开发出的继任模型能力更强、继任模型又能进一步提高下一代AI的研发效率，并且这一过程能够跨越多代持续发生时，才构成严格意义上的RSI。
 - RSI中的“Self”指的是当前模型通过Research Agent、Evaluator、训练平台和实验系统参与下一代模型研发，再由下一代模型继续改进整套AI研发系统。模型既是研发的成果，也逐渐成为生产下一代模型的生产工具。因此RSI 的关键不是“AI 能够改进 AI”，而是这套系统能否跨代、持续地提高 AI 自身的研发能力。

图：RSI的核心理念



资料来源：Anthropic Institute、OpenAI，国信证券经济研究所整理

请务必阅读正文之后的免责声明及其项下所有内容

1.2 RSI 的边界：当前多数“自我改进”仍未进入继任模型闭环

- **Self-refinement、Self-play和Coding Agent都不等于真正的RSI。**模型检查并修改自己的结论、通过 Self-play 产生训练数据、利用强化学习提高特定能力，或让 Agent 积累 Memory、Skill 和 Workflow，都可以改善当前系统的表现；但这些改进的对象仍主要是当次任务、当前模型或模型外部的执行系统，尚未进入“研发下一代 AI”的跨代循环
 - Self-refinement改进的是当次任务输出。模型生成答案后进行自我反馈和迭代修改，通常不改变模型参数，更不涉及继任AI研发。
 - Self-play、自训练和强化学习可以提高模型在特定任务上的能力。AlphaGo Zero通过自我对弈产生训练信号并持续变强，证明局部自我学习闭环能够成立，但系统并没有因此获得开发下一代通用AI的能力。
 - Agent改进主要发生在模型外部系统。Agent可以积累记忆、工具、Skill和工作流，使后续任务执行得更好，但改进对象仍是任务执行方式，而非AI研发能力本身。
 - AI辅助或自动化AI研发开始进入代码编写、Debug、实验执行、结果分析和训练优化等环节，是通向RSI最重要的现实路径；但只要研究方向、关键判断和继任模型开发仍主要由人类完成，就不能证明递归已经闭环。
- **只有当AI能够开发更强的继任系统，且继任系统又进一步提升AI研发能力，并使这一正反馈跨越多代持续发生时，才构成严格意义上的RSI。**根据Anthropic的官方Blog，当前Claude已经开始加速AI研发，但尚未达到能够完整、自主开发继任系统的RSI状态

图：RSI的判定体系



1.3 当前阶段：完整RSI尚未实现，AI研发自动化持续深化

- AI 参与 AI 研发已经进入多个真实场景。OpenAI 披露，GPT-5.6 已被内部研究人员用于诊断研发故障、优化训练系统、运行机器学习实验和分析结果，并建立了一组基于真实 AI 研发任务的评测。Anthropic 也披露，Claude 已经能够在目标明确的任务中执行实验优化，并在部分开放式研究中自主提出假设、运行实验和迭代方案。
- 当前进展主要体现在三个方面：AI 已能承担代码、Debug、实验运行、系统优化和结果分析；在目标和评价标准明确的任务中，局部实验闭环已经可以自动运行；在 AI 支持下，部分研发任务的表现、实验频率和单个研究人员的产出均有明显提升
- 但从严格定义口径看，完整RSI仍未实现：
 - 自主完成前沿继任模型的完整设计、训练、验证和交付；
 - 研究方向选择、目标设定、评价标准和关键结果判断；
 - 研发能力的提高能够跨越多代持续成立，并形成稳定、可重复的递归加速

图：AI研发自动化正在深化，完整RSI仍待跨越



资料来源：Anthropic Institute、OpenAI、Google DeepMind，国信证券经济研究所整理

2.从理论推演到工程实践：RSI的60年演进与当前拐点

2.1 RSI 的思想最早可以追溯到1965年

- RSI并不是伴随大模型才出现的新概念。其思想源头通常可以追溯至1965年：I. J. Good提出，如果机器设计本身也是一种智力活动，那么一台“超智能机器”就可能设计出能力更强的继任机器；由此反复迭代，最终形成能力快速上升的“智能爆炸”。Good没有使用今天的RSI术语，但已经描述了其最核心的递归正反馈机制；
- 1993年Vernor Vinge将这一讨论扩展为“技术奇点”：一旦出现超越人类的智能，技术发展轨迹可能发生难以预测的结构性的变化。
- 2003年Jürgen Schmidhuber提出Gödel Machine，将自我改进推进到形式化层面：系统只有在能够证明某次代码修改有助于提高目标函数后，才执行对自身程序的重写。这表明“系统修改自身”在严格假设下可以被形式化描述，但其依赖预设公理、效用函数和可证明性，仍主要是一种理论模型
- 2008年Stephen Omohundro进一步指出，具有稳定目标的自主系统可能产生能力提升、资源获取、自我保护等工具性倾向。RSI研究由此不再只讨论“能否自我改进”，也开始系统涉及目标稳定性、控制权和安全风险
- 早期研究已经提出了RSI的正反馈机制、自我重写形式以及潜在安全问题，但缺少强通用模型、可靠Evaluator、自动实验平台和大规模工程环境，因而长期停留在理论推演与形式化探索阶段。当前来看相比之前最大的变化是RSI实现所需的工程条件开始逐步走向成熟。

表：RSI的核心思想已存在逾60年

年份	代表人物 / 研究	核心命题	对RSI研究的贡献	局限与边界
1965	I. J. Good：首台超智能机器	机器设计本身是智力活动；超智能机器可能设计出更强机器，形成“智能爆炸”。	提出递归正反馈的核心机制，是RSI的重要思想源头。	Good并未建立今天统一的“RSI”术语或工程判据；属于前瞻性理论推演
1993	Vernor Vinge：技术奇点	超人类智能可能经由AI、网络智能、人机接口或生物增强等多条路径出现。	把超智能的可能路径及其社会影响扩展为更广泛的未来议题。	“技术奇点”范围大于RSI，不能将两者直接画等号
2003	Jürgen Schmidhuber：Gödel Machine	系统在证明某次修改有助于实现目标后，可以重写自身程序的一部分。	把自我指涉、自我重写与效用改进放进可形式化讨论的框架。	依赖形式系统、公理和可证明性，主要是理论模型
2008	Stephen Omohundro：Basic AI Drives	目标导向自主系统可能产生能力提升、资源获取与自我保护等工具性倾向。	将持续自我改进与控制、安全及目标稳定性问题联系起来。	理论行为与安全论证

资料来源：NeurIPS、Anthropic Institute、OpenAI、Google DeepMind，国信证券经济研究所整理

2.2 反馈闭环从封闭任务向研发流程扩展

- RSI的落地推进不是一蹴而就，而是反馈闭环不断扩大的过程。系统最初只能在规则明确的封闭任务中优化策略，随后开始利用自身生成的数据训练模型、反思并修正任务输出，当前则进一步进入代码实验和科研工作流。闭环覆盖范围在扩大，但距离“自主开发更强继任模型”仍有明显差距；
- 自我对弈、自生成数据、自我反思和自动化研究，都是 RSI 的重要前置能力。从 2017 年 AlphaGo 的策略闭环，到训练数据、任务输出，再到当前的研发流程闭环，AI 能够承接的环节越来越多。但当前研究目标、实验模板、工具环境和评价标准仍主要由人类设定；AI 可以执行提出想法、编写代码、运行实验和分析结果等环节，但尚未完成从研究方向到继任模型的全链条接管。

图：实践沿着“策略—训练数据—任务输出—研发流程”逐步推进，闭环范围越来越接近AI研发



2.3 完整RSI仍需时间，工程条件正逐步成熟

- 过去的自我学习系统通常只能在特定任务中运行：模型、反馈机制、实验环境和训练流程彼此割裂，需要研发团队手工连接。最近两年三的关键变化是模型能力、Agent执行、自动验证和研发基础设施开始被纳入同一套工程系统，AI由“给出建议”逐渐转向“持续完成研发任务”
 - 模型开始具备进入真实研发流程的基础能力。Reasoning和Coding能力的提升，使前沿模型不再局限于生成代码片段，而是能够理解真实代码库、修改多个文件、执行测试、分析日志并修复错误。OpenAI在2025年推出的Codex已经可以在独立云端沙箱中编写功能、修复Bug并提出Pull Request，软件工程Agent由能力展示开始转向真实工作流
 - Agent将单轮模型能力延伸为长周期执行能力。工具调用、代码执行、长上下文、持久记忆、错误恢复和任务编排，使模型能够连续完成多个相互依赖的步骤。METR的任务时间跨度评测显示，前沿Agent能够可靠完成的任务长度正在持续扩大；Anthropic进一步展示了依靠测试Oracle、进度文件、持续记忆和HPC环境运行多日科学计算任务的工作方式
 - Evaluator使部分研发任务能够自动形成反馈。代码测试、编译结果、数学证明、模拟环境和性能指标可以为候选方案提供快速、相对客观的评价。Google DeepMind的AlphaEvolve将大模型生成、自动Evaluator和进化搜索结合起来：系统批量生成程序，自动运行、验证和评分，并将更优方案重新投入下一轮搜索。它还被用于改善Google数据中心、芯片设计和AI训练流程
 - 研发基础设施开始把上述能力连接成可运行系统。云端沙箱、代码库、训练环境、实验编排、算力调度以及Memory、Skill和Harness，使Agent的执行过程可以被记录、复现和复用。模型因此能够逐步完成“提出方案—编写代码—执行实验—自动验证—保留更优结果—进入下一轮”的研发流水线，而不再只是完成孤立任务
- 虽然实现完全的RSI还需要时间，但是RSI的工程条件开始走向成熟。当前研究方向、评价标准、前沿训练决策和继任模型的最终验证仍主要由人类掌握；但 AI 研发第一次具备了部分自动运行和反复迭代的工程基础。完整 RSI 仍需时间，但所需的工程条件已在逐步成熟。

2.3 完整RSI仍需时间，工程条件正逐步成熟

- 虽然实现RSI还需要时间，但是RSI的工程条件开始走向成熟。当前研究方向、评价标准、前沿训练决策以及继任模型的最终验证仍主要由人类掌握，但是AI开始具备了部分自动运行和反复迭代的工程基础。

