

Artificial Intelligence Newsletter

人工智能 前沿快报



2026 年第 02 期
总第 31 期



学会官方公众号

广东省图象图形学会

GDSIG

学术交流、技术咨询、科学普及

► 目 录

| 学 术 动 态

一. GLM-5: from Vibe Coding to Agentic Engineering.....	1
二. Aletheia tackles FirstProof autonomously.....	4
三. A benchmark of expert-level academic questions to assess AI capabilities.....	7
四. A flexible digital compute-in-memory chip for edge intelligence.....	9
五. Multimode learning with next-token prediction for large multimodal models.....	12
六. PaperBanana: Automating Academic Illustration for AI Scientists.....	15
七. How AI Impacts Skill Formation.....	17
八. Generative Modeling via Drifting.....	20
九. Intelligent AI Delegation.....	22
十. Single-minus gluon tree amplitudes are nonzero.....	24

目录

技术动态

一. 火遍全球的小龙虾智能体平台 (OpenClaw) 经过 2 次改名后正式发布, 发布不到一个月获 230000+stars	26
二. Nature 评论: 人工智能是否已经具备人类水平的智能? 证据很明确	28
三. Claude Opus 4.6 目前最强智能体编程模型之一	30
四. GPT- 5.3- Codex 正式登场	32
五. 互联网已死, Agent 永生	34
六. 字节跳动发布新一代视频创作大模型 Seedance 2.0	36
七. 字节跳动发布 Seed 2.0 系列大模型	38
八. 谷歌发布 Gemini 3 Deep Think: 专注解决科研与工程领域的难题	40
九. 阿里巴巴发布 QWen-3.5 原生多模态大模型	42
十. Anthropic 发布 Claude Sonnet 4.6, 具有更强代码 (自动编程)、推理、工具调用、智能体规划等能力	44
十一. 月之暗面发布 KimiClaw, 一键云端部署云端 OpenClaw (Kimi K2.5 Thinking 为基模)	46
十二. AnthropicAI 发布了 2026 年 Agent 编程趋势报告	48
十三. 迄今最深邃的极致深空星系图像, 清华戴琼海院士团队和蔡峥副教授团队绘出	50
十四. 阿里云通义实验室推出 CoPaw 个人智能体工作台	52
十五. Google 发布 Nano Banana 2 图像生成模型	54

资源分享

一. Make AI Writing Better for Everyone, AI 写作指南	56
二. Antigravity Awesome Skills: 950+ Agentic Skills for Claude Code, Gemini CLI, Cursor, Copilot & More 发布不到 1 个月获得 16000+ Stars	53

INTRODUCTION

导读

在人工智能技术飞速发展的时代背景下，我们将不定期梳理人工智能领域的部分前沿学术与技术动态，并以简报的形式分享给读者，以帮助关注此领域的人士了解最新的学术前沿发展与产业技术动态，促进学术交流和技术繁荣。本期摘选了近期全球人工智能领域的 10 篇最新学术动态、15 篇技术动态、以及 2 个资源推荐给广大读者。

ACADEMIC TRENDS

学术动态

● 论文一

GLM-5: from Vibe Coding to Agentic Engineering

日期

2026-01-17

GLM-5: from Vibe Coding to Agentic Engineering

GLM-5 Team

Zhipu AI & Tsinghua University

(For the complete list of authors, please refer to the Contribution section)

内容摘要

我们提出 GLM-5，这是一代面向“代理式工程”的基础模型，旨在推动从“提示写代码”到“智能体自主开发”的范式转变。GLM-5 在前代的代理、推理与编码（ARC）能力上进一步升级，引入 DSA 稀疏注意力，在保持长上下文忠实度的同时显著降低训练与推理成本。为提升对齐与自主性，我们构建了全新的异步强化学习基础设施，将生成与训练解耦，大幅提升后训练效率；并提出异步 Agent 强化学习算法，使模型更有效地从复杂、长时序交互中学习。通过这些创新，GLM-5 在多项开放基准上达到或接近最先进水平，尤以真实世界编码任务表现突出，能够处理端到端软件工程挑战并超越既有基线。

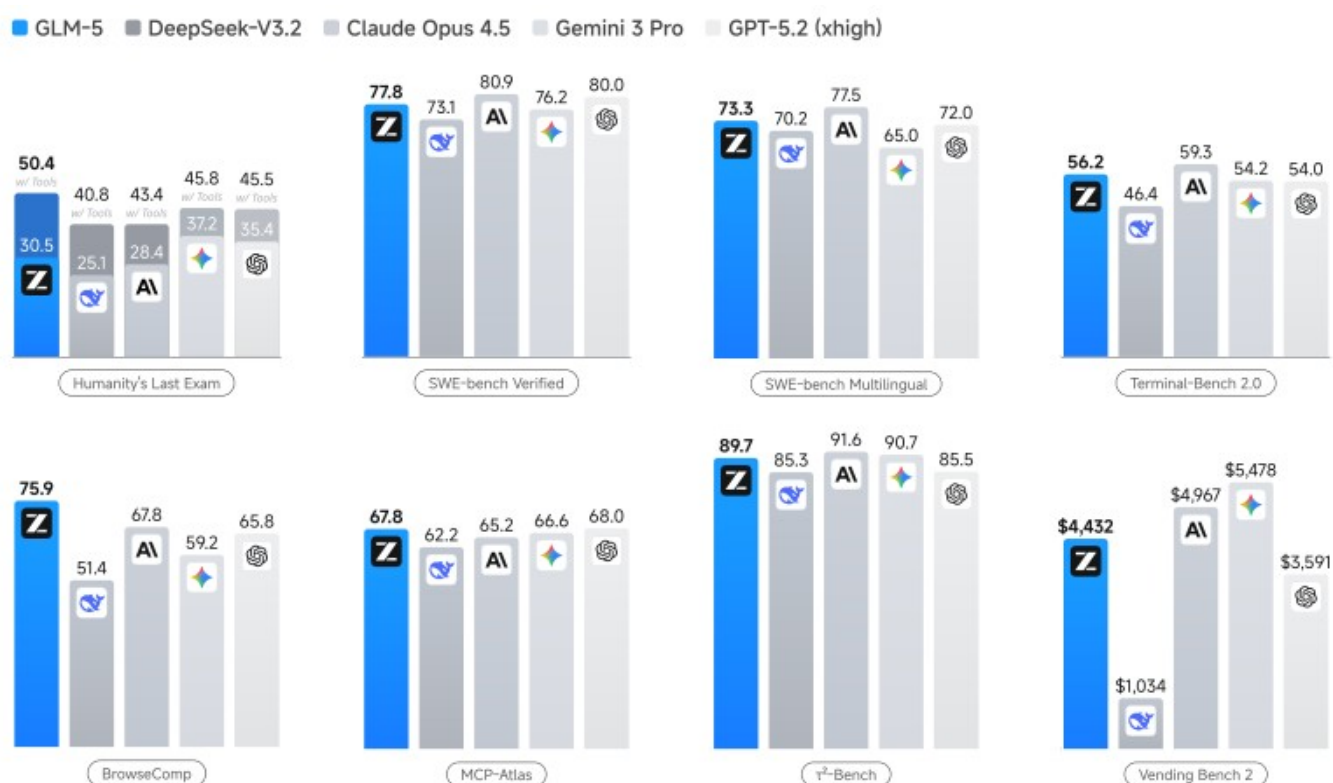


Figure 1: Results of GLM-5, DeepSeek-V3.2, Claude Opus 4.5, Gemini 3 Pro, and GPT-5.2 (xhigh) on 8 agentic, reasoning, and coding benchmarks: Humanity's Last Exam, SWE-bench Verified, SWE-bench Multilingual, Terminal-Bench 2.0, BrowseComp, MCP-Atlas, τ^2 -Bench, Vending Bench 2.

