

汽车行业深度报告

AI 系列深度 15 期: Transformer 如何开启 AI 第二次爆发?

增持 (维持)

2026 年 08 月 05 日

证券分析师 黄细里

执业证书: S0600520010001

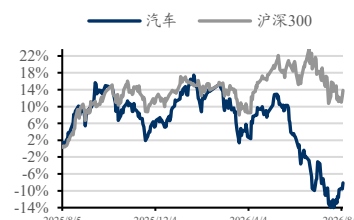
021-60199793

huangxl@dwzq.com.cn

投资要点

- **AI 第二次爆发的本质,是模型能力形成方式从“人工设计”转向“规模化学习”。**AI1.0 时代深度学习仍带有“一任务一模型”特征,CNN、RNN 等架构分别服务于图像、文本等专用任务,能力主要停留在基础感知层面。本轮 AI2.0 以 Transformer 等通用底座为核心,叠加数据、算力与自监督训练,使模型学习统计结构与通用表征,能力来源由强任务先验和强归纳偏置,转向弱归纳偏置下的规模化学习与涌现。
- **Transformer 成为 AI2.0 关键基础架构:** 自注意力机制解决循环网络串行计算和长程依赖瓶颈,更适合 GPU 并行算力与大规模训练;弱归纳偏置为跨模态迁移提供统一接口,使多类输入逐步进入同一底座体系。2017 年以来,主干架构总体稳定,能力提升更多来自 MoE、长上下文、多模态融合、后训练与推理机制,而非底层架构切换。
- **从技术演进看,大模型可概括为“模型能力奠基—人机对齐与行为优化—推理与外部系统赋能”三阶段,并由能力广度和使用效率两条支线贯穿。**第一阶段通过预训练、规模扩张和 Scaling Laws 奠定参数内能力;第二阶段通过 SFT、RLHF、DPO 等对齐技术提升可用性;第三阶段依托思维链、RAG、工具调用、强化学习推理和任务路由,将能力增量迁移到运行时系统。
- **GPT 系列提供了大语言模型落地产业端的典型路径:** GPT-1 至 GPT-3 验证 decoder-only 架构和规模化泛化能力,InstructGPT 与 ChatGPT 补齐交互入口,GPT-4/GPT-4o 推动模型走向多模态交互,o1/o3/o4-mini/GPT-5 进一步整合推理时计算、工具使用、多模态任务和实时路由。关键技术通常先在研究与工程验证中形成,再通过产品化入口实现规模扩散。
- **Scaling Law 回答“规模如何转化为能力”,也反映 AI2.0 资源优化边界的外扩。**早期 scaling 围绕参数量、数据量与训练算力展开,随后纳入能力涌现、计算最优配比、高质量数据、对齐、部署和推理成本等约束。2024 年以来,MoE 稀疏激活、测试时计算、RL reasoning、扩散 Transformer 与视频生成等方向,使 scaling 扩展为训练、推理、数据、架构、对齐与多模态的全链条资源优化框架。
- **从产业形态看, AI 第二次爆发的外在结果是数字 AI 任务边界初步收敛。**Transformer 已参与文本、代码、知识库、工具、多模态输入和数字工作流,推动架构、任务与模态三重收敛。产业价值重心也从“设计何种架构”,转向“谁能把规模、数据、对齐、推理与系统工程能力做到位”。但推理阶段增量目前更多体现在语言、数学、代码等可验证任务中,向视觉生成、复杂多模态交互与真实物理闭环迁移仍需验证。
- **我们认为,数字 AI 底座模型发展路径已逐渐清晰、价值在于确定性,但其难以实现物理世界的建模闭环;我们认为物理 AI 将成为物理世界的 AI 解决方案,从而在 AI 的当前发展阶段成为增量更大的方向,智能驾驶作为最先跑通闭环的代表场景,我们认为应重点关注。**
- **风险提示:** 技术迭代不及预期的风险,大模型厂商 ARR 增速放缓的风险,云服务商资本开支不及预期的风险,智能驾驶及机器人商业化落地不及预期的风险。