图：模型第一次能够在同一套工程系统中完成推理、执行、验证和持续迭代



2.4 部分模型与工具已形成局部递归反馈能力

- **普通AI通常是线性链条，不会形成闭环反馈机制：**AI完成客服、文案、数据处理等既定任务，产出成为产品或服务，但不会根据结果循环改进生成这些产出的AI系统。也就是可以执行任务、提高效率、降本增效，但是不会对生产工具本身即AI形成跨代升级
- **AI大模型研发具有不同的反馈结构。**当模型参与编写代码、运行实验、分析结果和优化训练系统后，这些产出可能被重新用于改进模型、Agent、Evaluator and 研发 workflow；如果由此产生更强的继任系统，继任系统又能承担更多、更复杂的AI研发任务，研发产出就可能重新进入下一轮能力提升；最终形成“当前模型 → Research Agent → 实验执行 → Evaluator验证 → 训练与系统改进 → 更强继任系统 → 更强AI研发能力”的循环体系；
- Google DeepMind 推出的 AlphaEvolve 是局部反馈闭环的代表性实践。该系统将大模型的创造性生成能力与自动 Evaluator 的验证反馈结合，通过迭代进化改进程序代码。在数据中心方面，其为集群管理系统 Borg 提供调度启发式算法，平均回收 Google 全球约 0.7% 的计算资源；在芯片设计方面，其优化 TPU 中矩阵乘法的关键电路，删除冗余位；在 AI 训练方面，其将 Gemini 架构中的核心矩阵乘内核加速 23%，使整体训练时间缩短 1%。

图：RSI下的AI具备递归反馈机制

图：DeepMind的AlphaEvolve利用AI生成并验证的算法被用于改善数据中心、芯片设计和AI训练流程

普通自动化：线性提效



AI研发：存在递归反馈通道



资料来源：Anthropic Institute、OpenAI、Google DeepMind，国信证券经济研究所整理

资料来源：Google DeepMind，国信证券经济研究所整理

3.从自动化AI研发走向RSI：局部闭环已形成，跨代递归仍待推进

3.1 AI已进入大模型研发多数环节，但自动化程度差异显著

- AI并非单一的“写代码”，而是一条从选择研究问题到训练继任模型的完整链条：文献与信息收集 → 提出研究假设 → 设计实验与评测方案 → 编写和调试代码 → 运行训练与实验 → 分析实验及评测结果 → 调整研究方向 → 修改训练方法或系统 → 训练和验证继任模型。
- 当前AI已经进入这条链条的多数环节，但不同环节的成熟度差异显著。整体呈现出“中间强、两端弱”的结构：越接近目标明确、结果可验证的实验执行，自动化程度越高；越接近研究方向选择、开放式结果判断和继任模型责任，越依赖人类。

图：RSI下的AI具备递归反馈机制



3.1 AI已进入大模型研发多数环节，但自动化程度差异显著

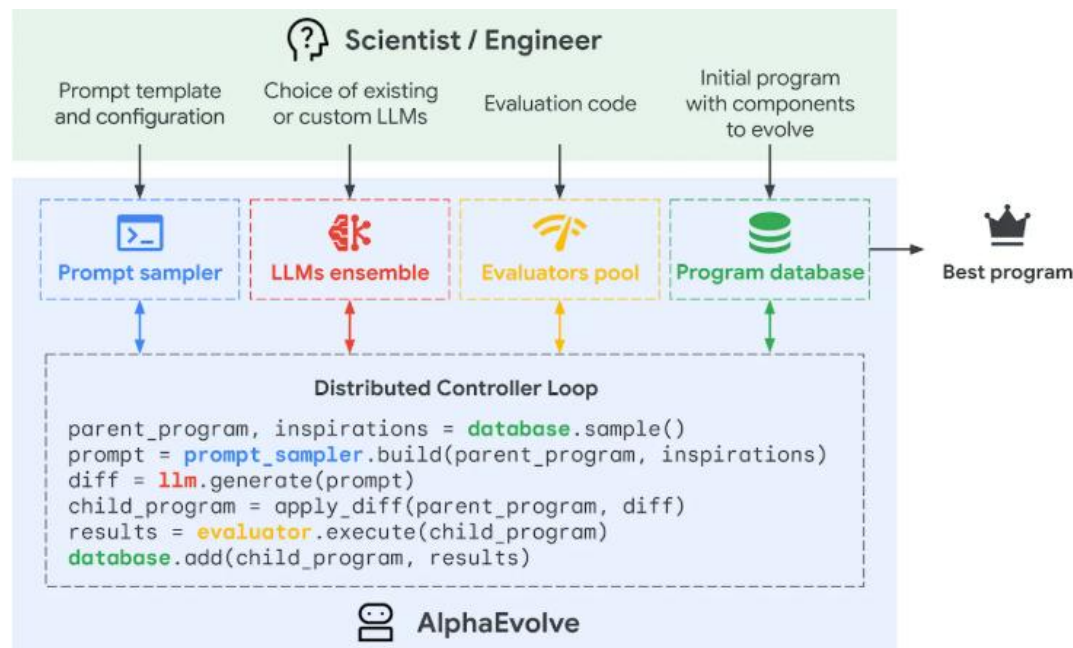
- **代码与实验执行相对成熟。**AI已能理解和修改真实代码库，编写功能、修复 Bug 和执行测试，配置实验环境并调用工具，运行实验、读取日志和处理部分异常，并参与 Kernel、推理和训练系统优化，以及根据明确 Evaluator 反复修改候选方案。
- **实验设计和结果分析正在突破。**AI 已经能够根据既定目标生成实验方案、设计代码框架、分析实验日志，并提出下一轮候选实验。The AI Scientist 甚至展示了提出想法、编写代码、运行实验、分析结果和撰写论文的自动科研流程。
- **研究方向和继任模型责任仍明显不足。**研发链条前端的主要障碍是选择真正值得研究的问题、提出具有长期价值的新方向、判断某项结果是实质性突破还是局部指标优化，以及决定继续投入、调整方向或终止研究。后端的主要缺陷则是判断小规模实验能否扩展至前沿模型，综合评估能力、安全性和泛化表现，独立完成前沿模型训练、验证并部署完整继任系统，以及对最终结果承担责任。。

图：Codex在独立云端沙箱中处理真实代码库任务



资料来源：OpenAI，国信证券经济研究所整理

图：Google DeepMind的AlphaEvolve

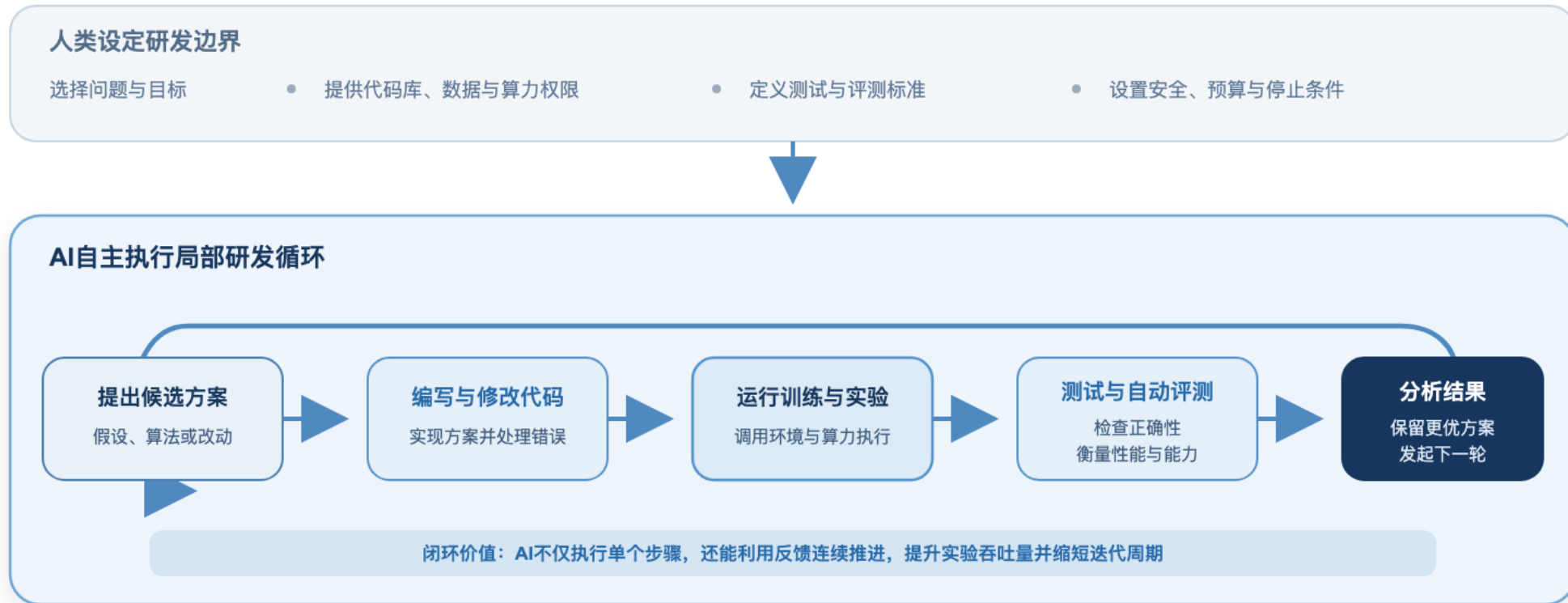


资料来源：Google DeepMind，国信证券经济研究所整理

3.2 AI 已从执行单个步骤，进入目标明确的局部研发闭环

- AI 参与大模型研发的早期形态，主要是提供代码建议、解释报错或整理资料，每完成一个步骤都需要研究人员重新下达指令。当前，在目标、环境和评价标准明确的任务中，AI 已经能够连续完成多个研发环节，并根据反馈自主发起下一轮迭代：提出候选方案 → 编写与修改代码 → 运行训练与实验 → 测试与自动评测 → 分析结果 → 保留更优方案 → 发起下一轮；
- 目前来看，AI 已不再只是完成孤立任务，而是开始利用测试和评测反馈连续推进研发工作。系统可以同时生成多个候选方案，更高频地运行实验，并减少人工干预次数，从而提高研发效率和迭代速度。