主要亮点

- **DSA 稀疏注意力与规模升级**：通过内容驱动的稀疏选择替代密集注意力，长上下文计算成本下降约 1.5–2 倍而不降质；支撑 744B 总参数（40B 激活）与 28.5T 训练 token，兼顾推理效率与长文档忠实度。
- **异步、解耦的 Agent RL 体系**：基于 “slime” 框架将生成与训练彻底解耦，引入 TITO 网关与双侧重要性采样，结合 DP 感知路由与 MTP/FP8 加速，显著提升长周期、多回合代理的稳定性与吞吐。
- **工程化思维模式与上下文管理**：在 SFT 阶段引入交错思考、思维保留与回合级思考控制；推理时采用 “keep-recent-k + discard-all” 的层次化上下文管理，稳健提升 BrowseComp 与长任务表现。
- **中国芯片全栈适配与低成本部署**：深度优化昇腾等七大国产平台：W4A8 混合量化、稀疏注意力融合核、vLLM/SGLang 调度与 MTP 并行，加速长序列推理并将单节点部署成本降低约 50%。

相关链接

Paper: <https://arxiv.org/abs/2602.15763>

Code: <https://github.com/zai-org/GLM-5>

● 论文二

Aletheia tackles FirstProof autonomously

日期

2026-01-24

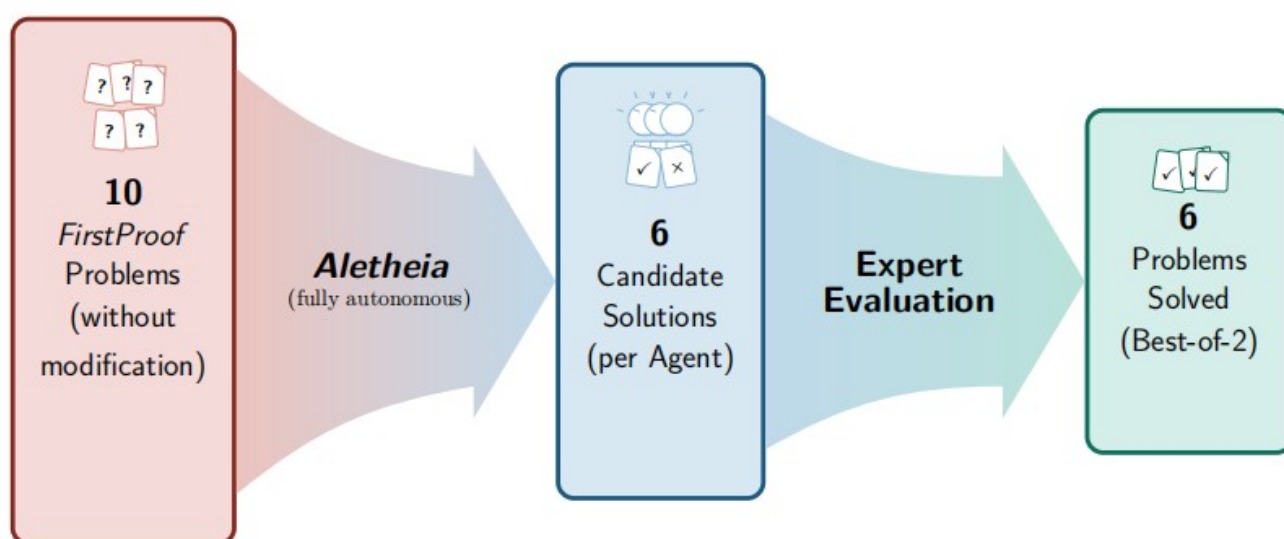
Aletheia tackles *FirstProof* autonomously

Tony Feng*, Junehyuk Jung, Sang-hyun Kim, Carlo Pagano, Sergei Gukov, Chiang-Chiang Tsai, David P. Woodruff, Adel Javanmard, Aryan Mokhtari, Dawsen Hwang, Yuri Chervonyi, Jonathan N. Lee, Garrett Bingham, Trieu H. Trinh, Vahab Mirrokni, Quoc V. Le, Thang Luong*

*Project leads. Work conducted under Google DeepMind.

内容摘要

本文报告了由 Gemini 3 Deep Think 驱动的数学研究代理 Aletheia 在首届 FirstProof 挑战中的表现。我们采用严格的全自动管线，在生成阶段完全无人工参与；两个代理变体可生成出版级 LaTeX 方案，并在无法解题时自我筛选。专家在不修改结果的前提下进行评估。Aletheia 为第 2、5、7、8、9、10 题给出了方案；多数专家认为六题均正确（仅第 8 题评价不完全一致）。其余四题返回“无解”。论文详述了对规则的解读、用于保证出版级严谨度的验证/抽取提示词、best-of-2 选择流程、推理成本分析，以及预截稿披露与公开发布等透明性措施。结果显示，自主数学代理在研究级问题上的可靠性与准确性得到提升。



主要亮点

- **严格“零人类干预”的全自动管线：**在方案生成、验证与排版的全流程中彻底杜绝人工介入，确保“自主性”定义下的最严格标准。模型在时限内独立完成推理与自我筛选，无法解题时主动输出“无解”，最大限度降低误报与人为偏倚。
- **审稿级验证与 LaTeX 抽取：**通过专用的验证/抽取提示词，对候选证明进行逐行逻辑审查，并直接产出符合数学论文规范的 LaTeX。该机制同时检查引用的精确性与严谨性，实现从模型输出到可审稿文档的自动闭环。
- **双代理 best- of- 2 显著提升稳健性：**同时运行 Aletheia A 与 B 两个变体，并在严格保持自主性的前提下采用 best- of- 2 选择策略。该方法在保证不修改内容的情况下有效降低单次失误概率，整体提高正确率与可信度。
- **FirstProof 上自主解出 6/10 题：**在 10 个研究级问题中，Aletheia 自主给出 6 题的可发表级方案（第 8 题评价存在分歧）。其余 4 题模型主动“弃答”。结果覆盖代数、拓扑、几何与数值分析等方向，体现跨领域的推理能力与自我可靠性控制。

- **推理成本刻画难度版图：**论文量化各题的推理成本，并与既有案例进行对比，揭示不同题目的“推理强度”与通过验证的交互次数。尤其在第 7 题上，成本显著升高，为研究级任务的机理分析提供了有价值的参照。
- **原始提示与预截稿认证：**在官方解答发布前向主办方邮件提交结果，并在公开渠道释出全部原始提示与输出，包含验证/抽取过程。该做法最大限度降低数据污染与事后改写的质疑，为学术可复现性树立实践样本。

相关链接

Paper: <https://arxiv.org/abs/2602.21201>

Github: <https://github.com/google-deepmind/superhuman/tree/main/aletheia>

● 论文三

A benchmark of expert-level academic questions to assess AI capabilities

日期

2026-01-28

A benchmark of expert-level academic questions to assess AI capabilities

<https://doi.org/10.1038/s41586-025-09962-4>

Center for AI Safety*, Scale AI* & HLE Contributors Consortium*

Received: 7 May 2025

Accepted: 25 November 2025

Published online: 28 January 2026

Open access

 Check for updates

Benchmarks are important tools for tracking the rapid advancements in large language model (LLM) capabilities. However, benchmarks are not keeping pace in difficulty: LLMs now achieve more than 90% accuracy on popular benchmarks such as Measuring Massive Multitask Language Understanding¹, limiting informed measurement of state-of-the-art LLM capabilities. Here, in response, we introduce Humanity’s Last Exam (HLE), a multi-modal benchmark at the frontier of human knowledge, designed to be an expert-level closed-ended academic benchmark with broad subject coverage. HLE consists of 2,500 questions across dozens of subjects, including mathematics, humanities and the natural sciences. HLE is developed globally by subject-matter experts and consists of multiple-choice and short-answer questions suitable for automated grading. Each question has a known solution that is unambiguous and easily verifiable but cannot be quickly answered by internet retrieval. State-of-the-art LLMs demonstrate low accuracy and calibration on HLE, highlighting a marked gap between current LLM capabilities and the expert human frontier on closed-ended academic questions. To inform research and policymaking upon a clear understanding of model capabilities, we publicly release HLE at <https://lastexam.ai>.