行业走势



相关研究

《AI 系列深度 14 期: 深度学习如何开启 AI 第一次爆发?》

2026-08-05

《AI 系列深度 13 期: 世界模型如何实现?》

2026-08-05

内容目录

1. AI2.0 的本质：从“人工设计”到“规模化学习” 4

2. Transformer 成为跨模态通用底座 4

3. Transformer 发展复盘 5

 3.1. 阶段一：模型能力奠基 6

 3.2. 阶段二：人机对齐、行为优化 6

 3.3. 阶段三：推理与外部系统赋能 7

 3.4. 模型能力的广度与效率 7

4. 复盘 GPT：大语言模型如何落地产业端 7

5. Scaling Law 复盘：从预训练幂律到全链条资源优化 11

6. 大语言模型或已成为数字 AI 任务的通解 13

7. 结论 13

8. 风险提示 14

图表目录

图 1: Transformer 结构示意.....5

图 2: Transformer 时代重要论文/成果发展时间轴.....6

图 3: GPT2018—2026 发展历程.....8

图 4: GPT-1 至 GPT-3 模型对比8

图 5: RLHF（基于人类反馈的强化学习）三阶段8

图 6: GPT-3.5/ChatGPT 正式进入产品化阶段.....9

图 7: GPT-4 支持多模态任务，可处理图片类信息输入.....10

图 8: o1 在推理密集型任务中的表现显著优于 GPT-4o.....10

图 9: GPT-5 的统一系统与推理路由机制10

图 10: 扩展规律——损失随算力/参数/数据的幂律下降12

图 11: 代表性模型规模演进（参数量 vs 训练数据量）12

图 12: Transformer 带来的模型三重收敛（专用模型→通用基础模型）13

表 1: 主流架构的归纳偏置.....4

表 2: 扩展规律（scaling law）的六阶段演进12

1. AI2.0 的本质：从“人工设计”到“规模化学习”

2012 年 AlexNet 在 ImageNet 数据集上取得突破，正式开启深度学习第一轮发展浪潮。该阶段具有鲜明的“一任务一模型”特征：CNN、RNN 等网络架构各自适配特定任务领域，模型核心价值主要集中在图像、文本等基础感知环节。受架构能力约束，当时模型尚难形成流畅的人机交互能力，这也是 To C 产品未能大规模落地的重要原因。

AI 第二轮产业变革的核心逻辑出现根本变化：模型能力不再主要依靠人工定制，而是依托海量数据与规模化训练自主形成。上一轮深度学习的模型性能高度依赖人为设计的网络结构，研发人员依据图像局部特征、时序依赖等现实规律设计架构。本轮大模型以 Transformer 等统一通用架构为基础，叠加算力提升与数据规模扩张，通过自监督训练从海量样本中学习统计结构与通用表征，使模型能力形成方式由过去依赖强任务先验和强归纳偏置，转向弱归纳偏置下的规模化学习与能力涌现。

各类主流网络架构的归纳偏置强度整体呈递减趋势：CNN 绑定较强的空间特征约束，RNN 内置时序递归逻辑，而 Transformer 主要依靠注意力机制与位置编码搭建通用框架，人为预设约束相对最少，模型效果更多由数据驱动。对比 AI 1.0 与 AI 2.0 可以看到，长期维度下，依托通用框架叠加算力扩容的技术路线，逐步优于人工注入领域知识的定制化方案。

表1：主流架构的归纳偏置

架构	归纳偏置	代表能力	与算力/数据的关系
卷积网络 (CNN)	强空间先验： 局部性、平移 等变	图像等感知任 务	先验强、样本效率高，但通用性 与上限受限
循环网络 (RNN)	强序列先验： 递归、时序依 赖	语言等序列建 模	顺序计算难并行，长程依赖建模 偏弱
Transformer	弱先验：注意 力+位置编码	跨模态通用建 模	可并行、可充分利用大算力与大 数据

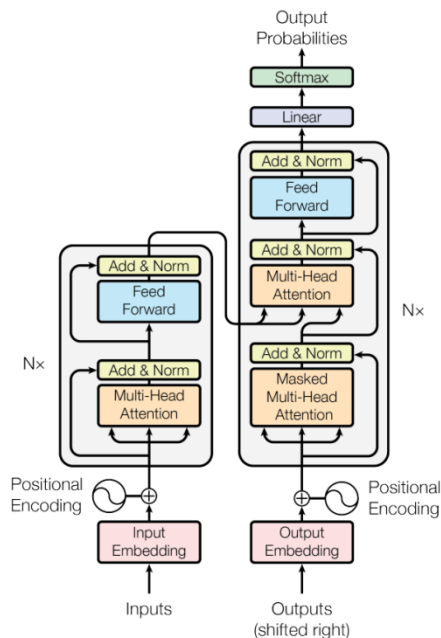
数据来源：《Attention Is All You Need》，《ImageNet Classification with Deep Convolutional Neural Networks》等其他文献，东吴证券研究所

2. Transformer 成为跨模态通用底座

Transformer 是 AI2.0 的起点。相较循环网络，其主要突破两个瓶颈：1）循环网络需按时间步串行计算，难以并行；自注意力允许序列内各位置并行计算，能够更充分利用 GPU 算力；2）循环网络在远距离信息传递中存在衰减，全局自注意力使任意两个位置可直接交互，从而缓解长程依赖问题。