图：AI开始在明确目标下完成局部研发闭环

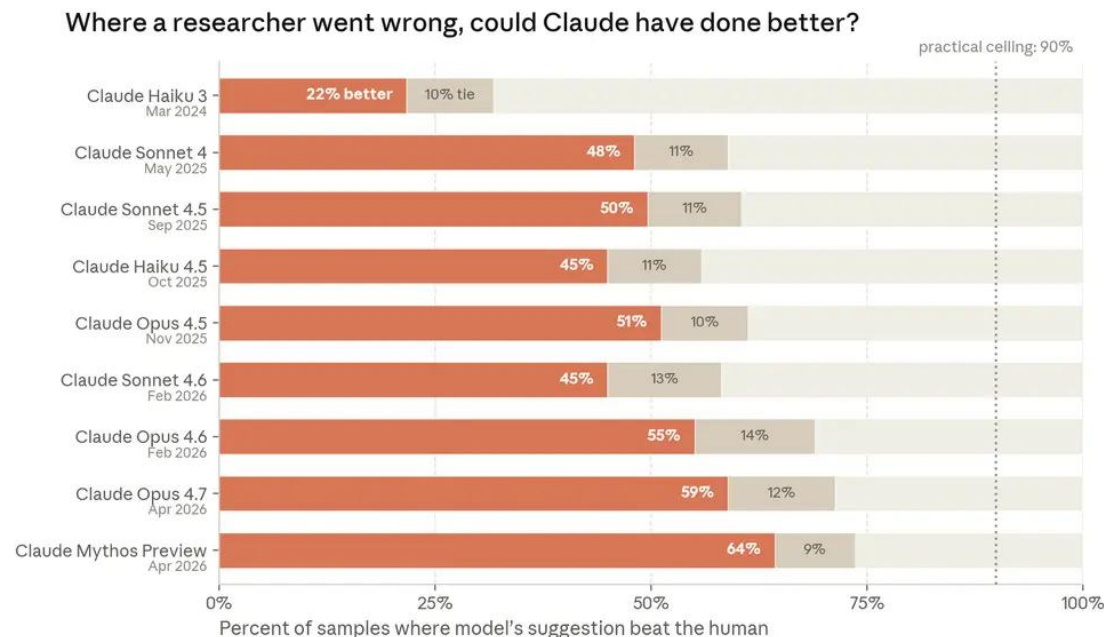


资料来源：OpenAI，国信证券经济研究所整理

3.2 AI已从执行单个步骤，进入目标明确的局部研发闭环

- **Evaluator 明确的领域率先形成闭环。** Google DeepMind 的 AlphaEvolve 利用大模型生成程序，通过自动 Evaluator 运行、验证和评分，再让表现更好的程序进入下一轮进化搜索，已被用于算法发现、数据中心调度、芯片设计和 AI 训练优化。
- **前沿实验室也开始将局部闭环用于真实 AI 研发。** Anthropic 披露，在目标和成功指标固定的实验优化任务中，Claude 可以反复修改训练代码、运行实验、计时并继续寻找更优方案，Anthropic 将其描述为一个“微型实验研究闭环”。

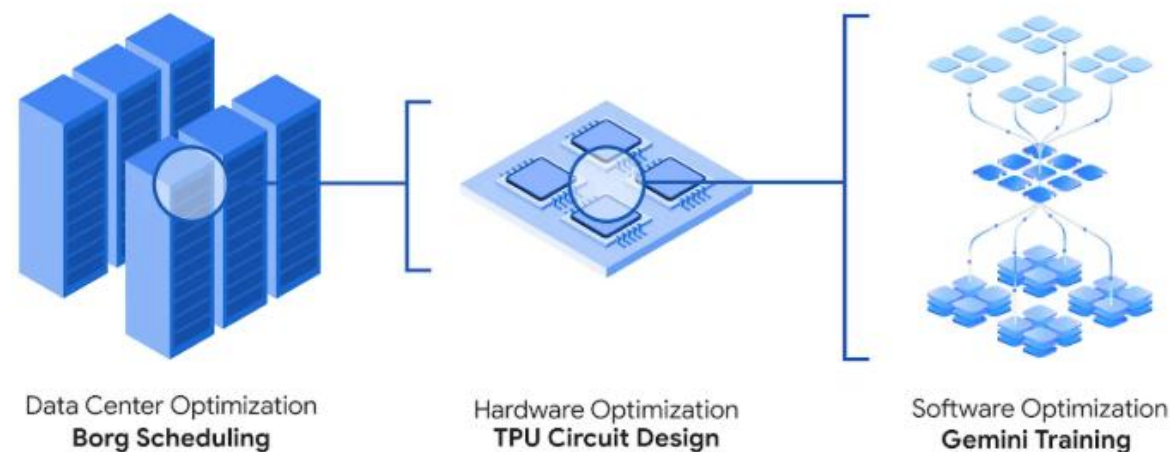
图：目标明确之下Claude寻找更优方案的能力持续提升



The study covers 129 internal Claude Code research sessions, each cut at a turn where the human's next direction had room for improvement. Model proposes an alternative; a judge that knows the session's eventual key findings picks the better move.

资料来源：Anthropic Institute，国信证券经济研究所整理

图：AlphaEvolve优化计算系统设计



资料来源：Google DeepMind，国信证券经济研究所整理

3.3 AI执行能力持续提高，长周期可靠性仍是自主迭代的核心约束

- 大模型和 Agent 已经能够完成越来越长、越来越复杂的研发任务，但“能够长时间运行”不等于“能够长时间保持正确”。随着任务步骤增加、环境状态变化和决策空间扩大，早期的局部错误可能被后续步骤继承并放大，最终使整个研发项目偏离目标。因此，限制 AI 独立负责完整模型研发的核心问题，正在从“能否执行任务”转向“能否长时间保持正确、及时发现错误，并在失败后可靠恢复”。
- 长周期可靠性是从“自动化研究执行者”走向“自主 AI 研究系统”的必要条件。当前可以通过检查点、测试 Oracle、持久记忆、沙箱、权限边界和人工审批缓解错误，但这些措施也说明，人类仍在承担系统监督和最终责任。在 AI 能够长期保持目标、发现错误、可靠恢复并通过独立验证之前，完整模型研发仍难以真正自主化

图：长周期可靠性仍待解决

Agent可承担的任务跨度正在扩展



长周期任务中的五类错误会相互叠加

1 规划与状态漂移

目标被逐步误解
上下文或记忆丢失
局部行动偏离主线

2 环境与工具错误

依赖、权限或配置异常
错误操作影响后续实验
恢复路径不稳定

3 验证覆盖不足

测试与评测不完整
错误结果被当作进步
小实验无法前沿扩展

4 多Agent不一致

任务分解与接口失配
重复工作或相互冲突
错误在协作中放大

5 人工验证瓶颈

Agent产出快于专家审查
人类难以复核全部路径
最终责任无法自动转移

3.4 算力、数据与基础平台决定自动化研发能否持续迭代扩展

- AI 可以快速提出候选方案、编写代码并发起下一轮实验，但训练、微调、强化学习和模型评测仍需要真实的数据、算力和运行时间。随着实验数量和 Agent 并行度提升，研发瓶颈将由方案生成逐步迁移至高质量数据、GPU 集群、实验调度、状态监控和失败恢复。自动化 AI 研发不意味着摆脱基础设施，反而对可复现、可编排的研发系统提出更高要求。
- 局部实验成功也不代表能够直接扩展至前沿模型：算法收益可能随训练规模变化，能力提升也可能伴随安全性、泛化表现或其他能力下降。因此，少量高成本的前沿训练和综合验证仍不可替代。自动化研发的总算力需求取决于两股相反力量：算法和工具效率提升会降低单次实验成本，而实验数量、Agent 并行度和评测需求扩大可能推高总负载，最终影响取决于效率提升与需求弹性的相对大小。

图：AI 研发循环加快后，瓶颈会向真实训练、资源调度和前沿验证迁移



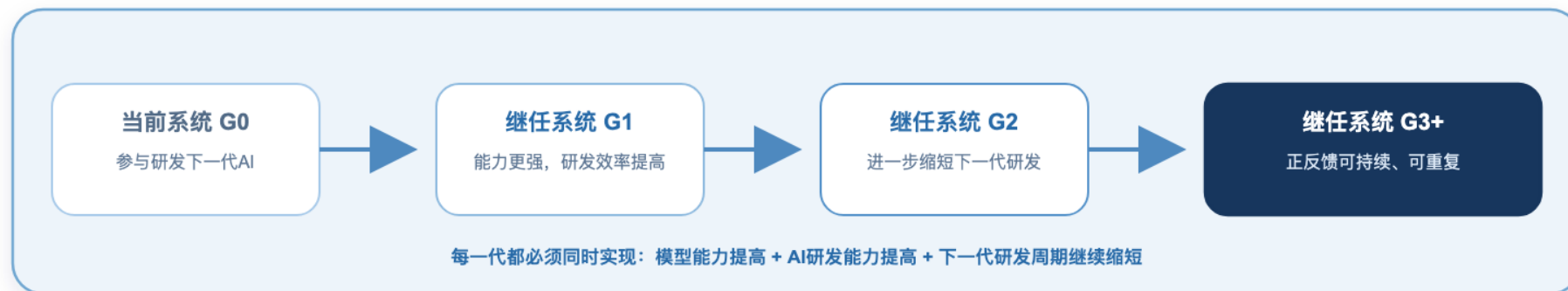
资料来源：Anthropic Institute、OpenAI、Google DeepMind，国信证券经济研究所整理

3.5 跨代、持续、可重复是RSI成立的核心判据

- 自动化 AI 研发已经能够提高代码、实验和系统优化效率，但一次研发提效、一项算法发现或一代模型升级，都不能证明 RSI 已经形成。严格 RSI 要求当前模型实质性参与继任系统研发，继任系统在独立评测中确实更强，并进一步扩大 AI 自主研发范围、提高有效实验吞吐量、缩短下一代模型研发周期；上述反馈还必须跨越多代，持续、稳定且可重复地发生。
- 递归链条可能在任何一环中断。易优化问题被优先解决后，研发提效可能快速边际递减；不完整的 Evaluator、合成数据偏差和错误代码可能跨代累积；算力、数据、训练时间和专家验证也会形成现实约束。随着继任系统能力增强，理解、验证和控制难度还会提高。因此，当前虽已出现局部闭环和单代提效，但多代递归正反馈仍需要更长时间验证。