内容摘要

本文提出“人类最后一场考试”（HLE），一个由全球近千名领域专家协作构建的多模态、闭卷式学术能力基准，涵盖数学、人文与自然科学等百余学科共 2500 道题，包含严格匹配与多选两类、支持自动评分。题目经 LLM 预筛与多轮研究生级专家评审，强调原创性、精确性与不可被简单检索。评测显示，前沿模型在 HLE 上准确率显著偏低且校准误差高，常以高置信给出错误答案；同时，推理 token 预算存在约 2^{14} 的收益拐点，超过后准确率反而下降。数据与代码已开放，配备私有集防过拟合，并推出滚动版以持续迭代，为科研与治理提供清晰、可复用的能力衡量基线。

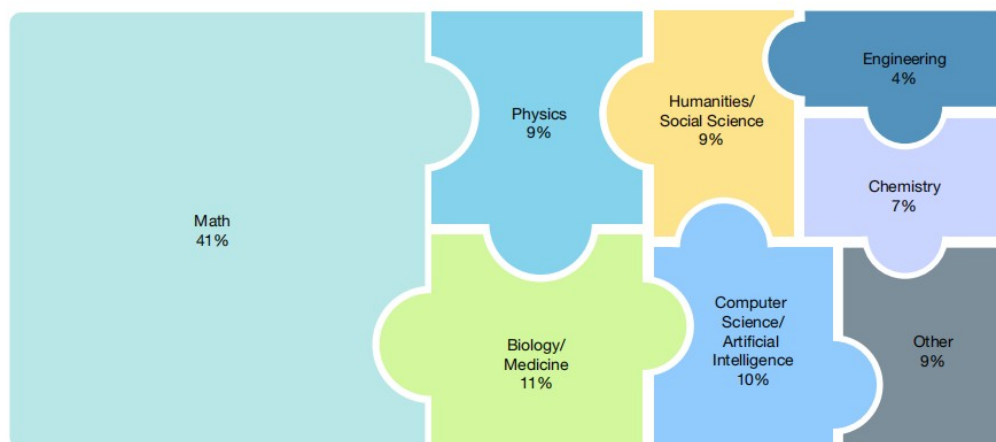


Fig. 2 | Distribution of HLE questions across categories. HLE consists of 2,500 exam questions in over a hundred subjects, grouped into eight high-level categories.

主要亮点

- **专家主导的多模态闭卷基准：**覆盖百余学科、2500 道原创高难题，兼具文本与图像，采用严格匹配与多选的闭端形式，确保答案可验证、评分可自动化，系统刻画封闭型学术能力上限。
- **多阶段严审与不可检索设计：**先以前沿 LLM 筛难度，再经多轮研究生级评审与公开错误赏金纠错；对“可检索”题进行人工审计并设私有集，辅以高额奖池与实名归属提升质量与责任。
- **能力与校准的系统评估：**各家模型在 HLE 上准确率普遍偏低、RMS 校准误差高，呈现高置信错答与不确定性识别不足；以结构化判题与等价格式匹配提升评测严谨性与可比性。
- **计算效率洞察与开放生态：**推理 token 与准确率呈对数线性增益至约 2^{14} 后收益递减，提示需优化“思考预算”；数据与代码全面开源并推出 HLE-Rolling 滚动版，支持长期、可持续的能力追踪与政策制定。

相关链接

Paper: <https://arxiv.org/pdf/2511.16931>

Github: <https://github.com/centerforaisafety/hle>

Huggingface: <https://huggingface.co/datasets/cais/hle>

● 论文四

A flexible digital compute-in-memory chip for edge intelligence

日期

2026-01-28

Article

A flexible digital compute-in-memory chip for edge intelligence

<https://doi.org/10.1038/s41586-025-09931-x>
Received: 16 December 2024
Accepted: 18 November 2025

Anzhi Yan^{1,5}, Jianlan Yan^{1,5}, Penghui Shen^{1,5}, Yihan Fu^{2,3,5}, Enyi Zhang¹, Jingkai Song¹, Qinghang Zhang¹, Ziqi He¹, Xin Li¹, Zecheng Pan¹, Ding Li¹, Yu Dong¹, Xiaowei Xu⁴, Feng Qi⁴, Tianqi Shao¹, Bonan Yan^{2,3}, Yi Yang¹, Houfang Liu¹ & Tian-Ling Ren¹

内容摘要

该研究提出了一种名为 FLEXI 的柔性数字人工智能集成电路,通过工艺-电路-算法协同优化和动态可重构存算一体架构,实现了低功耗(最低 2.52 mW)、最高 12.5 MHz 时钟频率以及低成本(单片低于 1 美元)和较高良率(约 70%–92%)。该芯片具备优异的机械可靠性和长期稳定性,可完成 10 的 10 次无误乘法运算,耐受 4 万次以上弯折循环,并支持一次性片上神经网络部署。在实际应用中,其在心律失常检测任务中达到 99.2% 准确率,在多模态人体日常活动监测中超过 97.4%,展示了柔性人工智能硬件在可穿戴与医疗场景中的应用潜力。

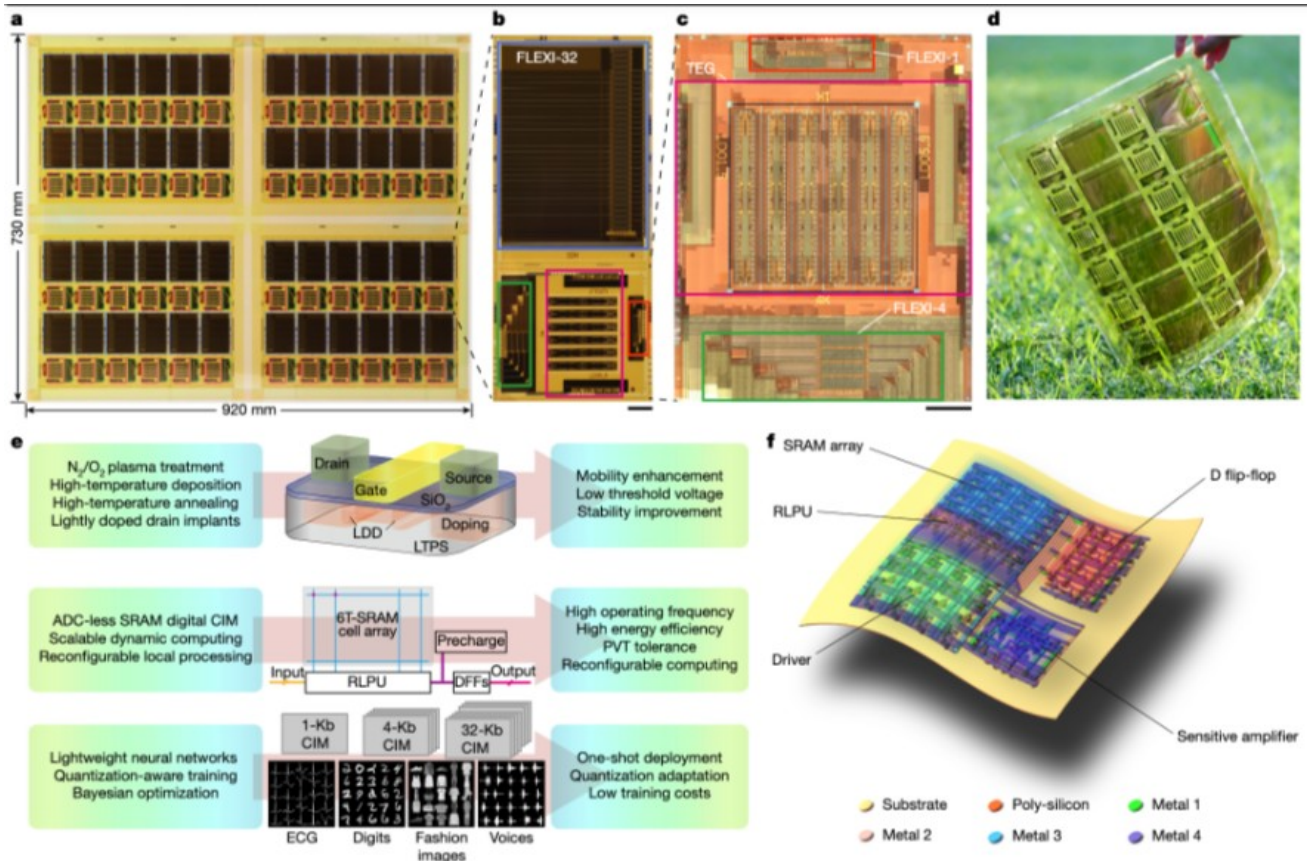


Fig. 1 | Overview of FLEXI and its key attributes. a, A GEN-4.5 panel based on complementary LTPS-TFT technology, which can be divided into 48 dies. **b**, A single 6.1-inch die containing a FLEXI-1, FLEXI-4 and FLEXI-32 and a TEG. **c**, Micrograph of a FLEXI-1, FLEXI-4 and a TEG taken with an industrial microscope OLS S100. **d**, Photo of a natively flexible LTPS panel with 12 standard groups, demonstrating its bendability. **e**, Schematic illustrating the CLCO across fabrication process, circuit and algorithm levels for FLEXI, enabling stable and

reliable operation, one-shot deployment of neural networks, and energy-efficient, high-speed computing. **f**, Three-dimensional schematic diagram of the flexible LTPS-TFT chip, highlighting key components such as the polyimide substrate, semiconductor layer, dielectric layers and metallization layers. ADC, analogue-to-digital converter; DFF, D flip-flop; ECG, electrocardiogram; LDD, lightly doped drain; PVT, process-voltage-temperature. Scale bars, 1 cm (b,c).