同时，Transformer 能够成为大模型通用基础架构，除并行计算优势外，另一关键原因在于其弱归纳偏置设计，天然支持跨模态迁移与复用。

图1: Transformer 结构示意图



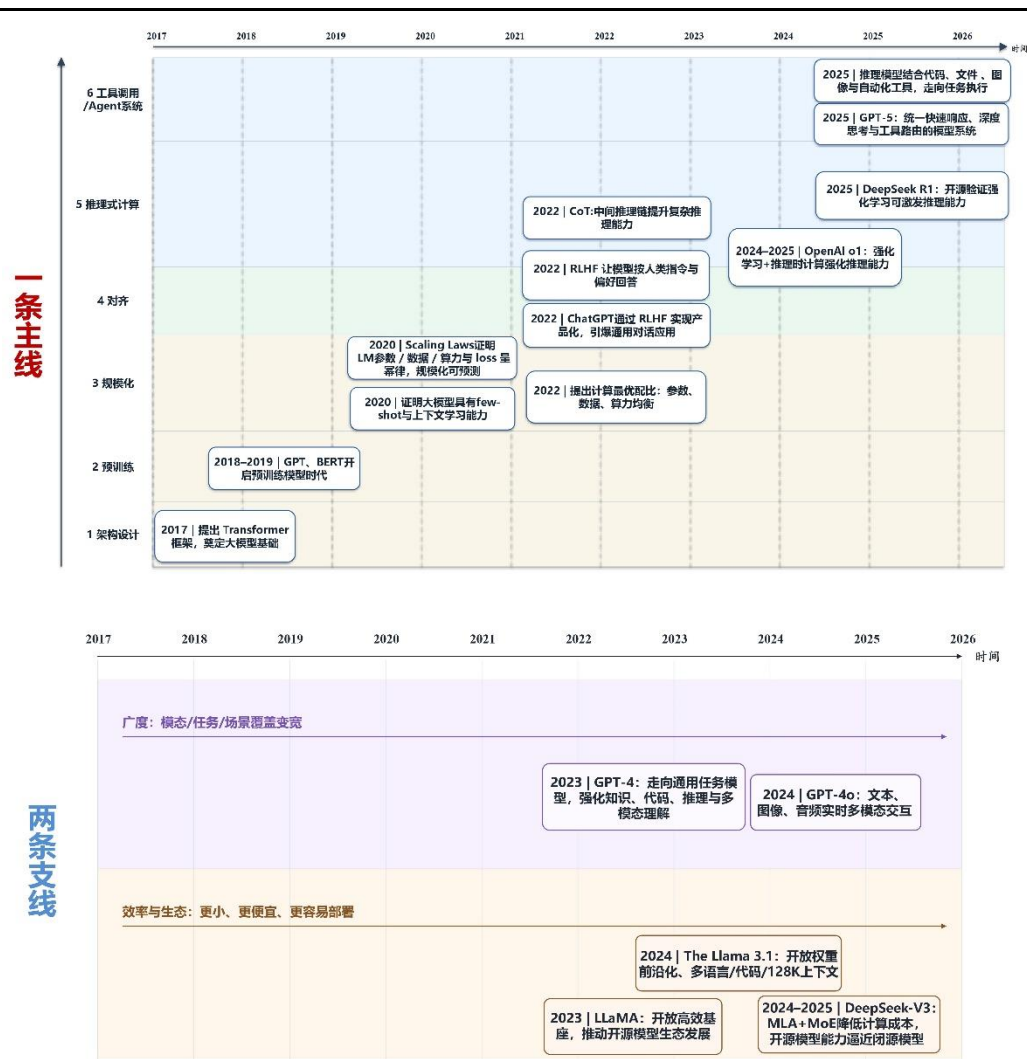
数据来源：《Attention Is All You Need》，东吴证券研究所

2017 年以来，主流大模型底层主干架构保持高度稳定，Transformer 仍是前沿大模型最核心的架构基础。近年来模型能力提升更多来自 MoE 稀疏化、注意力机制优化、长上下文扩展、多模态 token 融合，以及以强化学习和 test-time compute 为代表的推理型后训练机制，而非底层架构范式的全面切换。

3. Transformer 发展复盘

结合核心论文，我们认为大模型演进框架可以总结为：一条核心发展主线搭配两条并行优化支线。主线为模型能力升级重心持续向外转移，分为三个递进阶段；两条支线分别对应能力拓展广度、训练推理使用效率。该划分标准区分能力优化所处层级，分为架构设计、预训练、规模扩张、对齐调优、推理增强、智能体系统等各类技术进展。

图2: Transformer 时代重要论文/成果发展时间轴



数据来源:《Attention Is All You Need》,《Language Models are Few-Shot Learners》等其他文献, 东吴证券研究所

3.1. 阶段一: 模型能力奠基

本阶段核心思路是依托模型参数承载全部基础能力，覆盖底层架构设计、预训练、规模扩容三大环节。2017 年 Transformer 问世，搭建起支持并行运算的注意力基础框架；2018-2019 年 GPT-1、BERT 落地，确立预训练+微调的主流训练模式，实现海量文本数据的特征学习；2020 年 Scaling Laws 理论与 GPT-3 落地，证明模型性能可随规模提升稳定增长，形成可量化的工程落地路径；2022 年 Chinchilla 进一步明确算力约束下，参数量与训练数据的最优配比。

整体来看，该阶段核心目标是将通用基础能力固化至模型参数，奠定大模型“预训练+规模扩张”的底层发展范式，模型性能上限直接由算力、数据投入规模决定。

3.2. 阶段二: 人机对齐、行为优化

本阶段不会为模型新增底层原生能力，而是校准模型现有输出逻辑，使其贴合人类指令与使用偏好。代表性技术包括 FLAN 指令微调、InstructGPT 提出的 SFT 监督微调-奖励模型-强化学习 RLHF 三段式流程、Constitutional AI、DPO 等，ChatGPT 的推出则是该阶段技术落地、走向商业化的标志性节点。

我们认为，模型具备基础能力、可正常落地使用、输出行为可控是三个独立维度，对齐技术主要解决后两大痛点。正是依托对齐相关技术，预训练模型中潜藏但无法直接落地的能力得以释放，成为 AI2.0 从实验室技术转向大众产品爆发的核心转折点。