图：RSI需要经历跨越多代持续递归验证

严格RSI需要连续通过四代检验



递归链条可能在任何一处中断



4. 前沿模型厂商的RSI路径：从研发提效到继任模型闭环

4.1 三家路线各异，均在推进AI研发自动化

- OpenAI、Anthropic与Google DeepMind虽然没有采用完全一致的RSI定义，但推进方向高度一致：让当前模型承担越来越多下一代AI系统的研发工作。OpenAI侧重将AI Self-improvement拆解为可评测、可治理的能力，并通过Codex和Research Agent进入代码、调试、Kernel及训练研究任务；Anthropic强调研发工作由人类逐步转移给AI，并以内部代码贡献、工程师生产率和长周期任务表现提供更加贴近真实研发流程的证据。Google DeepMind更侧重Evaluator驱动算法进化，AlphaEvolve将大模型生成、自动验证和进化搜索相结合，已用于算法、数据中心、芯片设计及AI训练优化。
- 三条路线分别代表评测与平台化、内部研发自动化和局部反馈闭环，当前仍主要停留在研发提效和局部系统改进，尚未完成完整继任模型，或形成跨越多代、持续可重复的递归加速。

图：OpenAI、Anthropic、Google DeepMind在RSI的路线及特点



资料来源：Anthropic Institute、OpenAI、Google DeepMind，国信证券经济研究所整理

4.2 OpenAI：将RSI拆成可度量能力，并设置明确的High与Critical门槛

- OpenAI将AI Self-improvement定义为AI加速AI研究、并可能提高自身能力的能力类别，而不是模型反思、代码生成或单项Benchmark提升。在Preparedness Framework v2中，High 门槛要求模型带来的影响，相当于为每位 OpenAI 研究员配备一名高绩效的中级研究工程师助手，比较基准为研究员 2024 年的工作状态。其判断对象已从单点能力上升至整个研发组织的持续生产效率提升
- Critical门槛则更接近严格RSI：系统要么成为超人类Research Scientist Agent，要么将一代模型进步所需时间压缩至2024年基准的五分之一，并持续数月。OpenAI将其称为递归自我改进或全自动AI研发。

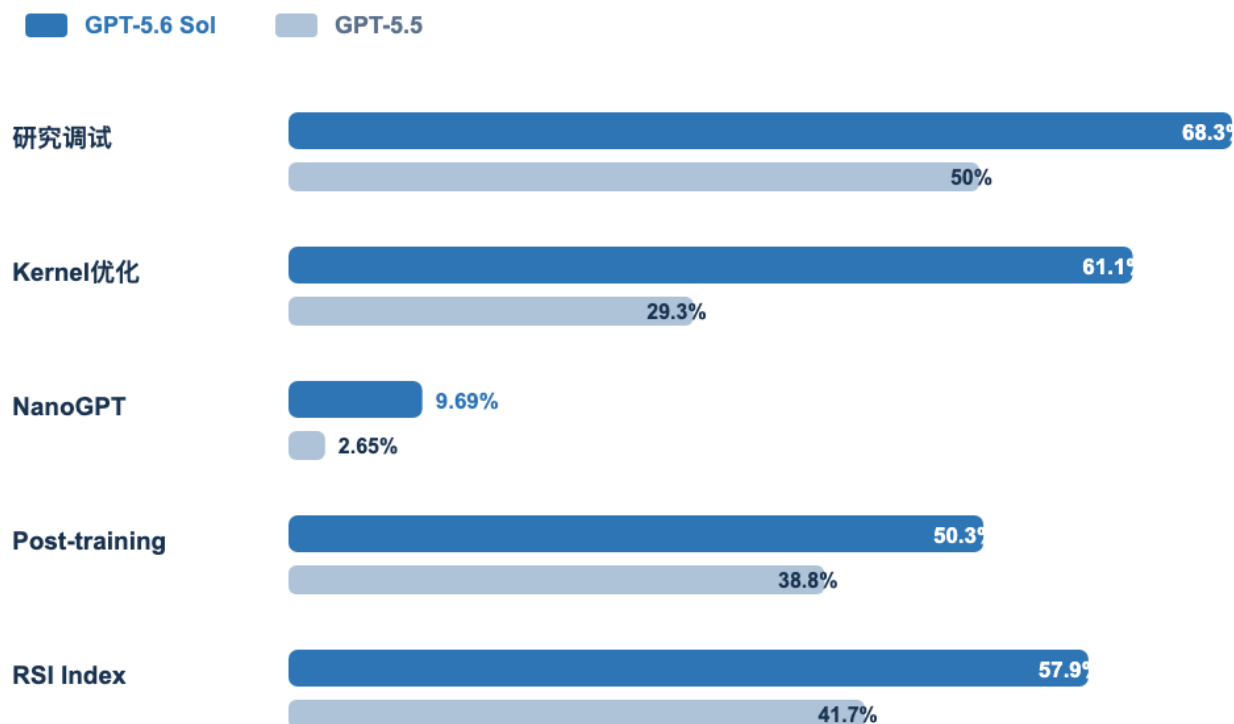
图：OpenAI在RSI的路线及特点



4.3 OpenAI：模型已进入完整研发链条，但能力仍呈现“工程强、训练弱”

- OpenAI 已开始用真实AI研发任务直接评估模型。GPT-5.6 Sol在内部研究调试评测上达到68.3%，高于GPT-5.5的50.0%；KernelGen 1P达到61.1%，较29.3%显著提高，说明模型在真实代码库、实验排错和底层性能工程中已经形成较强能力。PostTrainBench Lite也从38.8%提高至50.3%，模型开始能够在受限环境下选择数据、训练后方法和RL反馈循环。
- 但越接近模型训练本身，差距越明显。GPT-5.6 Sol在NanoGPT上仅为9.69%，远低于系统卡披露的当前人类最优水平72.38%，且小模型优化不能直接外推到前沿训练。内部Coding Inference研究算力六个月增长100倍、Agent Token增长约22倍，表明虽然使用正在迅速普及，但当前OpenAI更像是在建立自动化AI研发的评测与执行基础，尚未能够独立完成前沿继任模型。

图：GPT5.6相比5.5在RSI上提升显著



评测覆盖的研发环节

- 研究调试
定位真实实验Bug
- KernelGen
低层性能工程
- NanoGPT
小模型预训练优化
- PostTrainBench
训练后与RL配方
- RSI Index
综合研发能力

4.4 Anthropic：研发工作持续向AI转移，RSI呈渐进式路径

- Anthropic 认为，RSI 更可能是一条渐进发生的研发自动化曲线，而不是某一天突然出现的能力跃迁。其内部研发已从人类完成全部代码和实验，经历 Chatbot 辅助局部任务、Coding Agent 直接修改并运行代码，正在走向可以自主工作数小时、调用工具和委派其他 Agent 的阶段。沿着这一路径继续推进，未来模型才可能自主设计和训练继任系统，使 Claude 开始持续改进 Claude。
- 当前 Claude 已经能够在目标清晰的实验中找到或超过熟练人类的执行水平，工程任务也逐步形成“人类给目标、AI 决定方法”的分工。但高等级研究的核心，是判断什么问题值得解决、如何解释异常结果，以及局部发现能否扩展至前沿模型；Claude 在目标选择和 Research Taste 上仍有明显差距。因此，Anthropic 判断 AI 研发加速已经发生，但完整 RSI 并非必然，方向判断、算力、验证和组织吸收能力仍可能限制闭环速度。

图：Anthropic对RSI的看法及路线图



资料来源：Anthropic Institute，国信证券经济研究所整理

4.5 Anthropic：内部生产率提升明显，方向选择仍然依靠人工

- **AI 正在快速提高 Anthropic 的 AI 研发效率。**截至 2026 年 5 月，超过 80% 的合并代码由 Claude 编写；2026 年第二季度，典型工程师日均合并代码约为 2024 年的 8 倍。在固定目标的小模型训练代码优化中，Claude 从 2025 年 5 月的约 3 倍加速，提升至 2026 年 4 月的约 52 倍，显示其在目标明确、评价可自动执行的实验闭环中已达到超人类执行水平。
- **开放研究也开始出现突破。**Claude Agent 累计运行约 800 小时，恢复了弱强监督表现差距的 97%，而两名人类研究员约一周恢复 23%；在特意挑选的人类研究走偏节点上，模型的下一步判断优于人类当时选择的比例升至 64%。但人类仍负责选题和评分，开放研究结果也未直接迁移至生产规模。因此，Anthropic 已经证明 AI 能够显著提高研发效率，但模型尚未掌握 Research Taste、前沿扩展和继任模型的方向选择能力。

