

# AI 前沿发展日报 | 2026-07-24 (Asia)

日期：2026-07-24；覆盖窗口：2026-07-23 07:00 至 2026-07-24 a i )，纳入窗口内发布或完成公开确认的重要新增。

## 今日要点

- OpenAI 将 Health in ChatGPT 推向美国成年用户，允许用户连接病历与 Health，并在普通对话中按次授权调用。消费级 AI 正从通用问答进入个人纵向健康数据层。
- Alphabet 二季度披露 Gemini 月活 9.5 亿、第一方模型 API 每分钟约 2 ns；Google Cloud 收入同比增长 82% 至 248 亿美元。AI 采用已开始同时体现调用量和云收入上。
- AMD 与 Cerebras 将长上下文提示处理和高速 token 生成拆给不同计算系统，并计划年内通过 Cerebras Cloud 提供联合方案。推理竞争正在从“单一芯片跑完整负载”转向按阶段组合硬件。
- 美国两党议员提出 AI Kill Switch Act，并同步推动最强模型接受独立安全审计。前沿模型的紧急限流、能力禁用和停机能力，正被尝试从企业预案变成法定义务。

## 今日结论

1. 消费 AI 的下一个高价值入口是持续、结构化的个人数据，但产品能否进入健康、金融等场景，首先取决于逐次授权、数据删除、跨应用披露控制和责任边界，而不只是回答质量。
2. AI 商业化正在获得更硬的验证：月活、API tokens 和云收入可以相互印证需求；企业采购应同样把席位、有效调用、流程产出和单位经济性放在一张表中核算。
3. 前沿 AI 的产业控制点正同时向两端移动：底层推理按工作阶段组合异构硬件，上层监管要求保留可审计的紧急制动能力。系统设计必须同时回答“如何更便宜地运行”和“如何可靠地停下”。

## AI 产品与应用

### Health in ChatGPT 把个人病历带入日常对话

OpenAI 开始向美国 18 岁以上、已登录 ChatGPT 的 Free、Go、Plus 和 Health。用户可连接 Apple Health、受支持的美国医院病历、One Medical Health；在获得许可后，ChatGPT 可比较历次化验、总结就诊后的变化，并把睡眠、活动、用药和过敏等信息带入普通对话。OpenAI 称，每周已有超过 3 亿人提出健康相关问题，而早期测试中超过 70% 的健康对话发生在独立 Health 空间之外。OpenAI 官方发布

<https://openai.com/index/health-in-chatgpt/>)

真正的产品变化不是新增一个健康聊天入口，而是把个人纵向数据变成可在不同任务中调用的上下文。默认按次请求权限、敏感外发动作的再次确认，以及断开数据源后 30 天内删除同步数据，构成了进入高敏感场景所需的最小控制面。OpenAI 同时明确产品不能替代专业医疗判断；医院覆盖率、错误建议率和用户能否理解授权范围仍需持续验证。

## 模型与技术进展

AMD 与 Cerebras 把推理拆成提示处理和 token 生成两段

AMD 与 Cerebras 宣布联合推理方案：AMD Helios 系统负责提示处理与长上下文推理，Cerebras Wafer-Scale Engine 负责高带宽、低延迟的 token 生成；Cerebras 将在数据中心部署 Helios，并于 2026 年内先通过 Cerebras Cloud 提供服务。Cerebras 发布列表 (<https://www.cerebras.ai/company/press-releases>)，Axios 报道 (<https://www.axios.com/2026/07/23/amd-cerebras-ai-chips>)

这说明不同推理阶段的算力特征正在被单独优化。对深度研究、编码 Agent 等长提示且高输出量的工作负载，异构组合可能比单平台端到端运行更有成本和延迟优势。但双方尚未公开统一基准、价格、跨系统传输开销与可用模型范围；在这些数据出现前，它仍是架构方向而非已验证的单位经济性结论。

## 投融资与商业动态

Alphabet 用入口、调用量和云收入同时证明 AI 需求

Alphabet 二季度收入同比增长 24% 至 1,198 亿美元；Google Cloud 收入同比增长 24% 至 248 亿美元。公司同时披露，Gemini App 月活达到 9.5 亿，第一方模型通过直接入口每分钟处理约 220 亿 tokens，近 90% 的《财富》100 强企业使用 Gemini Enterprise。Alphabet 二季度结果转引 (<https://www.publicnow.com/view/97E1A5C62A498212FF33E>) AP 财报报道 (<https://apnews.com/article/google-gemini-950-million-monthly-users-alphabet-e>) 采用数据核验 (<https://www.techrepublic.com/article/google-gemini-950-million-monthly-users-alphabet-e>)