3.3. 阶段三：推理与外部系统赋能

本阶段聚焦模型部署使用环节，依托运行时配套系统补充额外能力。2022 年思维链技术落地，通过分步推理优化模型解题逻辑，能力提升不再单纯依赖模型参数；RAG 检索增强、Toolformer、ReAct 等方案，将外部知识库、工具能力与模型解耦；2024-2025 年 o1、DeepSeek-R1 等工作将强化学习用于长链条推理与可验证任务，打开推理阶段算力优化的新方向；o3、GPT-5 更进一步整合推理逻辑、工具调用、任务调度路由，产品形态从单一独立模型升级为一体化智能系统。

3.4. 模型能力的广度与效率

两条并行优化支线贯穿整个发展阶段，两条支线不会改变能力优化的核心层级，而是决定同一阶段下，模型可覆盖场景范围、落地成本高低：

能力拓展广度：Codex 实现代码领域适配，Flamingo、GPT-4o、Gemini 打通图像、音频、视频多模态输入，超长上下文技术将单轮可处理信息量级提升至百万 token；

落地使用效率：MoE 稀疏架构（Switch Transformer、Mixtral、DeepSeek-V3）在固定算力下拓展有效参数规模，LoRA/QLoRA、FlashAttention、RoPE 等技术降低训练与推理部署成本，LLaMA、Mistral、通义千问、DeepSeek 等开源模型完善产业生态。

站在产业视角来看，当前行业底层架构逐步趋同，预训练方案也趋于同质化，行业竞争核心已不再是底层架构设计优劣，而是企业在模型外部体系的搭建能力，具体涵盖对齐数据集与优化方案、推理架构与工具链路、全流程训练部署工程化能力，行业价值创造重心，正从底层架构研发转向全链路工程落地。

4. 复盘 GPT：大语言模型如何落地产业端

我们选择 GPT 作为大模型落地的产品代表，从产业视角复盘大语言模型由学术成果走向规模化应用的完整周期。2018—2026 年，GPT 大致经历基础语言模型、产品化、多模态、推理模型四个阶段。每一次跃迁，均对应 Transformer 能力被进一步产品化、系统化的关键节点。

图3：GPT2018—2026 发展历程

发展阶段	时间	代表版本	阶段发展
基础语言模型阶段	2018-2022	GPT-1, GPT-2, GPT-3, InstructGPT	从"证明Transformer解码器可通过大规模预训练学习语言能力", 发展到"通过规模化提升通用能力", 再到"通过指令微调与RLHF让模型听懂并执行人类指令"
产品化阶段	2022-2023	ChatGPT, GPT-3.5	GPT从研究模型进入大众产品形态, 核心变化是从"能补全文本"变成"能持续对话、理解用户意图并完成任务"
多模态阶段	2023-2024	GPT-4, GPT-4 Turbo, GPT-4o	GPT从文本模型扩展为图文音多模态智能系统, 能力从语言理解扩展到视觉、语音、图像理解与工具协作
推理模型阶段	2024-2026	o1, o3, o4-mini, GPT-5, GPT-5.5	GPT从"聊天助手进一步转向"能思考能调用工具、能完成复杂任务"的统一智能体系系统

数据来源：OpenAI 官方网站，东吴证券研究所整理

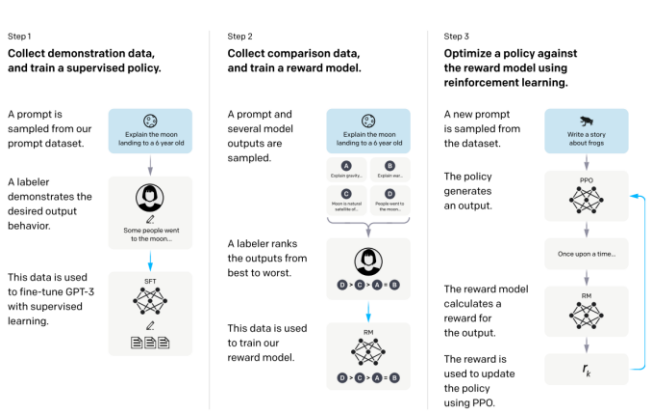
基础语言模型阶段（GPT-1 至 GPT-3、InstructGPT）的核心特征，是 decoder-only 架构确立与规模化泛化能力显现。OpenAI 基本确立以自回归 Transformer 为核心的大语言模型路线，并依托参数、数据、算力扩张持续提升模型能力：GPT-1 验证预训练表征可迁移至下游任务，GPT-2 展现无需微调的 zero-shot 泛化能力，GPT-3 进一步体现 few-shot 与上下文学习能力，并在工程层面开始以 scaling law 指导规模扩张。InstructGPT 在预训练模型基础上引入人工示范、人工偏好数据与 RLHF，使模型输出更贴合人类意图，形成 SFT—奖励模型—强化学习（PPO）的经典三阶段流程。

图4：GPT-1 至 GPT-3 模型对比

	GPT 1	GPT 2	GPT 3
最终模型参数量	117M	1.5B	175B
Transformer层数	12	48	96
d_model	768	1600	12288
注意力头数	12	-	96
上下文长度	512 tokens	1024 tokens	2048 tokens
下游任务适配方式	fine-tuning	zero-shot	few-shot / prompt-based
标志性突破	预训练+微调范式 零样本多任务能力 in-context learning		

数据来源：《Improving Language Understanding by Generative Pre-Training》，《Language Models are Few-Shot Learners》等其他文献，东吴证券研究所