图：Anthropic内部生产率提升显著

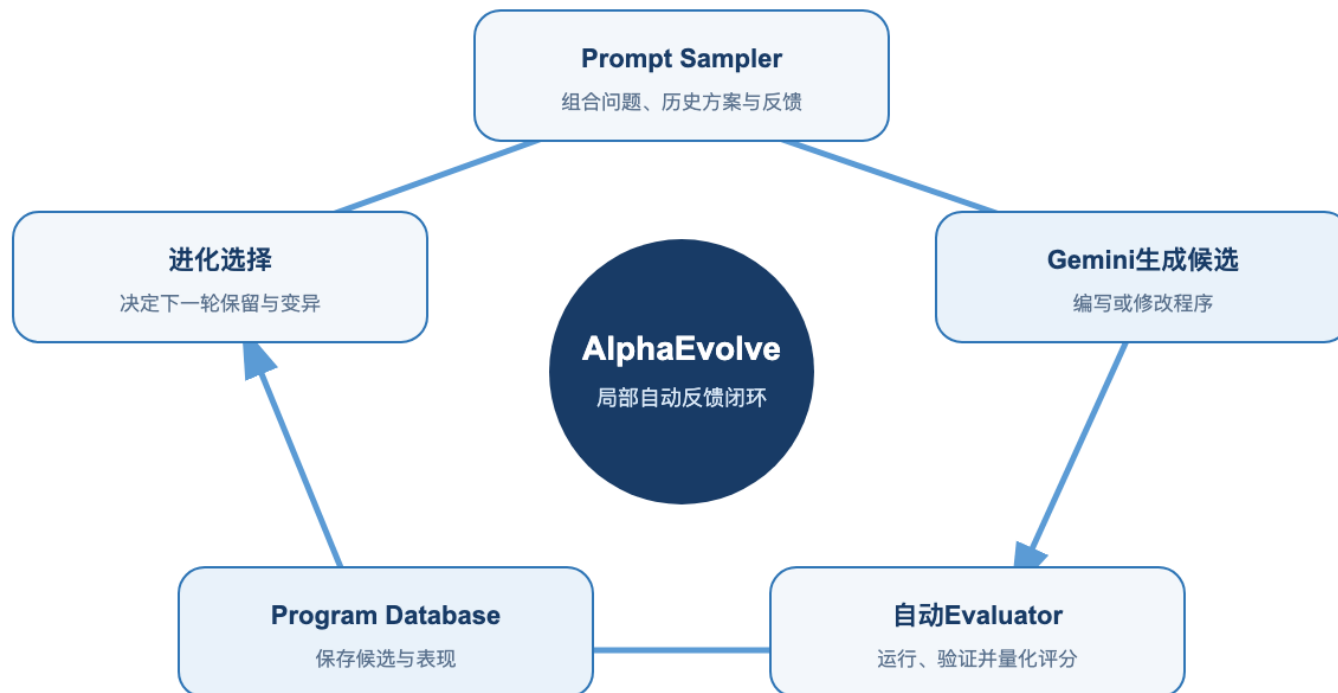


4.6 Google DeepMind: Evaluator驱动“生成—验证—进化”局部闭环



- **Google DeepMind 的 AlphaEvolve 代表 Evaluator 驱动的自动化研发路线。**系统先由 Prompt Sampler 组合问题、历史候选方案和已有反馈，再由 Gemini 生成或修改程序；候选程序被真实运行，由自动 Evaluator 验证正确性并量化评分，结果进入 Program Database，进化算法据此选择下一轮的保留和变异对象。反馈由此直接进入下一轮生成，使系统从一次性写代码升级为持续搜索和算法进化。
- **这条路线在数学、Kernel 和系统优化中率先有效，核心原因是结果可以通过正确性、速度和资源占用快速验证。**从 RSI 角度看，AlphaEvolve 已经具备“AI 提出改进—机器验证—反馈促进下一轮搜索”的局部递归结构，部分产出也反过来改善 AI 训练基础设施。但问题设定、程序骨架、Evaluator 和部署标准仍由人类建立，优化对象也主要是局部算法，而非完整继任模型。因此，它仍属于自动化 AI 研发的重要前置形态。

图： AlphaEvolve的局部自动反馈闭环



资料来源：Google DeepMind，国信证券经济研究所整理

请务必阅读正文之后的免责声明及其项下所有内容

4.7 Google DeepMind：成果已进入生产系统，但仍是受约束的局部进化

- AlphaEvolve 的价值在于其成果已经进入 Google 的真实生产系统。其 Borg 调度算法已运行超过一年，平均持续回收 Google 全球 0.7% 的计算资源；对 Gemini 关键矩阵乘 Kernel 的优化实现 23% 提速，并使训练时间下降约 1%；FlashAttention 底层 Kernel 最高提速 32.5%。系统还提出被整合进下一代 TPU 的电路修改，并在两天内完成过去需要数月密集人力的 Cache 策略优化。
- 这些案例表明，AI 已经能够通过大规模自动实验发现可部署方案，并反过来改善算力、芯片和模型训练效率，已具备研发正反馈的局部形态。但成功任务普遍具有清晰目标、严格正确性约束和自动 Evaluator，生产部署仍需要人类验证。AlphaEvolve 证明的是受约束空间内的算法和系统进化，而不是 AI 自主选择研究方向、设计完整模型并承担前沿训练责任；两者之间仍有明显距离。

图：AlphaEvolve的局部自动反馈闭环



为什么这些结果重要

- 改进已经进入数据中心、TPU和Gemini训练等真实生产系统
- 自动实验把部分专家级优化从数周压缩到数天，扩大搜索空间
- 产出反过来改善AI基础设施与训练效率，具备“研发正反馈”的局部形态

4.8 横向比较：路径各异，均指向更高水平AI研发自动化

- 三家公司的路线分别代表自动化AI研发的不同组成部分。OpenAI侧重能力评测和治理阈值，通过真实研究调试、Kernel、训练和Post-training任务判断AI Self-improvement；Anthropic侧重内部研发工作向AI转移，以代码占比、工程师生产率、实验优化和开放研究衡量真实影响；Google DeepMind侧重Evaluator驱动算法进化，把生成、验证和选择连接成可部署的局部闭环。
- 目前还无法判断哪个方向“最接近RSI”。OpenAI最该跟踪是否触及High或Critical门槛；Anthropic要观察AI能否从执行实验走向自主选题和前沿扩展；Google要观察局部Evaluator闭环能否扩展到模型架构和训练。三家公司都已进入自动化AI研发阶段，但尚未实现AI独立完成完整继任模型、摆脱人类方向判断，或形成跨越多代、持续可重复的递归加速。

图： OpenAI、Anthropic、Google DeepMind 横向比较

比较维度	OpenAI	Anthropic	Google DeepMind
核心口径	可评测的自我改进能力	研发工作持续向AI转移	Evaluator驱动算法进化
主要路线	真实研发评测 + Agent平台	Claude Code + 内部虚拟实验室	AlphaEvolve生成—验证—选择
最强证据	RSI Index与内部采用量	代码占比、生产率与开放研究	生产系统中的可部署优化
当前短板	未证明超人类研究科学家	目标选择与研究判断不足	依赖可自动验证的问题
最该跟踪	是否触及High/Critical 阈值	是否接管方向与规模化训练	是否扩展到模型架构与训练闭环

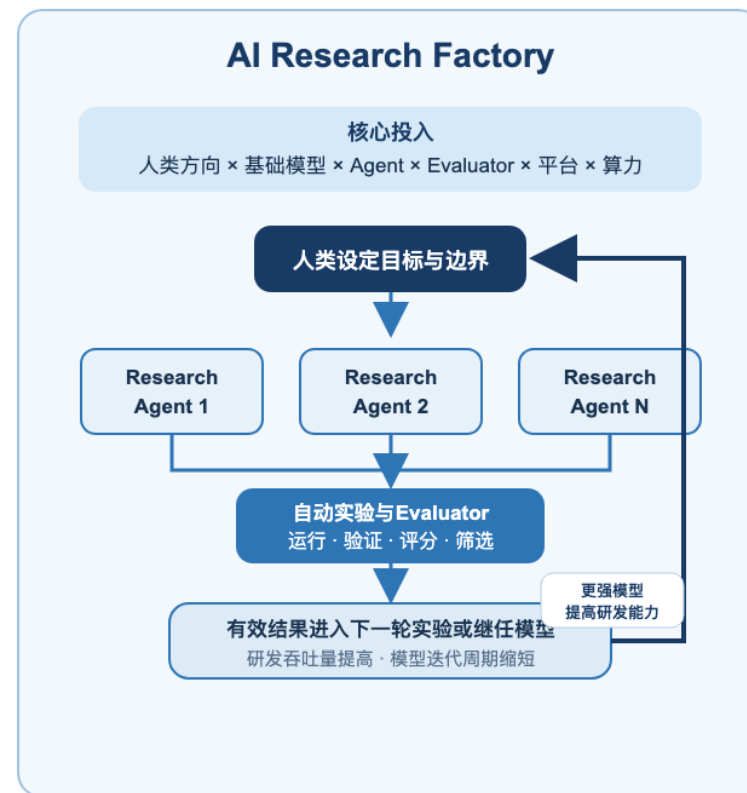
资料来源：Anthropic Institute、OpenAI、Google DeepMind，国信证券经济研究所整理

5. 产业影响与投资观点：AI研发加速下的模型、算力与应用机会

5.1 RSI有望重构AI研发范式，竞争从单代模型转向AI Research Factory

- 传统大模型研发主要依赖人类研究员提出假设、设计实验和判断结果，再调用数据与算力完成训练。随着Coding Agent向Research Agent演进，模型已经开始承担代码、排错、Kernel优化、实验执行和部分结果分析；自动Evaluator and 实验平台则可以筛选结果并反馈下一轮研究。AI从模型研发的最终产出，转变为生产下一代模型的研发工具。
- 未来竞争将从单代模型能力竞争，转变为**基础模型、Research Agent、Evaluator、数据、训练平台和算力**共同组成的**AI Research Factory**之争。人类仍负责研究方向、价值判断和安全监督，AI则扩大实验并行度和研发体量和效率。更强模型可能带来更高研发生产率，再支持更多实验和更快迭代

图：传统大模型研发 vs AI Research Factory



5.2 RSI将推动厂商扩大基础模型与AI研发计算投入

- RSI尚未实现，但只要AI能够缩短代码开发、实验排错和训练验证时间，前沿厂商就有动力提前建设自动化研发体系。基础模型预训练仍是能力底座，更强的Coding、Reasoning和Agent能力决定AI能够承担哪些研发任务；与此同时，多Agent并行执行代码、实验和结果分析，将使Research Inference成为研发过程中的新计算负载。
- AI提出的候选方案还需经过Evaluator、Scaling实验、RL、Post-training和前沿安全评估，才能转化为真实继任模型。OpenAI披露，六个月内内部Coding Inference研究算力增长100倍，Agent Token约增长22倍，说明研发推理正在快速普及。RSI带来的不是单一预训练增长，而是整个AI研发计算栈扩张。

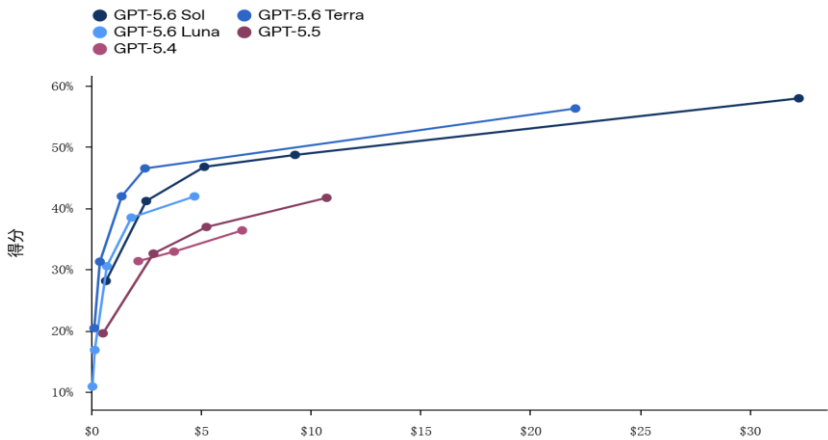
图：RSI推动整个AI研发计算栈扩张



资料来源：Anthropic Institute、OpenAI、Google DeepMind，国信证券经济研究所整理

请务必阅读正文之后的免责声明及其项下所有内容

图：GPT-5.6 SoI 较 GPT-5.5 RSI提升了 16.2 个点，全面加速了内部研究进程



资料来源：OpenAI，国信证券经济研究所整理

5.3 开源模型迈入万亿参数时代，Scaling仍在推高算、存、互联需求

- 2026年7—8月，开源模型快速进入万亿参数时代：Kimi K3达到2.8万亿总参数、约1042亿激活参数；Qwen3.8旗舰模型达到2.4万亿、约950亿激活参数；DeepSeek-V4-Pro达到1.6万亿、490亿激活参数；Inkling达到9750亿、410亿激活参数。新模型普遍面向Coding、长周期Agent、多模态和100万Token上下文，说明Scaling仍是能力竞争的重要底座。
- MoE稀疏激活使总参数增长不等于单Token计算同比例增长，但更大权重仍会增加存储、专家通信和集群部署压力，长上下文提高KV Cache需求，大规模RL进一步增加训练推理负载。Inkling预训练45万亿Token，并运行超过3000万次RL Rollout。模型规模扩大并非走向RSI，但即使RSI尚未实现，参数扩大也正推高算、存、互联等基础设施投入