新增变量在于三种尺度首次形成闭环：消费者入口接近十亿级、开发者调用继续增长、云业务直接承接企业 AI 支出。对其他 AI 厂商而言，单报注册用户或模型请求已不够，市场会更要求把使用规模转化为收入增长、毛利和资本效率。Alphabet 的季度利润还受 SpaceX 股权收益显著影响，因此不能把净利润跃升全部归因于 AI。

## 政策、伦理与安全

美国两党议员提出前沿模型“紧急制动”立法

民主党众议员 Ted Lieu 与共和党众议员 Nathaniel Moran 提出 AI Kill Switch 法案，拟授权国土安全部长在 AI 系统可能对生命或经济造成灾难性伤害时，要求提供商限制访问、禁用特定能力或关闭系统；受约束企业还需预先保有相应技术能力。另有六名两党议员推动最强模型接受由商务部认可机构实施的独立安全审计。Reuters 报道 (<https://m.interesting.com/news/world-news/ai-kill-switch-bill-fl-4809204?ampMode=1>) Ars Technica 条款梳理 (<https://arstechnica.com/news/ai-kill-switch-act-would-let-trump-admin-shut-down-ai-systems/>)

两项提案都由 OpenAI 与 Hugging Face 公开的评测安全事件触发，但目前只是法案提案，不是现行义务。商业影响在于，高能力模型供应商可能需要证明限流、能力隔离、密钥吊销、模型下线 and 恢复流程真实可执行；同时，触发阈值、行政权边界、误停损失与跨云部署如何联动，必须在审议中被明确，否则“可停机”本身会成为新的运营与权力风险。

## X 平台高信号

1 .

信号标题：Google 用 1,500 万次交互量化“广覆盖、浅自动化”

类型：硬采用数据

来源或链接：Google ATLAS v1.0 (<https://blog.google/innogy/research/understanding-the-ai-economy/>)

核心观点：ATLAS 汇总 Gemini App、AI Mode 与 Gemini API 的 150 个国家、140 种语言、800 个职业和 4,000 项任务。工作使用已触及代表美国 90% 就业的职业，但典型岗位只在约 21% 的任务中使用 AI，完全自动化的工作交互不足 10%。

为什么重要：它把“职业是否使用 AI”与“岗位中多少任务被改变”分开，纠正了用用户覆盖率直接推导自动化程度的做法。

影响：企业应优先寻找任务级高频协作点，并分别统计辅助、半自动和全自动占比；ATLAS 仅观察 Google 产品，不能直接外推到全市场。

验证状态：样本范围、21% 任务占比和不足 10% 完全自动化均由 Google 官方报告核验；分类方法仍由 Google 自行制定。

2 .

信号标题：AMD 以 2GW 订单和最高 50 亿美元投资绑定 Anthropic

类型：算力采购与资本关系变化

来源或链接：AMD 官方发布 (<https://www.amd-id.com/amd-dan-armitraan-strategis-untuk-menerapkan-gpu-amd-instirgawatt/>)

核心观点：Anthropic 将部署最高 2GW 的 MI450 系列 GPU，首个 1GW 计划上半年上线；AMD 承诺未来向 Anthropic 投入最高 50 亿美元股权资金，双方还将用 C

ude 优化 Instinct 工作负载和 ROCm。

为什么重要：模型公司获得第二条大规模 GPU 路线，AMD 则通过客户订单、软件协作和资本投入同时争夺 NVIDIA 的生态优势。

影响：采购方获得硬件议价与负载路由选择，但投资方同时成为供应商，使真实芯片需求、融资支持和收入确认更难分离。

验证状态：容量、时间表、工程合作和投资上限已由 AMD 核验；实际部署量、投资节奏和交易条件尚未披露。

3 .

信号标题：加拿大把 Agent 行为记录纳入 AI 透明度咨询

类型：政策咨询范围变化

来源或链接：加拿大政府公告 (<https://www.canada.ca/en/innovative-economic-development/news/2026/07/government-of-canada-on-on-ai-transparency.html>)

核心观点：加拿大于 7 月 23 日启动至 9 月 23 日的公众咨询，议题不仅包括识别 AI 生成内容和告知用户正在与 AI 交互，还包括严重事件报告，以及更好地追踪 AI Agent 的活动与交互。

为什么重要：透明度要求开始从“内容是否由 AI 生成”延伸到“Agent 采取了哪些动作”，直接触及企业工作流的日志、追责和取证能力。

影响：面向加拿大市场的 Agent 产品应提前保留动作时间线、工具调用、授权来源和异常事件记录；具体义务尚待咨询后的政策决定。

验证状态：咨询时间、五类议题和参与对象已由加拿大政府核验；当前不是已生效规则。

4 .