图5：RLHF（基于人类反馈的强化学习）三阶段



数据来源：《Training Language Models to Follow Instructions with Human Feedback》等其他文献，东吴证券研究所

产品化阶段（ChatGPT/GPT-3.5）的核心变化，是基础语言模型正式转向对话式 AI

助手产品。GPT-3.5 是底层模型系列，ChatGPT 则是在其基础上经过后训练与产品化封装形成的对话应用，补齐模型从“能续写”到“能对话、能理解用户意图并完成任务”的关键环节，也成为大语言模型进入大众应用市场的标志性拐点。

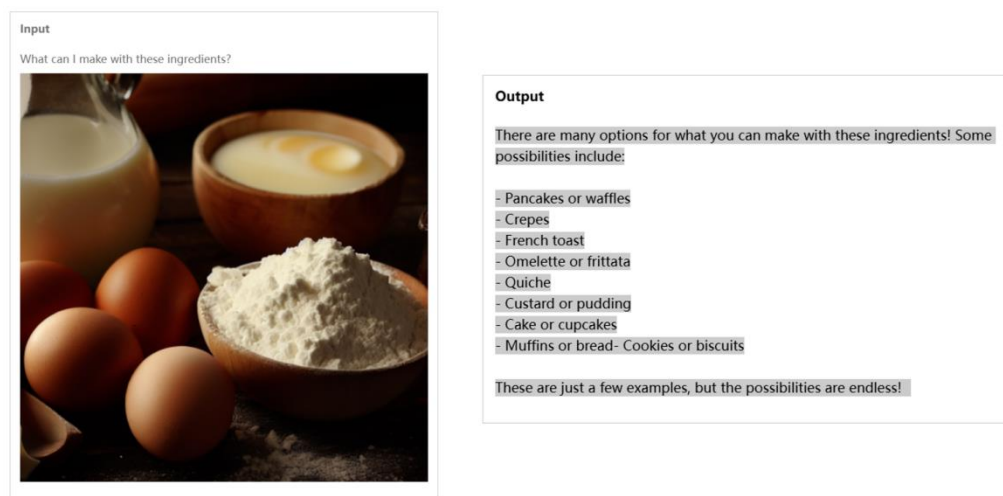
图6： GPT-3.5/ChatGPT 正式进入产品化阶段

	GPT-3	GPT-3.5/ChatGPT
模型定位	基础语言模型	对话式 AI 助手
训练目标	标准自回归语言建模，即预测下一个 token	在语言建模基础上，强化指令跟随、对话生成和人类偏好对齐
任务适配方式	主要依赖 prompt、zero-shot / few-shot / in-context learning	更强的自然语言指令跟随、多轮上下文理解
后训练方式	仅在预训练基础上通过 prompt 或特定场景微调适配任务	SFT + RM + PPO，面向助手行为优化
对话能力	非对话优化，可通过 prompt 模拟聊天	面向多轮自然对话优化，形成 ChatGPT 产品形态
产品形态	主要以 API / 文本补全模型形式使用	ChatGPT 网页产品 + GPT-3.5 Turbo API，进入大众用户和开发者生态

数据来源：OpenAI 官方网站，东吴证券研究所整理

多模态阶段（GPT-4/GPT-4 Turbo/GPT-4o）中，OpenAI 的优化重点从单纯提升文本生成与指令跟随，进一步转向输入模态扩展。GPT-4 引入图文输入，使模型能够在文本与图像共同输入下完成视觉问答、图表解释、截图分析与文档理解，ChatGPT 由此具备“看图理解”的多模态能力；GPT-4o 则进一步走向原生端到端实时多模态交互，推动模型产品形态由文本助手向综合交互入口升级。