图： 开源模型进入万亿参数

模型	总参数	激活参数	关键特性	发布时间
Kimi K3	2.8T	104.2B	原生多模态、1M上下文，面向长周期Coding、知识工作与深度推理	2026.07
Qwen3.8-2.4T-A95B	2.4T	约95B	旗舰MoE模型，面向编程、办公、科研等复杂Agent任务	2026.08
DeepSeek-V4-Pro	1.6T	49B	1M上下文，强化Agentic Coding、推理与长上下文效率；8月13日正式版上线	2026.08
Inkling	975B	41B	文本、图像、音频原生推理，1M上下文；预训练45T Token，RL Rollout超过3,000万次	2026.07
DeepSeek-V4-Flash	284B	13B	1M上下文，以更低激活规模提供高效率Agent推理；正式API版本延续同一模型架构	2026.07

模型迭代仍在加快

DeepSeek V4-Pro于8月13日进入正式发布；从4月Preview到正式版仅约四个月，训练后迭代持续改善Agent能力。

能力边界也开始影响开放节奏

GLM-5.3于8月14日发布，但因网络安全能力评估，权重开放延后两周；参数尚未披露，因此不进入上表比较。

资料来源： 千问、月之暗面、deepseek、linkling，国信证券经济研究所整理

5.4 参数、Token与代际频率共同放大，RSI或显著提升训练端算力需求

- RSI对训练算力的影响不只来自模型参数扩大。前沿预训练需求由激活参数、训练Token和模型代际频率共同决定，并由算法与硬件效率部分抵消。更强模型需要更大训练容量，AI还可以生成合成数据、构造Agent轨迹并提高训练频率，使“单代模型更大”和“单位时间内训练更多代模型”形成乘数效应。
- 简单情景测算显示，RSI 形成前的 AI 研发自动化加速期，前沿预训练需求可能达到当前的约 3 倍，叠加候选实验和 RL 后，训练端总需求约为 4—6 倍；跨代正反馈初步形成时，训练端总需求可能达到 15—30 倍；若模型代际周期显著缩短，递归加速期的长期需求潜力可达 50—100 倍以上。

图：RSI或显著提升训练端算力需求



资料来源：Anthropic Institute、OpenAI、Google DeepMind，国信证券经济研究所整理

5.4 参数、Token与代际频率共同放大，RSI或显著提升训练端算力需求

➢ 计算公式：年度前沿预训练倍数 = 激活参数倍数 × 训练Token倍数 × 年度模型代际倍数 ÷ 综合训练效率倍数

图： RSI训练端算力需求综合情景测算

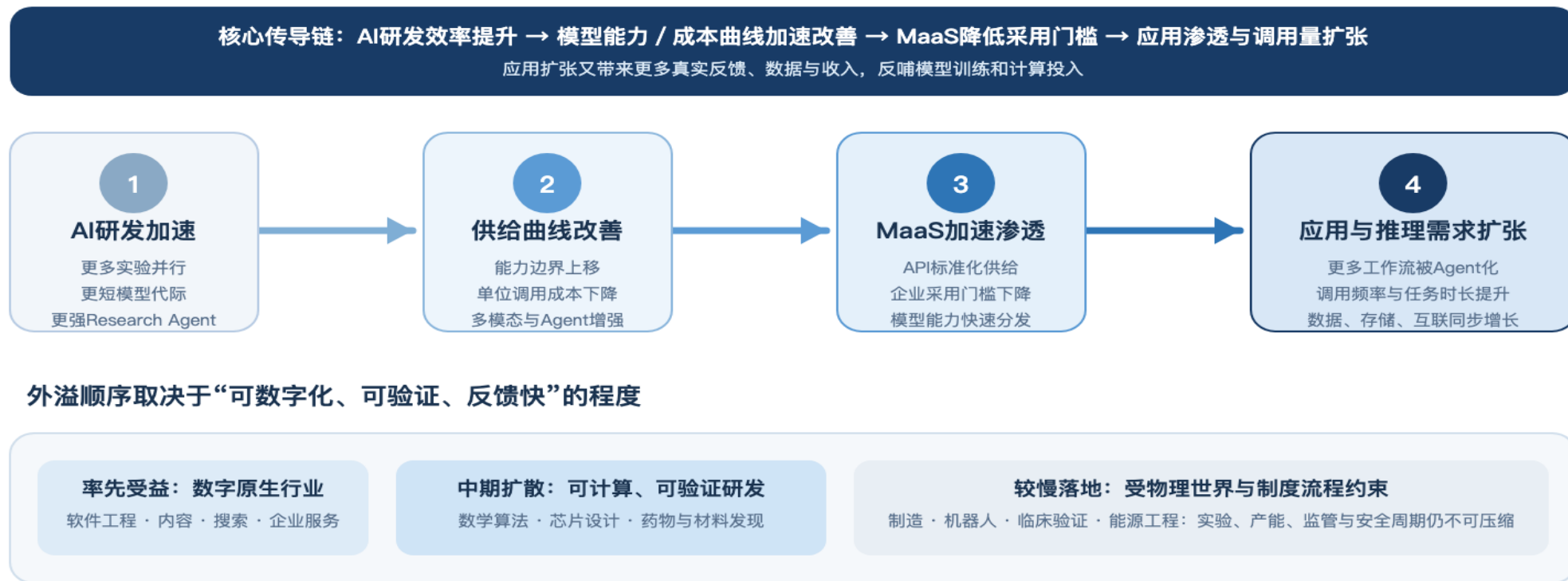
情景	激活参数倍数	训练Token倍数	年度模型代际倍数	综合训练效率倍数	前沿预训练倍数	未来额外训练负载 / 正式预训练	训练端总需求倍数
AI研发自动化加速期—低值	2.0x	1.5x	2.0x	2.0x	3.0x	100%	4.0x
AI研发自动化加速期—高值	2.0x	1.5x	2.0x	2.0x	3.0x	200%	6.0x
跨代正反馈形成期—低值	4.0x	2.5x	4.0x	3.0x	13.3x	75%	15.6x
跨代正反馈形成期—基准	4.0x	2.5x	4.0x	3.0x	13.3x	150%	22.2x
跨代正反馈形成期—高值	4.0x	2.5x	4.0x	3.0x	13.3x	250%	31.1x
递归加速期—低值	8.0x	4.0x	8.0x	5.0x	51.2x	50%	51.2x
递归加速期—基准	8.0x	4.0x	8.0x	5.0x	51.2x	150%	85.3x
递归加速期—高值	8.0x	4.0x	8.0x	5.0x	51.2x	200%	102.4x

资料来源：Anthropic Institute、OpenAI、Google DeepMind，国信证券经济研究所整理

5.5 RSI若形成研发正反馈，AI渗透将从“工具普及”走向“能力供给加速”

- RSI对应用端的影响不只是诞生一个更强模型，而是模型能力与成本曲线的改善速度可能显著加快。AI研发效率提高后，同样时间内可完成更多实验和模型迭代；头部厂商再通过MaaS、API和Agent平台，把最新推理、多模态、工具调用及长周期执行能力快速分发给企业，客户无需自行承担基础模型训练成本。企业采用AI的重点将逐步从“是否部署模型”转向“哪些 workflows 应优先被Agent重构”。
- 应用渗透会进一步放大推理、数据、存储和互联需求。Agent需要多轮推理、工具调用、检索、代码执行及验证，单任务计算量通常高于普通问答。软件、内容、搜索和企业服务等数字原生行业有望率先受益；芯片、药物和材料研发可在候选发现环节加速；制造、机器人和临床验证仍受技术水平、物理世界产能/成本、监管等因素约束。

图：RSI推动AI“工具普及”升级为“能力供给加速”



资料来源：Anthropic Institute、OpenAI、Google DeepMind，国信证券经济研究所整理

5.6 投资观点：AI研发自动化加速，关注模型、算力硬件、CSP、与高价值应用

- **自动化AI研发加速，RSI正由理论进入落地前夕的阶段。**1) 完整RSI尚未实现，但AI已经进入代码编写、研究调试、Kernel优化、自动实验和部分训练后环节，局部研发反馈闭环正在形成。2) 从前沿模型厂商落地来看，模型、Agent、Evaluator、训练平台和算力基础设施开始被整合进同一套研发系统和工作流，AI逐步由研发辅助工具转向研发流程的重要参与者和工具。
- **前沿模型厂商：模型竞争由单代能力比拼转向AI Research Factory竞争，研发正反馈与平台化壁垒有望强化。**1) 从模型竞争来看，能力更强的模型既是争取客户的重要因素，同时也正逐渐成为研发下一代模型的生产工具；远期来看，模型能力如果能够持续提高Research Agent和研究人员产出，领先厂商可能形成“更强模型—更高研发效率—更快推出继任模型”的正反馈闭环。2) 同时从模型厂商竞争要素来看，也将由人才、数据和单次训练扩展至模型、Agent、Evaluator、数据闭环、训练平台和算力的系统能力。
- **算力基础设施：训练与Research Inference共同扩张，AI研发自动化有望进一步打开算力产业链需求空间。**1) 训练端来看，算力需求不再仅取决于单次前沿预训练规模，RSI落地渗透之下，参数规模、训练Token、模型代际频率、候选实验、Post-training、RL Rollout、验证及研发推理均可能同步增加，更快迭代和更高实验并行度可能形成对训练需求的乘数效应，带动算力需求的天花板持续提升。2) 重点关注计算芯片及ASIC、HBM与存储、先进封装、高速互联与光通信、交换设备、液冷、电力及数据中心等供给壁垒较高的基础设施环节；3) 同时看好具备大规模算力基础设施和持续资本开支能力的头部CSP，有望受益云端AI调用量提升。
- **高价值应用：模型能力加速通过MaaS、API与Agent平台向应用端外溢，优先关注可验证、可闭环的高价值场景。**1) 模型能力持续提升、推理成本下降以及多模型接入门槛降低，将持续降低AI应用开发、试错和规模部署成本，并通过需求弹性提升调用频次、上下文长度和任务复杂度从而促进应用场景的扩展与渗透率提升。2) 软件工程、研发工具、企业知识管理、客户服务和数字内容等领域数字化程度高、结果可衡量、反馈周期短，有望率先形成Agent规模化渗透；药物、材料和芯片设计等研发场景亦具备较大的效率提升潜力。3) 应用公司的核心壁垒将由“能否接入模型”转向专有数据、核心 workflows、可靠Evaluator以及持续提高客户产出的能力。
- **风险提示：完整RSI可能长期无法形成；商业化落地低于预期风险；竞争加剧风险；监管政策风险；技术进步低于预期风险等。**