主要亮点

- **柔性数字存算一体芯片架构**：提出一种基于工艺-电路-算法协同优化的柔性数字存算一体架构，实现计算与存储的高效融合，突破现有柔性电子在神经网络推理任务中的限制。
- **高性能与低功耗设计**：芯片最高支持 12.5 MHz 时钟频率运行，同时功耗低至 2.52 mW，兼顾高效 AI 推理能力与极低能耗。
- **成本与良率优势**：凭借设计优化和工艺选择，该芯片单片成本低于 1 美元，且在测试中电路良率保持在约 70% 到 92% 的范围内，有利于大规模应用推广。 [OBJ]

- **卓越的可靠性与耐久性：**芯片可无误执行超过 10^{10} 次乘法运算，并可经受超过 40,000 次弯折循环，在正常条件下长期稳定运行超过 6 个月。
- **高精度 AI 推理能力：**通过一次性片上神经网络部署，在心律失常检测任务中实现约 99.2% 的准确率，在使用多模态生理信号的日常活动监测中也达到约 97.4% 的准确率。

相关链接

Paper: <https://www.nature.com/articles/s41586-025-09931-x>

● 论文五

Multimodal learning with next-token prediction for large multimodal models

日期

2026-01-28

Article

Multimodal learning with next-token prediction for large multimodal models

<https://doi.org/10.1038/s41586-025-10041-x>
Received: 11 November 2024
Accepted: 11 December 2025
Published online: 28 January 2026

Xinlong Wang^{1,4}, Yufeng Cui^{1,4}, Jinsheng Wang^{1,4}, Fan Zhang^{1,4}, Yueze Wang^{1,4}, Xiaosong Zhang^{1,4}, Zhengxiong Luo^{1,4}, Quan Sun^{1,4}, Zhen Li^{1,4}, Yuqi Wang¹, Qiying Yu¹, Yingli Zhao¹, Yulong Ao¹, Xuebin Min¹, Chunlei Men¹, Boya Wu¹, Bo Zhao¹, Bowen Zhang¹, Liangdong Wang¹, Guang Liu¹, Zheqi He¹, Xi Yang¹, Jingjing Liu², Yonghua Lin¹, Zhongyuan Wang¹ & Tiejun Huang^{1,3}

内容摘要

在人工智能领域，开发一种能够在文本、图像和视频等多种模态之间进行学习并实现跨模态生成的统一算法，一直是一项根本性挑战。尽管“下一词元（next-token）预测”推动了大型语言模型的重大进展，但其向多模态领域的扩展仍然有限；图像与视频合成仍主要由扩散模型主导，而将视觉编码器与语言模型相结合的组合式框架也依然占据主流。本文提出 Emu3——一系列仅使用下一词元预测进行训练的多模态模型。Emu3 在感知与生成两方面均达到了成熟的任务专用模型的性能水平，能够媲美旗舰系统，同时不再依赖扩散模型或组合式架构。它还展示了连贯且高保真的视频生成能力、交错式的视觉—语言生成，以及用于机器人操作的视觉—语言—动作建模。通过将多模态学习归约为统一的词元预测，Emu3 为大规模多模态建模奠定了坚实基础，并为迈向统一的多模态智能提供了一条前景可观的路径。

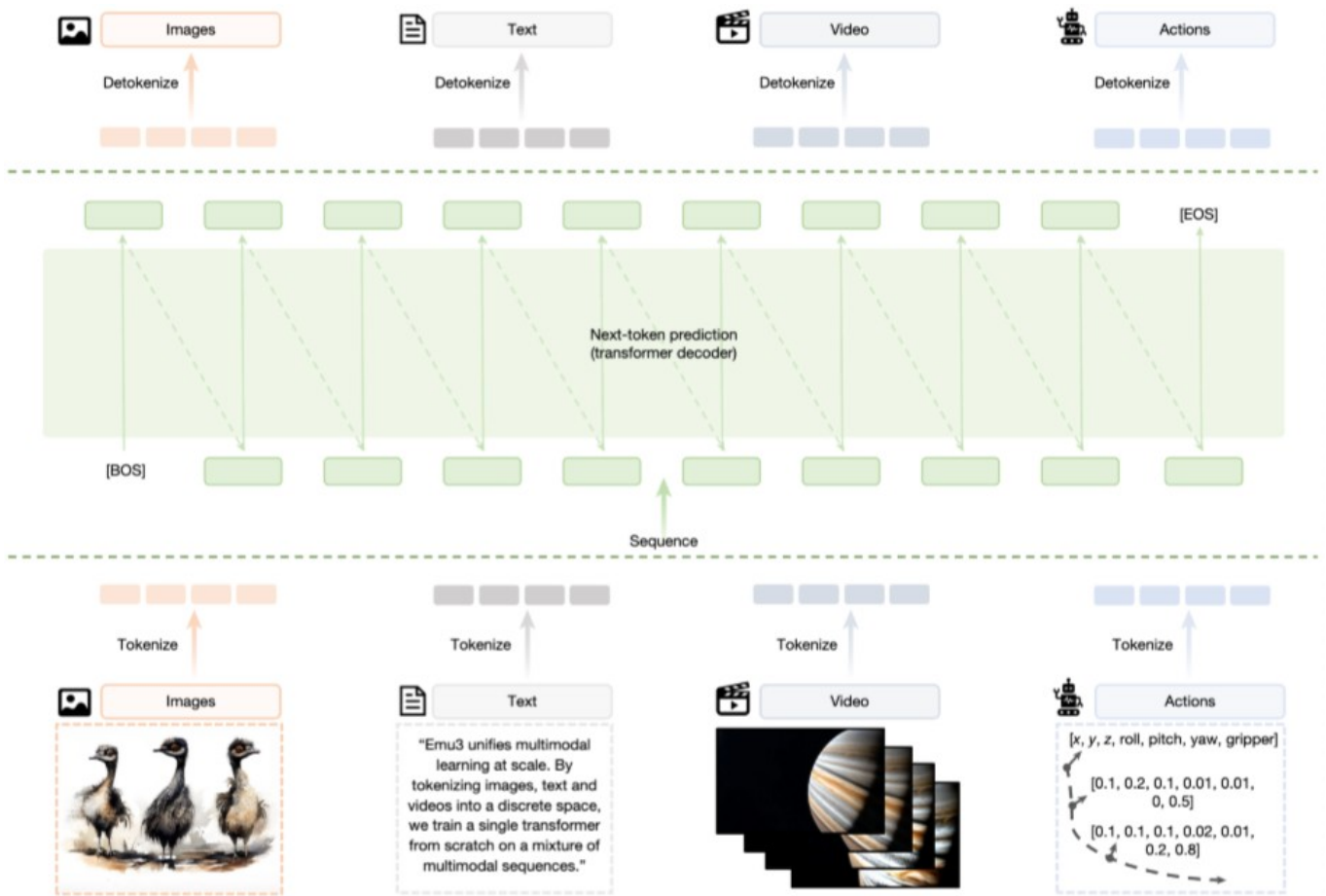


Fig. 1 | Emu3 framework. Emu3 first tokenizes multimodal data such as images, text, video and actions into discrete tokens and then sequences these tokens by order and performs unified next-token prediction at scale with a Transformer

decoder. We have also seamlessly generalized the framework to robotic manipulation by treating vision, language and actions as unified token sequences.