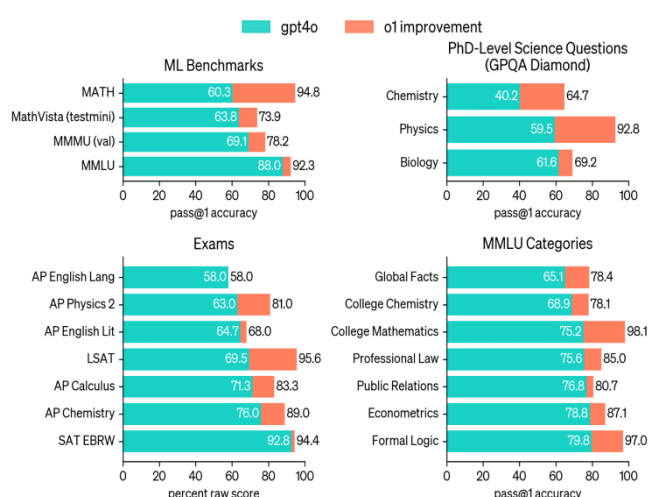
图7: GPT-4 支持多模态任务, 可处理图片类信息输入



数据来源: OpenAI 官方网站, 东吴证券研究所整理

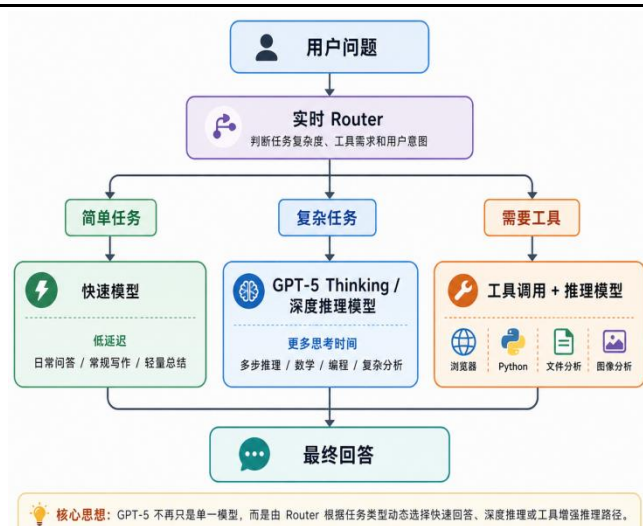
推理模型阶段 (o1/o3/o4-mini/GPT-5) 的核心突破, 是推理时计算量扩展。o1 通过强化学习优化模型解题过程, 使模型在回答前先完成问题拆解、步骤规划与推理校验, 而非直接快速生成答案; o3、o4-mini 将推理能力进一步延伸至工具使用与多模态任务; GPT-5 则在性能升级基础上, 将快速模型、深度推理模型与实时路由机制整合为统一系统, 以适配快速问答、复杂推理以及不同成本、延迟约束下的多类使用场景。

图8: o1 在推理密集型任务中的表现显著优于 GPT-4o



数据来源: OpenAI 官方网站, 东吴证券研究所整理

图9: GPT-5 的统一系统与推理路由机制



数据来源: OpenAI 官方网站, 东吴证券研究所整理

复盘 GPT 迭代历程可以发现, 其产业落地关键节点与前文划分的三大演进阶段高度一致: 基础大模型发展阶段, 对应模型内部原生能力搭建; 以 RLHF、ChatGPT 为代表的产品化阶段, 对应人机对齐行为优化; 多模态大模型与推理模型落地, 则对应推理

与外部系统赋能阶段，并同步带动能力广度、落地效率两条支线持续展开。

通过对 GPT 发展的复盘，我们发现从 RLHF 到 ChatGPT，从长链推理到 o 系列模型，技术路线往往先在研究与工程验证中形成，再通过产品化入口实现规模扩散。当前，大模型能力提升的重心已不再局限于预训练阶段的参数扩张，而是逐步转向由后训练对齐、测试时计算、工具调用、外部数据与自动化工作流共同构成的系统化能力扩展。由此，围绕经典技术论文梳理行业演进脉络，能够帮助我们识别下一阶段产业落地的潜在方向与节奏。

5. Scaling Law 复盘：从预训练幂律到全链条资源优化

扩展规律（scaling law）回答的是“规模如何转化为能力”。这一规律伴随大模型发展不断被修正与拓展，逐步形成覆盖训练、数据、对齐、推理与多模态的全链条框架。

阶段 0（2017—2019）中，相关研究开始发现模型误差会随数据、模型与算力扩张呈稳定幂律下降，为“深度学习性能可被外推”提供了关键思想基础。

阶段 1（2020）确立预训练 scaling law。Kaplan 等将语言模型 loss 抽象为参数量、数据量与算力的幂律函数，GPT-3 进一步证明 scale 不仅能够降低 loss，也会带来 few-shot 与上下文学习能力。

阶段 2（2021—2022）转向能力 scaling 与涌现争论。Gopher、PaLM 将 scaling 研究从 loss 曲线推向多任务能力谱系，“涌现能力”概念由此形成，随后也引发“涌现或由指标选择造成”的讨论。

阶段 3（2022—2023）进入计算最优 scaling。Chinchilla 修正早期“只堆参数”的路径，指出在给定算力约束下，参数量与训练 token 数应同步增长，并以 700 亿参数 Chinchilla 在多项评测上优于 2800 亿参数 Gopher 作为验证。

阶段 4（2023—2024）进一步纳入数据、对齐与部署约束。高质量 token 成为新的瓶颈，数据质量与重复、指令微调、奖励模型过优化、推理成本等因素均被纳入 scaling 变量。

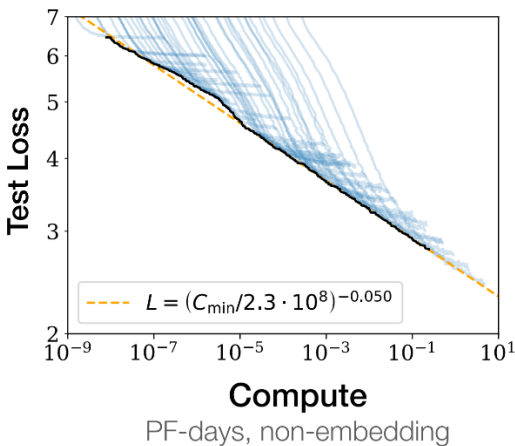
阶段 5（2024—2026）扩展出新的 scaling 轴。MoE 稀疏激活、测试时计算、RL reasoning、扩散 Transformer 与视频生成等方向，共同推动 scaling law 从“预训练大模型规律”，演化为覆盖训练、推理、数据、架构、对齐与多模态的全链条资源优化框架。