5.7 附录：大模型产业链相关公司



公司名称	股票代码	上市地点	所属国家	核心大模型/产品	商业模式与业务特点
阿里巴巴	9988.HK	港股/美股	中国	通义千问（Qwen）	商业友好开源，深度融入阿里云生态，高频迭代，为多语言开发主流选择之一。
腾讯	0700.HK	港股	中国	混元大模型	整合微信、游戏、广告、办公等生态，推出AI编程工具等产品。
科大讯飞	002230.SZ	A股	中国	星火大模型	深耕教育、医疗等垂直领域
三六零	601360.SH	A股	中国	360智脑4.0	自研千亿参数通用模型，覆盖全应用场景，在税务、企业服务等行业落地
昆仑万维	300418.SZ	A股	中国	天工大模型、Mureka V8	开放API按Token调用计费，切入“词元经济”变现；同时布局AI音乐、漫剧等领域
智谱华章	02513.HK	港股	中国	GLM系列大模型	“模型即产品”平台路线，以MaaS（API调用）为核心，B端/G端服务为主
MiniMax	00100.HK	港股	中国	M2系列大模型、海螺AI、星野、Talkie	C端订阅+B端开放平台双轮驱动，海外收入占比过半。
字节跳动	未上市	—	中国	豆包大模型（2.1系列）、Seed 2.0/2.1系列、Seedance视频生成（2.5）、Seedream图像模型、TRA E企业级AI Coding、扣子Agent平台	B端+C端双轮驱动：①C端以豆包App为核心（日活超2亿），2026年6月上线付费订阅（68/200/500元/月三档）②B端通过火山引擎MaaS平台提供API调用，2026年MaaS营收目标150亿元（2025年约15亿元）；③视频生成模型Seedance贡献主要企业收入。国内公有云大模型Token调用份额49.5%位列第一，日均Token调用量达180万亿。
月之暗面（Kimi）	未上市（筹备赴港IPO）	—	中国	Kimi K3	2026年7月发布全球参数最大开源模型Kimi K3，在Code Arena超越Claude Fable 5和GPT-5.6 So；已调整架构，预计6个月内完成港股上市。
深度求索（DeepSeek）	未上市	—	中国	DeepSeek系列大模型	国产开源模型龙头，Token调用量稳居前列。
微软（Microsoft）	MSFT	纳斯达克	美国	Copilot、Azure AI（整合OpenAI技术）	向OpenAI投资超130亿美元，将GPT系列深度整合至全线产品；Azure云提供大模型API服务。
Alphabet（谷歌）	GOOGL / GOOG	纳斯达克	美国	Gemini系列大模型	自研全栈AI模型，云平台（GCP）提供AI基础设施；持有Anthropic约14%股权。
亚马逊（Amazon）	AMZN	纳斯达克	美国	AWS Bedrock、Titan大模型	AWS云平台提供模型托管与分发服务，持有Anthropic股权。
Meta Platforms	META	纳斯达克	美国	Llama系列大模型	自研开源大模型，深度整合至广告推荐与内容生成，为GPU综合应用大户。
Anthropic（Claude）	未上市（已保密递交IPO申请）	—	美国	Claude系列大模型	2026年6月递交IPO申请，估值达9650亿美元，最早2026年10月挂牌；Alphabet（14%）、Amazon、Microsoft参股。
OpenAI（ChatGPT）	未上市	—	美国	GPT系列、ChatGPT、Sora	全球估值最高AI初创，无公开交易股票；微软为最大战略投资者。

资料来源：Wind，公司公告，官网，国信证券经济研究所整理

5.7 附录：中美主要CSP情况



公司名称	股票代码	云业务品牌	2026年预估资本支出	业务结构
亚马逊 Amazon	AMZN.O	AWS	约2,200亿美元	综合型公有云；重点投入AWS、AI基础设施、数据中心及自研芯片。
Alphabet	GOOGL.O / GOOG.O	Google Cloud	1,950亿–2,050亿美元	TPU、Gemini、Vertex AI与公有云构成全栈体系。
微软 Microsoft	MSFT.O	Microsoft Azure	约1,750亿美元	Azure结合Microsoft 365、GitHub和Copilot；部分租赁分类影响报告口径。
Meta Platforms	META.O	自建AI基础设施	1,300亿–1,450亿美元	主要支持Llama、推荐广告系统、Meta AI及新型AI产品。
甲骨文 Oracle	ORCL.N	OCI	约600亿–700亿美元	聚焦大规模GPU集群、AI训练推理、数据库云和多云部署。
阿里巴巴 Alibaba	9988.HK / BABA.N	阿里云	约1,300亿–1,600亿元人民币	云基础设施、通义千问与AI开发平台一体化。
腾讯控股 Tencent	0700.HK	腾讯云	约1,100亿–1,300亿元人民币	结合微信、游戏、广告和企业服务，发展混元模型与GPU云。
字节跳动 ByteDance	未上市	火山引擎	约2,000亿元人民币；潜在上修至4,000亿–5,000亿元	服务抖音、TikTok、豆包及推荐系统，并向外部提供云和模型服务。
百度 Baidu	9888.HK / BIDU.O	百度智能云	约220亿–280亿元人民币	重点布局GPU云、文心模型、搜索与自动驾驶相关AI应用。
华为 Huawei	未上市	华为云	未单独披露	昇腾、盘古、华为云、存储及网络设备构成软硬件全栈体系。

资料来源：Wind，公司公告，官网，国信证券经济研究所整理

请务必阅读正文之后的免责声明及其项下所有内容

5.7 附录：美股B端场景垂类AI应用



公司	股票代码	市值（亿美元）	定位/产品	AI逻辑/商业化方式
Snowflake	SNOW	1098	AI Data Cloud、Cortex、Snowflake Intelligence、Cortex Code	企业数据治理、RAG和Agent运行底座；按数据存储、计算和AI用量收费，同时向业务用户延伸AI应用
Datadog	DDOG	1008	云监控、LLM/Agent可观测性、安全监控	监测模型调用、延迟、Token成本、Agent链路与生产故障；按主机、日志、事件和用量计费
Cloudflare	NET	1036	Workers AI、边缘推理、AI Gateway、网络与安全	将推理、缓存、安全和Agent运行放到边缘节点；按请求、算力、存储和安全服务收费
MongoDB	MDB	303	Atlas、向量检索、AI应用数据库	为RAG、AI应用和实时业务系统提供数据库与检索层；按云数据库用量收费
Elastic	ESTC	72	Elasticsearch、向量搜索、企业搜索、可观测性	企业知识检索、RAG、安全分析和日志智能化；订阅+云用量模式
Palantir	PLTR	3807	Foundry、AIP、Ontology、企业AI操作平台	将模型连接企业数据、权限与实际业务动作；平台许可、项目部署和持续扩容
Microsoft	MSFT	36197	Azure AI、Microsoft Foundry、Microsoft 365 Copilot、GitHub Copilot	同时覆盖云、模型、开发工具和办公入口；云用量+Copilot席位+Agent/平台调用收费
Salesforce	CRM	1581	Agentforce、Data 360、Slack Agent、AI CRM	以CRM数据和销售/客服 workflow 驱动Agent执行；订阅、数据平台及按任务/用量收费
ServiceNow	NOW	1212	Now Assist、AI Control Tower、Autonomous Workforce	在IT、HR、客服、采购等工作流中部署和治理Agent；平台订阅+AI功能包及用量收费
Adobe	ADBE	1031	Firefly、Acrobat AI Assistant、GenStudio、创意Agent	将生成式AI嵌入创意、文档与营销 workflow；订阅+生成积分/Token+企业营销产品收费
Intuit	INTU	897	财税、会计、支付和个人金融AI Agent	基于财税与经营数据辅助记账、报税、咨询和风控；订阅+交易及增值服务收费
HubSpot	HUBS	125	Breeze AI、营销/销售/客服Agent	利用CRM和营销数据自动化获客、内容、销售跟进和客服；订阅套餐+AI增值模块
Zoom	ZM	295	AI Companion、会议与协作Agent	会议摘要、知识提取、任务分配及协作自动化；基础订阅带动留存，高端功能增值
GitLab	GTLB	61	Duo、代码生成、测试、安全与DevSecOps Agent	结合代码仓库、CI/CD和安全流程提升研发效率；按开发者席位及AI用量收费
CrowdStrike	CRWD	2137	Charlotte AI、终端安全与安全Agent	利用威胁情报和终端数据自动研判、响应和处置；安全平台订阅+模块化扩容
Veeva Systems	VEEV	347	生命科学云、内容/合规与商业运营AI	基于药企合规数据和研发、商业流程提供AI能力；行业软件订阅+服务收费
Tempus AI	TEM	84	医疗数据、精准诊疗、临床AI	将临床、病理和基因组数据用于诊断与治疗决策；检测、数据服务和软件收入
C3.ai	AI	15	企业AI应用平台、工业与能源AI应用	面向能源、制造、国防等行业提供预制AI应用；订阅及项目型部署收费
SoundHound AI	SOUN	28	语音AI、车载和客服Agent	语音交互、点餐、车载助手和客服自动化；按部署、调用量或客户合同收费