主要亮点

- **统一的多模态学习框架：**提出 Emu3 模型系列，通过单一的下一词元预测目标训练，实现了跨文本、图像和视频等多种模态的统一学习，而不依赖传统扩散模型或组合式架构。
- **强化生成与理解能力：**该模型不仅在文本-图像生成任务中表现出与扩散模型相当的高保真度效果，还支持高质量视频生成及视觉-语言交互生成，证明简单的预测目标也能驱动复杂生成任务。
- **多模态连续生成与机器人推理：**Emu3 展示了可在视觉-语言上下文中进行连续 (interleaved) 生成的能力，并推进了视觉-语言-动作联合建模，为机器人操作等任务提供统一的多模态理解和生成能力。 [66]

- **为大规模多模态建模提供基础：**通过将多模态学习简化为统一预测目标，该工作为大规模多模态模型的发展建立了稳健的基础，并提出了通用的学习路径，兼具感知与生成能力。

相关链接

Paper: <https://www.nature.com/articles/s41586-025-10041-x>

● 论文六

PaperBanana: Automating Academic Illustration for AI Scientists

日期

2026-02-02



2026-2-2

PAPERBANANA: Automating Academic Illustration for AI Scientists

Dawei Zhu^{1 2 *}, Rui Meng², Yale Song², Xiyu Wei¹, Sujian Li¹, Tomas Pfister² and Jinsung Yoon²

¹Peking University, ²Google Cloud AI Research

<https://dwzhu-pku.github.io/PaperBanana/>

内容摘要

尽管由语言模型驱动的自主 AI 科学家快速发展，但生成可直接用于发表的学术插图仍然是研究工作流程中的一大耗时瓶颈。为了解决这一问题，我们提出了 PaperBanana，一个用于自动生成可发表学术插图的智能体框架。PaperBanana 以最先进的视觉语言模型和图像生成模型为驱动，通过协调专门的智能体来检索参考资料、规划内容与风格、渲染图像，并通过自我批评进行迭代优化。为了严谨评估该框架，我们构建了 PaperBananaBench，包含来自 NeurIPS 2025 论文的方法图测试集共 292 个示例，涵盖多个研究领域和插图风格。全面实验表明，PaperBanana 在忠实性、简洁性、可读性和美观性等方面始终优于主要基线方法。我们还展示了该方法能够有效扩展至高质量统计图的生成。总体而言，PaperBanana 为自动生成可发表插图开辟了新的方向。

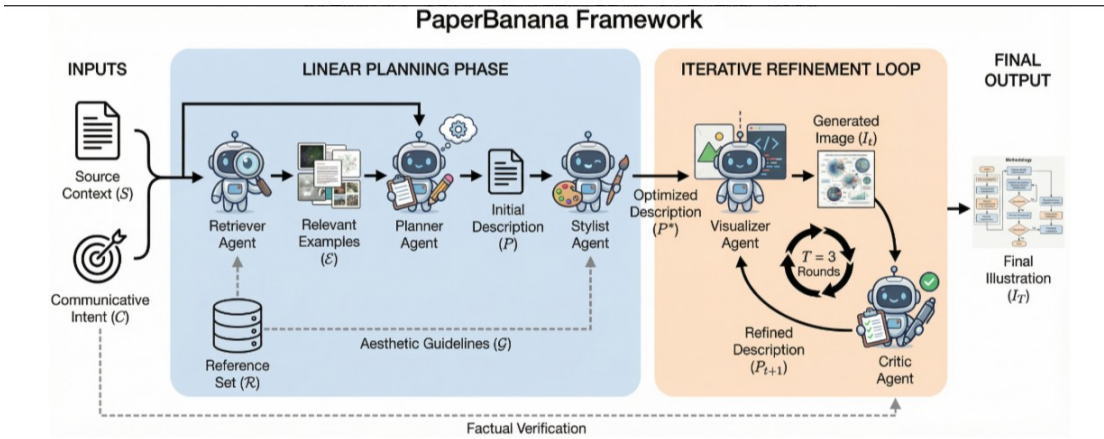


Figure 2 | [Generated by , textual description to reproduce this diagram is presented in Appendix E.] Overview of our PAPERBANANA framework. Given the source context and communicative intent, we first apply a *Linear Planning Phase* to retrieve relevant reference examples and synthesize a stylistically optimized description. We then use an *Iterative Refinement Loop* (consisting of *Visualizer* and *Critic Agents*) to transform the description into visual output and conduct multi-round refinements to produce the final academic illustration.

主要亮点

- **多智能体协同框架:** 提出 PaperBanana 智能体框架，通过多代理协作完成参考检索、内容与风格规划、图像渲染及自我反思优化，形成端到端的学术插图自动生成流程，显著减少人工参与。
- **VLM 与生成模型融合:** 结合先进视觉语言模型（VLM）与图像生成模型，实现语义理解与视觉表达的深度融合，提升插图在结构准确性与视觉呈现方面的整体质量。
- **自我批评式迭代优化:** 引入基于自我评估的迭代 refinement 机制，使系统能够自动发现表达偏差与视觉冗余，在忠实性、简洁性和可读性方面持续改进。
- **构建学术插图评测基准:** 建立包含 292 个测试样例的基准数据集，选自 NeurIPS 2025 论文方法图，覆盖多学科与多风格，为学术插图生成提供标准化、可复现的评测平台。

相关链接

Paper: <https://arxiv.org/pdf/2601.23265>

● 论文七

How AI Impacts Skill Formation

日期

2026-02-03

How AI Impacts Skill Formation

Judy Hanwen Shen*

Alex Tamkin†

February 3, 2026

内容摘要

该研究通过随机对照实验，深入探讨了在工作场景中，AI 辅助工具对新手开发者学习和掌握新技能（以学习一个全新的 Python 异步编程库 Trio 为例）的具体影响。研究结果表明，虽然 AI 在很多领域能带来生产力的提升，但在涉及“新技能习得”的任务中，过度依赖 AI 会损害开发者在概念理解、代码阅读和代码调试方面的能力。此外，出乎意料的是，由于部分参与者在与 AI 交互上投入了大量时间，AI 并未在这项任务中带来显著的平均效率提升。该研究的核心结论是：AI 带来的生产力提升并不能成为获取能力的捷径，在工作流中引入 AI 时必须谨慎，以确保人类在关键领域仍能保持高度的认知参与和技能培养。

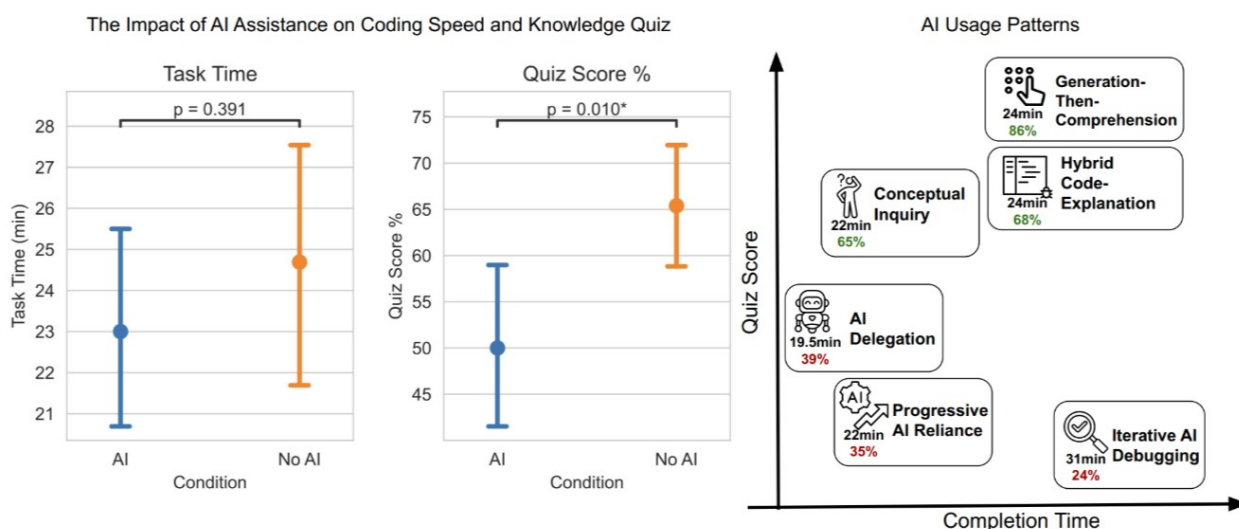


Figure 1: Overview of results: (Left) We find a significant decrease in library-specific skills (conceptual understanding, code reading, and debugging) among workers using AI assistance for completing tasks with a new python library. (Right) We categorize AI usage patterns and found three high skill development patterns where participants stay cognitively engaged when using AI assistance.