表2: 扩展规律（scaling law）的六阶段演进

阶段	时间	阶段名称	核心进展	代表成果
阶段 0	2017-2019	规模规律前史：性能可预测化	误差随数据/模型/算力呈幂律下降，性能可外推	Hestness 等（2017）
阶段 1	2020	预训练 scaling law 确立	loss 抽象为参数/数据/算力的幂律；GPT-3 证明 scale 带来 few-shot	Kaplan 等、Brown 等（GPT-3）、Henighan 等
阶段 2	2021-2022	能力 scaling 与涌现争论	重心由 loss 转向能力谱系；“涌现能力”及其“指标幻觉”之辩	Gopher、PaLM、Wei 等、Schaeffer 等
阶段 3	2022-2023	计算最优 scaling law	Chinchilla 修正“只堆参数”，参数与 token 按算力共同增长	Hoffmann 等（Chinchilla）、LLaMA
阶段 4	2023-2024	数据、对齐与部署约束 scaling	高质量 token 成瓶颈，数据质量/重复、指令微调、奖励模型纳入	Muennighoff 等、Chung 等、FineWeb、DataComp-LM
阶段 5	2024-2026	新 scaling 轴：多模态、MoE、推理时算力	进入多变量时代：MoE、测试时算力、RL reasoning、扩散/视频生成	o1、DeepSeek-R1、DiT、MoE scaling

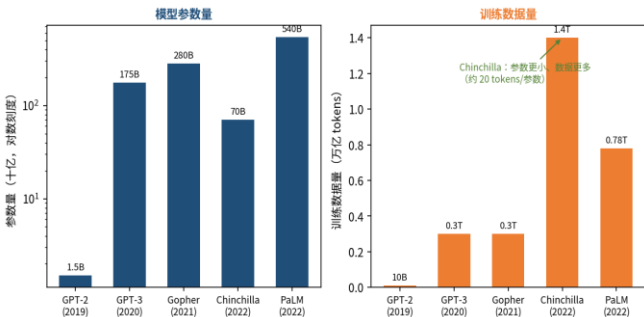
数据来源：《Scaling Laws for Neural Language Models》，《Training Compute-Optimal Large Language Models》等其他文献，东吴证券研究所

图10: 扩展规律——损失随算力/参数/数据的幂律下降



数据来源：《Scaling Laws for Neural Language Models》，东吴证券研究所

图11: 代表性模型规模演进（参数量 vs 训练数据量）



数据来源：《Language Models are Few-Shot Learners》等其他文献，东吴证券研究所

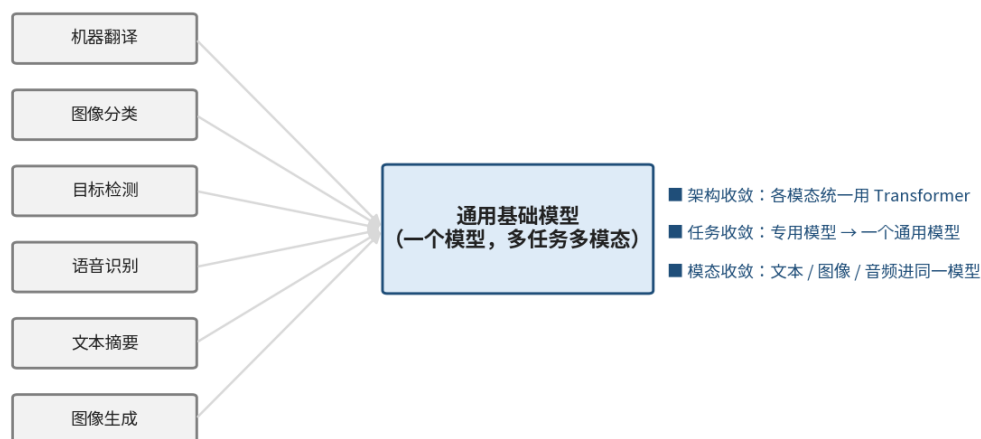
扩展规律六阶段演进，与三大演进阶段在逻辑上互为表里：当 scaling 从“预训练参数与 token”扩展到“数据、对齐、推理时算力与多模态”，本质上对应能力升级重心从模型内部原生能力搭建，继续延伸至人机对齐行为优化和推理与外部系统赋能。换言之，scaling 的边界正在从训练阶段延伸至对齐与运行时。这进一步印证，AI2.0 后续能力增量与成本结构，将越来越多由模型外部工程能力决定。

6. 大语言模型或已成为数字 AI 任务的通解

较之 AI 1.0 时代不同任务匹配不同基座模型的产业形态，在 AI 2.0 时代，Transformer 已以基座模型或重要组件的角色参与到广泛的数字 AI 任务解决过程，使得 AI 任务在数字 AI 领域的解决方案实现初步收敛。

我们将其概括为三重收敛：**架构收敛**，即文本、图像、音频等各类模态逐步统一到 Transformer 底座；**任务收敛**，即从“一个任务一个模型”走向“一个通用模型应对多任务”；**模态收敛**，即文本、图像、音频等不同输入进入同一模型体系处理。三者共同构成 AI 从专用模型走向通用底座的外在表现。但是，我们认为大语言模型尚未展现出可作为物理 AI 任务通解的迁移能力。

图12：Transformer 带来的模型三重收敛（专用模型→通用基础模型）



数据来源：东吴证券研究所绘制

7. 结论

第一，关于本质。AI 第二次爆发的核心，是模型能力从依赖人工规则与任务专用

架构设计，转向依托通用架构、海量数据、参数规模与算力投入的共同扩展。相较于上一代任务专用模型，大模型不再主要依赖人为注入、且与具体任务深度绑定的强归纳偏置，而是通过统一底座模型在大规模数据中学习通用表征，并在不同任务、模态与场景中实现能力迁移。