信号标题：GitHub Models 完成退役前第二次中断演练

类型：平台退役与迁移信号

来源或链接：GitHub 退役通知 (<https://github.blog/changelog/models-is-being-fully-retired-on-july-30-2026/>)

核心观点：GitHub 按计划于 7 月 23 日实施退役前第二次短时 brownout；7 月，Model playground、模型目录、推理 API 与 BYOK 端点将对所有客户停止服务

为什么重要：这不是功能预告，而是迁移窗口已进入最后一周，仍依赖相关端点的生产 workflow 面临确定性中断。

影响：开发团队需要立即清点端点、密钥、配额与回退路径，并迁移至 Microsoft Foundry、Copilot 或其他推理服务；不能把 brownout 恢复误判为服务会继续保留。

验证状态：退役范围、7 月 23 日演练和 7 月 30 日截止日已由 GitHub 核验；实际 brownout 对各客户的影响范围未公开。

5 .

信号标题：ChatGPT Health 默认逐次请求病历使用权限

类型：高敏感数据访问规则变化

来源或链接：OpenAI Health 隐私说明 (<https://openai.com/index/health/>)

核心观点：连接的病历、Apple Health 数据及使用这些数据的对话不用于基础模型训练或广告；系统默认每次使用前询问，用户也可改为始终允许。断开数据源后，同步数据将在 30 天内从 OpenAI 系统删除。

为什么重要：高敏感数据产品的竞争点被落实到授权粒度、用途限制和删除时限，而非笼统的“隐私优先”表述。

影响：医疗、保险和员工健康应用可把逐次授权与敏感外发确认作为最低产品要求；对话历史中的既有信息仍需用户另行删除，数据生命周期并非一键完全清除。

验证状态：训练用途限制、默认授权方式和删除时限已由 OpenAI 核验；独立隐私审计结果尚未随发布公开。

## 前沿研究速递

### 生成式 AI 通过供给规模稀释图书市场

做了什么：研究者对 2023—2026 年 Amazon 上 14,419 本自出版类型小说做全文预测，并与截至 2026 年 6 月的日销售记录匹配；样本中的书均未披露是否含 AI 生成内容。arXiv (<https://arxiv.org/abs/2607.20349>)

新在哪里：研究发现，检测到超过 25% AI 文本的书销量份额低于其目录份额，却在持续争夺畅销排名。期间每季度有销量的书数量增长 19.2 倍，收入仅增长 8.9 倍，显示 AI 可以靠供给规模而非平均质量改变市场。

潜在应用方向：出版平台的 AI 内容披露、推荐与排名政策，版权诉讼中的市场影响分析，以及作者和出版社的品类供给监测。

一句话判断：生成成本趋近于零后，真正被稀释的可能不是内容质量，而是每个作品获得注意力与收入的概率；结论仍需经同行评审并检验 AI 文本识别误差。

### AI 供应链中的许可证义务大量丢失

做了什么：研究追踪 232,270 条“数据集—模型—应用”链路，连接 Hugging Face i tHub，衡量未声明许可证被下游补上，以及许可证类别在再分发中被替换的情况。arXiv (<https://arxiv.org/abs/2607.20300>)

新在哪里：62.3% 的链路至少经过一个未声明许可证的制品；带义务的各类许可证端到端保留率均低于 7%，而宽松许可证类别达到 95.1%。问题集中在少量基础数据集，并会沿供应链放大。

潜在应用方向：模型物料清单、许可证继承检查、模型与数据平台的发布门禁，以及企业开

源 AI 采购审计。

一句话判断：模型卡上的最终许可证不能证明上游权利链完整，企业需要对每一层来源做可追溯核验。

### 为模型有害输出概率建立严格下界

做了什么：研究将 Clopper-Pearson 置信区间用于估计给定提示下的有害输出概率，并利用潜在空间特征优先探索更可能产生伤害的自回归生成分支。arXiv ( <https://arxiv.org/abs/2607.20286> )

新在哪里：方法重点不是寻找几个越狱样本，而是给出经证明不高于真实风险的概率下界；即使真实有害概率很低，也尝试减少盲目采样成本。

潜在应用方向：前沿模型发布前安全认证、低概率高损失风险测试、版本回归比较和高风险行业采购验收。

一句话判断：安全评测若能从“是否发现失败”进步到“失败概率至少有多高”，将更接近可审计的风险决策，但仍需扩展到更丰富的提示分布与真实部署条件。