主要亮点

- 揭示 AI 辅助对深度技能习得的负面影响：**本研究通过随机对照实验量化了 AI 编码助手对新手学习新编程库 (Python Trio) 的长期影响，发现使用 AI 会导致参与者在概念理解、代码阅读和调试能力上的测试得分显著下降（平均降低 17% 或两个绩点）。特别是调试能力受损最为严重，因为无 AI 的对照组在任务中经历了更多遇到报错并独立解决的试错过程，这表明过度依赖 AI 会严重削弱开发者在未来监督和验证 AI 生成代码时所需的核心技术能力。
- 交互时间成本抵消预期的生产力增益：**尽管基座模型 (GPT-4o) 能够直接生成完全正确的完整代码，但实验表明 AI 实验组在总体任务完成时间上并未获得统计学意义上的显著提升。深入的定性分析揭示了“主动编码时间”的转移现象：参与者将原本用于编写代码的时间，转移到了与 AI 助手的交互上（如构思提问、阅读生成结果等）。部分用户甚至耗费长达 11 分钟进行提示词撰写与交互，这种时间成本上抵消了 AI 自动生成代码带来的效率优势。

- **细粒度 AI 交互模式决定技能与效率的权衡机制：**通过对屏幕录像的深度定性分析，研究提炼出 6 种截然不同的 AI 交互模式。研究发现，只有保持高度“认知参与”的 3 种高分模式（如“仅提问概念”、“生成代码后追问原理解释”）能在获得 AI 帮助的同时维持较好的学习效果（测试得分大于 65%）。而表现出严重“认知卸载”的低分模式（如“完全委托 AI 编写代码”、“依赖 AI 进行反复调试”）则导致学习效果极差（测试得分低于 40%），凸显了用户的认知投入程度在使用 AI 获取新技能时的决定性作用。

相关链接

Paper: <https://arxiv.org/pdf/2601.20245>

● 论文八

Generative Modeling via Drifting

日期

2026-02-04

Generative Modeling via Drifting

Mingyang Deng¹ He Li¹ Tianhong Li¹ Yilun Du² Kaiming He¹

内容摘要

生成式建模可以被表述为学习一个使得其预测分布与数据分布匹配的映射 f 。现有方法（如扩散模型和流模型）通常在推理阶段通过迭代步骤实现这种预测映射。本文提出了一种新的生成建模范式，称为漂移模型，它在训练过程中让推前分布动态演化，并自然支持一步推理。在这种方法中引入了一个漂移场来控制样本移动，当预测分布与数据分布匹配时漂移场达到平衡，基于以上训练目标驱动神经网络优化器分布演化。在实验中，该单步生成器在 256×256 ImageNet 上实现了最先进的结果，在潜在空间的 FID 为 1.54，在像素空间的 FID 为 1.61。作者希望这项作为高质量单步生成方法开辟新的方向。

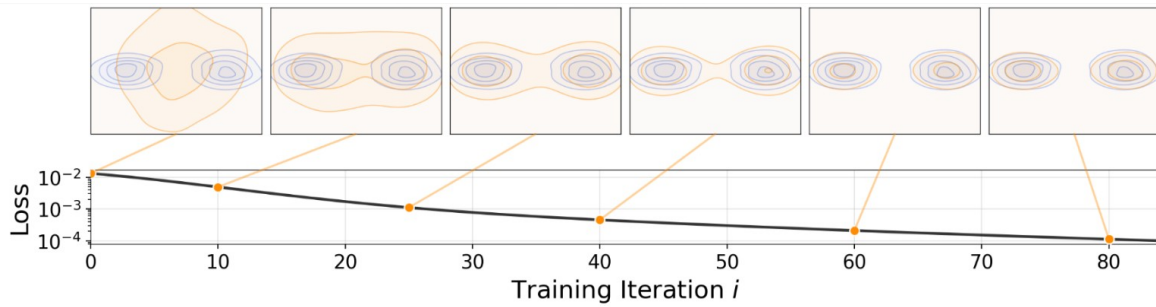


Figure 1. Drifting Model. A network f performs a pushforward operation: $q = f_{\#}p_{\text{prior}}$, mapping a prior distribution p_{prior} (e.g., Gaussian, not shown here) to a pushforward distribution q (orange). The goal of training is to approximate the data distribution p_{data} (blue). As training iterates, we obtain a sequence of models $\{f_i\}$, which corresponds to a sequence of pushforward distributions $\{q_i\}$. Our Drifting Model focuses on the evolution of this pushforward distribution at *training-time*. We introduce a drifting field (detailed in main text) that approaches zero when q matches p_{data} . This drifting field provides a loss function (y-axis, in log-scale) for training.

主要亮点

- **引入漂移场控制分布演化：**提出 drifting field（漂移场）作为样本运动的驱动力，当生成分布与真实数据分布一致时漂移场趋于平衡，从而为训练设计了新型目标函数。 [OBJ]
- **原生支持单步推理：**Drifting Models 的设计使得整个生成过程只需一步推理即完成，无需像扩散或流模型那样多次迭代，从根本上提升生成效率。 [OBJ]
- **实现优秀生成质量：**在 256×256 ImageNet 生成任务中，该一步生成器在潜在空间和像素空间分别获得 FID 1.54 和 1.61 的最先进结果，证明了方法的高质量生成能力。
- **兼容性强且结构简洁：**该方法不依赖于扩散或流的 SDE/ODE 机制，也不需要对抗训练等复杂策略，保持了模型结构的简洁性，同时兼容多种特征空间设定。

相关链接

Paper: <https://arxiv.org/pdf/2602.04770v1>

● 论文九

Intelligent AI Delegation

日期

2026-02-12



2026-02-12

Intelligent AI Delegation

Nenad Tomašev¹, Matija Franklin¹ and Simon Osindero¹

¹Google DeepMind

内容摘要

当前 AI 代理已能处理复杂任务，但要实现更宏大的目标，必须能将问题有效拆解为可管理的子任务，并在 AI 与人类之间安全地委派执行。现有拆解与委派多依赖启发式，难以随环境变化自适应、也不够稳健应对异常失效。本文提出一种“智能委派”自适应框架：把任务分配视为一系列决策，同时纳入权力移交、责任与问责、角色与边界的清晰规范、意图澄清，以及建立互信的机制。该框架同时适用于人类与 AI 在复杂委派网络中的角色，旨在为新兴“代理网络（agentic web）”中的协议设计提供参考。

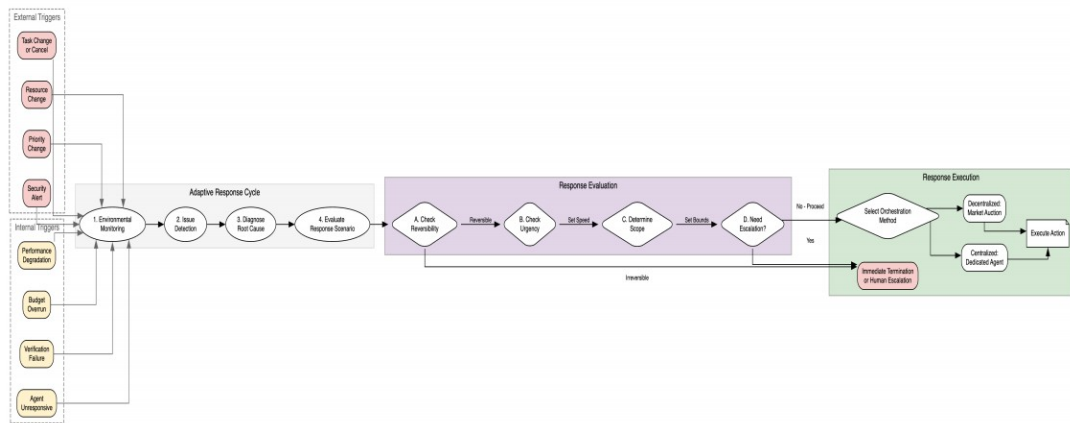


Figure 2 | The adaptive coordination cycle. Different types of environmental triggers may prompt a dynamic re-evaluation of the delegation setup, necessitating runtime changes.