第二，关于发展历史及未来方向。Transformer 时代的大模型演进，大体可以沿着“预训练与规模化奠基—人机对齐与行为优化—推理增强与系统化扩展”三大阶段展开。自 2017 年 Transformer 提出后，主干架构范式总体趋于稳定，后续边际进步更多来自规模化、数据工程、后训练对齐、测试时计算、工具调用与系统工程等环节。与之对应，扩展规律也从早期围绕预训练参数、数据 token 与训练算力的规模化，进一步延伸至数据质量、对齐方法、推理时算力、多模态输入输出与外部工具系统的全链条优化。相应地，行业价值重心正从单纯讨论“设计何种架构”，转向比较“谁能把规模、数据、对齐、推理与系统工程能力做到位”。算力与系统工程、数据工程、对齐优化、推理系统建设，将成为后续更需重点关注的环节。

需要提示的是，推理阶段算力与强化学习带来的能力增量，目前更多体现在语言、数学、代码等结果相对可验证的任务中；其能否稳定迁移至视觉生成、复杂多模态交互与真实物理闭环场景，仍需持续验证。

第三，关于大语言模型成为数字 AI 任务的通用底座。AI 第二次爆发的外在形态，是架构、任务与模态的持续收敛。大语言模型不仅提升了文本任务能力，也逐步成为连接代码、工具、知识库、多模态输入与数字工作流的统一交互入口。从产业角度看，在通用底座形成后，产业价值更容易向掌握基础模型能力、数据闭环、算力资源与生态入口的主体集中。

8. 风险提示

技术迭代不及预期的风险：数字 AI、物理 AI 技术进展低于预期，任务或应用场景拓宽速度低于预期。

大模型厂商 ARR 增速放缓的风险：若客户增长、续费率或客单价不及预期，模型厂商 ARR 增速可能放缓。

云服务商资本开支不及预期的风险：若自由现金流承压，或 AI 算力需求及投资回报低于预期，云服务商可能下调资本开支。

智能驾驶及机器人商业化落地不及预期的风险：若产品可靠性、成本下降或用户需求低于预期，相关产品商业化进程可能放缓。

免责声明

东吴证券股份有限公司经中国证券监督管理委员会批准，已具备证券投资咨询业务资格。

本研究报告仅供东吴证券股份有限公司（以下简称“本公司”）的客户使用。本公司不会因接收人收到本报告而视其为客户。在任何情况下，本报告中的信息或所表述的意见并不构成对任何人的投资建议，本公司及作者不对任何人因使用本报告中的内容所导致的任何后果负任何责任。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。

在法律许可的情况下，东吴证券及其所属关联机构可能会持有报告中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投资银行服务或其他服务。

市场有风险，投资需谨慎。本报告是基于本公司分析师认为可靠且已公开的信息，本公司力求但不保证这些信息的准确性和完整性，也不保证文中观点或陈述不会发生任何变更，在不同时期，本公司可发出与本报告所载资料、意见及推测不一致的报告。

本报告的版权归本公司所有，未经书面许可，任何机构和个人不得以任何形式翻版、复制和发布。经授权刊载、转发本报告或者摘要的，应当注明出处为东吴证券研究所，并注明本报告发布人和发布日期，提示使用本报告的风险，且不得对本报告进行有悖原意的引用、删节和修改。未经授权或未按要求刊载、转发本报告的，应当承担相应的法律责任。本公司将保留向其追究法律责任的权利。

东吴证券投资评级标准

投资评级基于分析师对报告发布日后 6 至 12 个月内行业或公司回报潜力相对基准表现的预期（A 股市场基准为沪深 300 指数，香港市场基准为恒生指数，美国市场基准为标普 500 指数，新三板基准指数为三板成指（针对协议转让标的）或三板做市指数（针对做市转让标的），北交所基准指数为北证 50 指数），具体如下：

公司投资评级：

买入：预期未来 6 个月个股涨跌幅相对基准在 15%以上；

增持：预期未来 6 个月个股涨跌幅相对基准介于 5%与 15%之间；

中性：预期未来 6 个月个股涨跌幅相对基准介于-5%与 5%之间；

减持：预期未来 6 个月个股涨跌幅相对基准介于-15%与-5%之间；

卖出：预期未来 6 个月个股涨跌幅相对基准在-15%以下。

行业投资评级：

增持：预期未来 6 个月内，行业指数相对强于基准 5%以上；

中性：预期未来 6 个月内，行业指数相对基准-5%与 5%；

减持：预期未来 6 个月内，行业指数相对弱于基准 5%以上。

我们在此提醒您，不同证券研究机构采用不同的评级术语及评级标准。我们采用的是相对评级体系，表示投资的相对比重建议。投资者买入或者卖出证券的决定应当充分考虑自身特定状况，如具体投资目的、财务状况以及特定需求等，并完整理解和使用本报告内容，不应视本报告为做出投资决策的唯一因素。

东吴证券研究所
苏州工业园区星阳街 5 号

邮政编码：215021

传真：（0512）62938527

公司网址：<http://www.dwzq.com.cn>