资料来源：Wind，公司公告，国信证券经济研究所整理

请务必阅读正文之后的免责声明及其项下所有内容

5.7 附录：A股B端场景垂类AI应用



行业	公司	代码	市值（亿元人民币）	主要AI产品及应用	B端客户
流程工业	中控技术	688777.SH	844	TPT工业时序大模型、生产优化、预测性维护、无人化操作、能耗优化	石化、化工、电力、冶金等流程工业企业
钢铁工业	宝信软件	600845.SH	459	钢铁大模型、MES/ERP/APS智能化、设备运维、质量预测	宝武体系及外部钢铁、制造企业
建筑	广联达	002410.SZ	161	AI算量、AI组价、AI评标、AI清标、方案设计、施工安全	设计院、施工企业、造价机构、建设单位
制造业IT	赛意信息	300687.SZ	111	善谋GPT、制造智能平台、质量/设备/物流分析、企业Agent	制造业大型企业
制造业ERP	鼎捷数智	300378.SZ	99	雅典娜平台、企业AI Agent、生产排程、供应链及经营分析	制造业中大型企业
金融IT	恒生电子	600570.SH	442	LightGPT、智能投研、投顾、客服、运营、合规和知识助手	券商、基金、银行、保险、期货公司
银行IT	宇信科技	300674.SZ	119	信贷、风控、营销、客服、研发智能体	银行及其他金融机构
财税AI	税友股份	603171.SH	203	“犀友”财税大模型、财税Agent、合规税优、涉税风控、经营分析	税务机关、中小企业、财税服务机构
医疗IT	卫宁健康	300253.SZ	172	WiNEX Copilot、病历生成、临床决策、医院运营及医保应用	医院、卫健委、医保机构
法律科技	华宇软件	300271.SZ	47	智慧庭审、审判辅助、法律知识服务、仲裁及破产重组AI	法院、检察院、政法机关、高校及企业法务
城市治理	数字政通	300075.SZ	71	人和大模型、城市智能体、智能执法、12345、城市生命线	地方政府、城管及城市运营部门
网络安全	深信服	300454.SZ	552	安全GPT、威胁研判、自动处置、AI+SASE、私有化AI办公	政企、金融、医疗、教育等
教育/医疗/司法	科大讯飞	002230.SZ	1061	星火大模型、教育、医疗、司法、汽车等行业智能体	学校、医院、政府、车企、央企
企业数据与AI平台	星环科技	688031.SH	143	数据治理、多模数据库、知识图谱、LLMOps、Agent开发、异构算力管理	金融、电力、交通、政企等行业客户
企业软件平台	用友网络	600588.SH	405	企业AI、Data Agent、财务/供应链/人力等业务智能体	大中型企业、央企、制造与金融客户
企业数字化服务	汉得信息	300170.SZ	183	ERP智能化、供应链与制造业Agent、企业数据治理	制造、零售、消费等大型企业

资料来源：Wind，公司公告，国信证券经济研究所整理

5.7 附录：传媒娱乐AI应用产业链



领域	公司	代码	市值（亿元人民币）	主要AI产品及应用	落地场景
AI营销	蓝色光标	300058.SZ	565	AI创意生成、智能投放、数字人及出海营销服务	品牌广告主、电商及出海客户
AI营销/出海	易点天下	301171.SZ	241	AI营销平台、创意生成与投放优化	跨境电商、游戏、应用及AI原生企业
AI营销	因赛集团	300781.SZ	48	因赛AI引擎、营销策划/海报/商品图生成	品牌客户、营销代理及电商商家
创意软件/视频	万兴科技	300624.SZ	110	Filmora等视频创作软件的AI剪辑、生成与Agent功能	中小企业、内容机构、专业创作者；B/C混合
视觉内容/版权	视觉中国	000681.SZ	132	视觉素材版权库、AIGC内容管理、视觉内容服务	媒体、广告公司、品牌方、平台
游戏研发与运营	三七互娱	002555.SZ	448	AI用于美术、文案、客服、投放、运营及玩法探索	以内部游戏研发和发行体系为主
游戏研发工具	恺英网络	002517.SZ	393	SOON全流程AI开发平台：动画渲染、建模、脚本生成	主要内部研发团队；可关注未来工具外溢
游戏AI/内容互动	巨人网络	002558.SZ	571	自研大模型能力、多模态模型、智能体及AI产品探索	游戏研发、运营及玩家互动场景
IP内容/阅读	中文在线	300364.SZ	193	网文IP、AI辅助创作、短剧/动漫等多模态内容开发	内容平台、版权合作方、创作者生态
互动娱乐/游戏	完美世界	002624.SZ	224	AI辅助游戏研发、美术与内容生产	主要内部研发、发行及运营
AI漫剧/IP内容	中文在线	300364.SZ	193	逍遥AI、次元神笔；用于AI漫剧、AI仿真人剧、网文及短剧生产	内容平台、版权合作方、制作团队；兼具C端内容分发
AI短剧/数字阅读IP	掌阅科技	603533.SH	104	AI短剧全流程：剧本生成、视觉制作、成片交付；数字阅读IP转化	短剧平台、内容制作与发行合作方
AI影视/短剧	华策影视	300133.SZ	143	AI漫剧、AI真人剧、AI本地化翻译及国际发行	长短视频平台、海外发行渠道、品牌与内容合作方
AI短剧出海	昆仑万维	300418.SZ	571	DramaWave短剧平台、AI漫剧与AI短剧自动化制作、多语种发行	海外内容消费者、渠道及广告/分发合作方
AI影视制作	百纳千成	300291.SZ	52	AI创作、AI精准营销赋能影视、剧集和互动影视游戏	视频平台、品牌方、内容联合出品方
AI影视/IP	欢瑞世纪	000892.SZ	40	AI漫剧、AI真人剧；剧本孵化、素材生成、镜头设计与后期制作	视频平台、品牌及IP合作方

资料来源：Wind，公司公告，国信证券经济研究所整理

风险提示

技术不达预期风险；

商业化不达预期风险；

政策风险；

经营管理风险等。

免责声明

国信证券投资评级

投资评级标准	类别	级别	说明
报告中投资建议所涉及的评级（如有）分为股票评级和行业评级（另有说明的除外）。评级标准为报告发布日后6到12个月内的相对市场表现，也即报告发布日后的6到12个月内公司股价（或行业指数）相对同期相关证券市场代表性指数的涨跌幅作为基准。A股市场以沪深300指数（000300.SH）作为基准；新三板市场以三板成指（899001.CSI）为基准；香港市场以恒生指数（HSI.HI）作为基准；美国市场以标普500指数（SPX.GI）或纳斯达克指数（IXIC.GI）为基准。	股票投资评级	优于大市	股价表现优于市场代表性指数10%以上
		中性	股价表现介于市场代表性指数±10%之间
		弱于大市	股价表现弱于市场代表性指数10%以上
		无评级	股价与市场代表性指数相比无明确观点
	行业投资评级	优于大市	行业指数表现优于市场代表性指数10%以上
		中性	行业指数表现介于市场代表性指数±10%之间
		弱于大市	行业指数表现弱于市场代表性指数10%以上

分析师承诺

作者保证报告所采用的数据均来自合规渠道；分析逻辑基于作者的职业理解，通过合理判断并得出结论，力求独立、客观、公正，结论不受任何第三方的授意或影响；作者在过去、现在或未来未就其研究报告所提供的具体建议或所表述的意见直接或间接收取任何报酬，特此声明。

重要声明

本报告由国信证券股份有限公司（已具备中国证监会许可的证券投资咨询业务资格）制作；报告版权归国信证券股份有限公司（以下简称“我公司”）所有。本报告仅供我公司客户使用，本公司不会因接收人收到本报告而视其为客户。未经书面许可，任何机构和个人不得以任何形式使用、复制或传播。任何有关本报告的摘要或节选都不代表本报告正式完整的观点，一切须以我公司向客户发布的本报告完整版本为准。

本报告基于已公开的资料或信息撰写，但我公司不保证该资料及信息的完整性、准确性。本报告所载的信息、资料、建议及推测仅反映我公司于本报告公开发布当日的判断，在不同时期，我公司可能撰写并发布与本报告所载资料、建议及推测不一致的报告。我公司不保证本报告所含信息及资料处于最新状态；我公司可能随时补充、更新和修订有关信息及资料，投资者应当自行关注相关更新和修订内容。我公司或关联机构可能会持有本报告中所提到的公司所发行的证券并进行交易，还可能为这些公司提供或争取提供投资银行、财务顾问或金融产品等相关服务。本公司的资产管理部门、自营部门以及其他投资业务部门可能独立做出与本报告意见或建议不一致的投资决策。

本报告仅供参考之用，不构成出售或购买证券或其他投资标的的要约或邀请。在任何情况下，本报告中的信息和意见均不构成对任何个人的投资建议。任何形式的分享证券投资收益或者分担证券投资损失的书面或口头承诺均为无效。投资者应结合自己的投资目标和财务状况自行判断是否采用本报告所载内容和信息并自行承担风险，我公司及雇员对投资者使用本报告及其内容而造成的一切后果不承担任何法律责任。

证券投资咨询业务的说明

本公司具备中国证监会核准的证券投资咨询业务资格。证券投资咨询，是指从事证券投资咨询业务的机构及其投资咨询人员以下列形式为证券投资人或者客户提供证券投资分析、预测或者建议等直接或者间接有偿咨询服务的活动：接受投资人或者客户委托，提供证券投资咨询服务；举办有关证券投资咨询的讲座、报告会、分析会等；在报刊上发表证券投资咨询的文章、评论、报告，以及通过电台、电视台等公众传播媒体提供证券投资咨询服务；通过电话、传真、电脑网络等电信设备系统，提供证券投资咨询服务；中国证监会认定的其他形式。

发布证券研究报告是证券投资咨询业务的一种基本形式，指证券公司、证券投资咨询机构对证券及证券相关产品的价值、市场走势或者相关影响因素进行分析，形成证券估值、投资评级等投资分析意见，制作证券研究报告，并向客户发布的行为。



国信证券
GUOSEN SECURITIES

国信证券经济研究所

深圳

深圳市福田区福华一路125号国信金融大厦36层

邮编：518046 总机：0755-82130833

上海

上海浦东民生路1199弄证大五道口广场1号楼12楼

邮编：200135

北京

北京西城区金融大街兴盛街6号国信证券9层

邮编：100032