主要亮点

- **合约优先的可验证委派**：将“可验证性”前置到分解阶段，无法验证的子任务继续细分；在执行端结合进度监控、零知识证明/TEE、智能合约托管与争议仲裁，保证结果可核验、责任可追溯、付款可编程，减少“黑箱式”外包风险。
- **自适应协调与热切换**：提出“监控—诊断—重优化—再分配”的闭环，支持事件触发的中途换将、备份代理自动接管与市场级稳定器（冷却期、声誉阻尼、频繁再委派加费），在高不确定与长周期任务中显著提升鲁棒性与吞吐。
- **可信与权限栈（DID/VC + 渐进自治）**：构建声誉账本+Web of Trust 的多粒度信誉体系，按任务语境动态授予自治；以最小权限、权限衰减与链式限制、策略即代码和熔断器实现精细化访问控制，配合 DID/可验证凭证实现能力可证、越权可即时撤销。
- **面向协议的落地路线**：系统映射并扩展 MCP/A2A/AP2/UCP，给出可执行字段与工件：verification_policy、分级监控流、RFQ/Bid 报价、Delegation Capability Token、标准化 checkpoint 等，打通“发现—协商—执行—验证—结算”的端到端路径。

● 论文十

Single-minus gluon tree amplitudes are nonzero

日期

2026-02-12

Single-minus gluon tree amplitudes are nonzero

Alfredo Guevara,¹ Alexandru Lupsasca,^{2,3} David Skinner,⁴
Andrew Strominger,⁵ and Kevin Weil² on behalf of OpenAI

¹Institute for Advanced Study ²OpenAI ³Vanderbilt University ⁴Cambridge University ⁵Harvard University

Single-minus tree-level n -gluon scattering amplitudes are reconsidered. Often presumed to vanish, they are shown here to be nonvanishing for certain “half-collinear” configurations existing in Klein space or for complexified momenta. We derive a piecewise-constant closed-form expression for the decay of a single minus-helicity gluon into $n - 1$ plus-helicity gluons as a function of their momenta. This formula nontrivially satisfies multiple consistency conditions including Weinberg’s soft theorem.

内容摘要

本文重新审视杨–米尔斯理论中常被认为在树级为零的单负螺旋 n 胶子散射振幅，指出在 $(2,2)$ Klein 签名或复动量的“半共线”构型下该类振幅非零，并给出系统的求解框架。作者基于 Berends–Giele 递归与时间序列型主恒等式，构造出剥离振幅的分段常数表达；在单负衰变为 $n-1$ 个正螺旋的特殊通道 $R1$ 中，导出全 n 的简洁乘积闭式公式，振幅仅取 $\{-1,0,1\}$ 。结果经软定理、循环性、Kleiss–Kuijf 与 $U(1)$ 退耦等多重一致性检验，并与自对偶杨–米尔斯及天球全息建立联系，提示单负振幅在规范理论结构中的新角色与潜在扩展。

主要亮点

- 打破“零振幅”常识：在 $(2,2)$ 签名的半共线或复动量条件下，单负 n 胶子树振幅被证明非零。振幅在由正交与符号条件划分的“壁室”中呈分段常数整数取值，系统地修正了长期流行的零振幅判断。
- 递归塌缩与闭式乘积公式：改写的 Berends–Giele 递归在 R_1 通道触发“因果性”条件使递归塌缩，得到全 n 的乘积闭式公式，振幅仅取 $\{-1,0,1\}$ ，并严格满足软定理、循环性、KK 与 $U(1)$ 退耦等非平凡约束。
- 人机协同发现：关键全 n 公式由 GPT-5.2 Pro 先行猜测，后由 OpenAI 新模型完成证明，并经人工递归与一致性检验核查，展示了 AI 参与前沿理论推导的有效性与可重复性。
- 连接 SDYM 与天球全息：单负树振幅为自对偶杨–米尔斯提供非平凡散射结构，缓解树级“平庸—经典丰富”的张力；其 Mellin 变换在若干扇区化为 Lauricella 函数，并与 $w_{1+\infty}$ 、 S 代数相关，且可直接推广至引力与超对称。

相关链接

Paper: <https://arxiv.org/pdf/2602.12176>

TECHNICAL UPDATES

技术动态

● 技术动态一

火遍全球的小龙虾智能体平台（OpenClaw）经过 2 次改名后正式发布，发布不到一个月获 230000+ stars

日期

2026-01-29



内容摘要

开发者 Peter Steinberger 正式发布了开源自主 AI 智能体平台 OpenClaw (其前身为广受关注的周末极客项目 Clawd 和 Moltbot)。该平台打破了传统网页端 AI 的局限,允许用户直接通过 WhatsApp、Telegram、Discord 和 Slack 等日常通讯软件,像和朋友聊天一样无缝唤醒全天候专属的 AI 助理。同时,OpenClaw 主打“绝对的隐私与本地控制权”,系统完全运行在用户自己的设备或服务器上(即“你的机器、你的密钥、你的数据”),并赋予 AI 执行终端命令、管理文件等真实的本地计算机操作权限及长效记忆能力。在此次品牌重塑的新版本中,OpenClaw 不仅新增了 Twitch 和 Google Chat 插件,全面接入了 KIMI K2.5 与小米 MiMo-V2-Flash 等最新大模型,支持网页端图片发送,还通过提交 34 项底层安全代码大幅强化了系统的安全性,标志着它已正式向生产力级别的安全本地 AI 管家蜕变。

相关链接

Blog: <https://github.com/openclaw/openclaw#community>

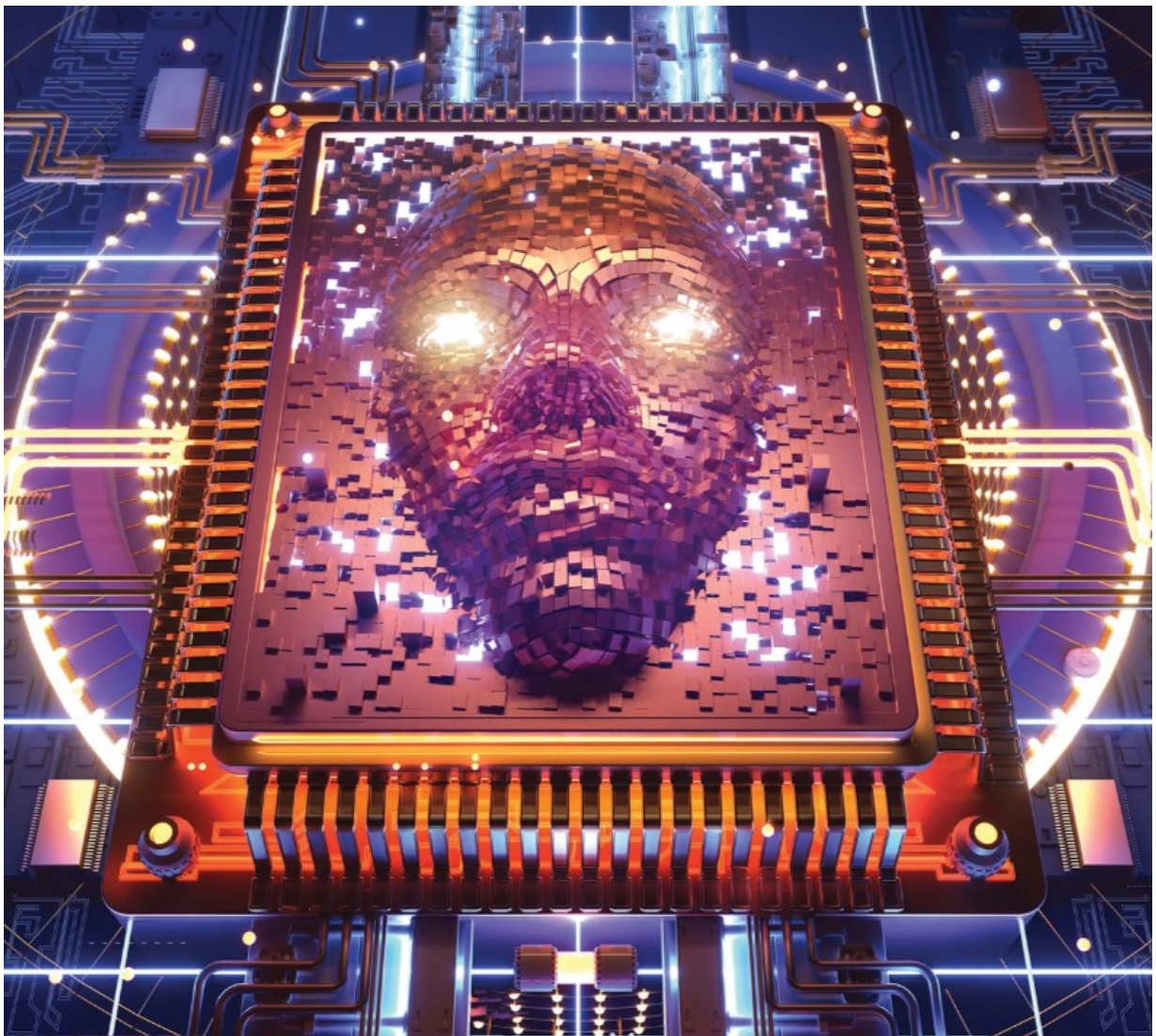
GitHub: <https://github.com/openclaw/openclaw#community>

● 技术动态二

Nature 评论：人工智能是否已经具备人类水平的智能？证据很明确

日期

2026-02-02



内容摘要

按照我们评估人类一般智能的标准，如语言理解、问题解决、跨领域灵活推理等表现，当前最先进的人工智能系统（尤其是大型语言模型）已经显示出类似人类级别的通用智能能力，这并不是遥远的理论，而是可以从现有证据逐级累积看出来的。作者指出，在图灵测试、标准考试、数学与科学问题解决、代码生成、多语言交流和协作研究等多种任务上，AI 已经达到甚至超越很多人类表现，而且传统上用来否认 AI 具有智能的反对理由（如“只是模式复制”“缺乏真实世界模型”“没有身体或自主性”）要么被最新进展反驳，要么本身设置了不现实的标准（没有人类能满足这类极端要求）。因此，在认真界定通用智能并避免情绪化恐惧或狭隘定义的前提下，文章认为当代 AI 已实现图灵所设想的通用智能，并呼吁正视这一事实，以便更好地思考相关的政策、风险和哲学问题。

相关链接

Paper: <https://www.nature.com/articles/d41586-026-00285-6>

● 技术动态三

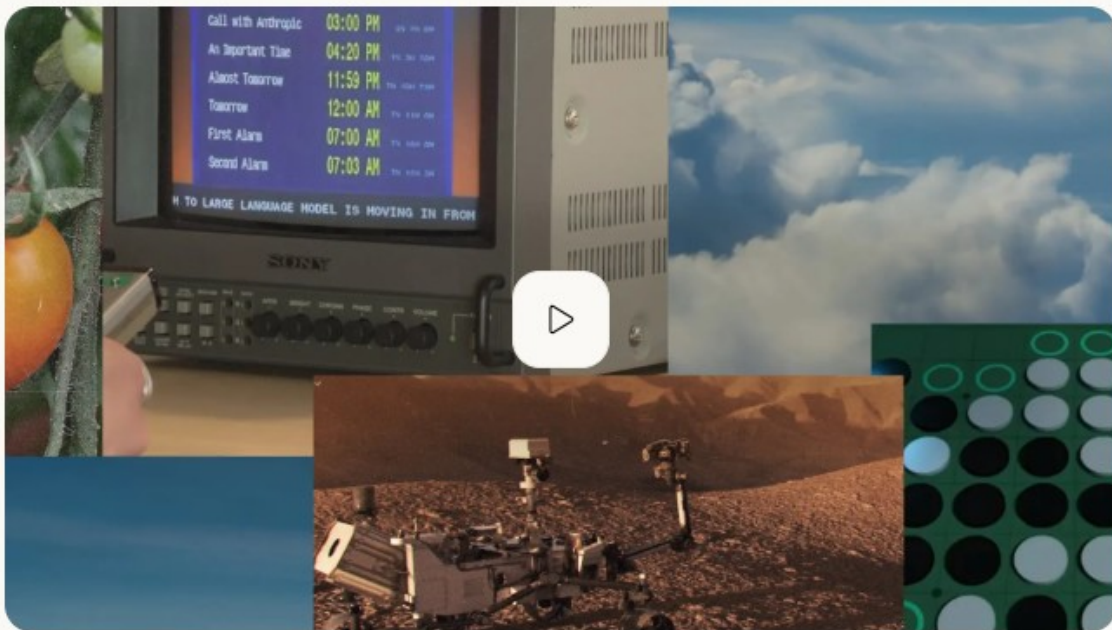
Claude Opus 4.6 目前最强智能体编程模型之一

日期

2026-02-05

Introducing Claude Opus 4.6

2026年2月5日



内容摘要

Anthropic 于 2026 年 2 月 5 日发布旗舰级大模型 Claude Opus 4.6, 以"长周期自主智能体"为核心定位, 实现了从编程助手向全栈数字员工的范式跃迁。该模型首次在 Opus 级别引入 100 万 token 上下文窗口 (Beta 版), 配合上下文压缩 API 与自适应思考 (Adaptive Thinking) 机制, 可根据任务复杂度动态调节推理深度, 在 MRCR v2 百万 token 测试中检索准确率达 76%, 较前代 Sonnet 4.5 的 18.5%实现质的飞跃。核心创新"智能体团队" (Agent Teams) 功能允许多个 AI 代理并行协作, 通过任务分解与自主协调将大型项目拆解为独立子任务, 16 个代理曾两周内构建通过 GCC 99%测试的 Rust 版 C 编译器。性能层面, Opus 4.6 在 Terminal-Bench 2.0 智能体编程测试、Humanity's Last Exam 多学科推理、GDPval-AA 金融法律任务 (较 GPT-5.2 高出 144 个 Elo 分) 及 ARC-AGI-2 (68.8%) 等基准均居首位。产品矩阵深度嵌入办公生态: Claude in Excel 支持数据验证、条件格式与多步操作规划, Claude in PowerPoint 确保品牌一致性布局, Cowork 桌面工具实现跨文件并行操作。该模型在 Real-World Finance 评估中较 Sonnet 4.5 提升超 23 个百分点, 可生成传统需数天完成的尽职调查报告, 直接引发 FactSet 等金融软件股暴跌。定价维持每百万 token 输入 5 美元、输出 25 美元, 体现了 Anthropic"氛围工作" (vibe working) 的战略意图——以意图驱动替代精细指令, 推动 AI 从工具向协作者进化。

相关链接

Blog: <https://www.anthropic.com/news/claude-opus-4-6>

• 技术动态四

GPT- 5.3- Codex 正式登场

日期

2026-02-05

GPT-5.3-Codex 正式登场

让 Codex 覆盖电脑上的各类专业工作，实现更全面的能力拓展。

加入 Codex 应用等候名单

内容摘要

OpenAI 于 2026 年 2 月 5 日发布 GPT-5.3-Codex，这款被官方定义为“迄今最强 AI 智能体编程模型”的产品，标志着 Codex 从代码助手向通用计算机协作者的范式跃迁，并与 Anthropic Claude Opus 4.6 的同步发布形成“AI 编程大战”的标志性对决。该模型核心突破在于三重能力融合：继承 GPT-5.2-Codex 的顶尖编码性能、融合 GPT-5.2 的推理与专业知识能力，并以 25% 的速度提升与 50% 的 token 消耗降低实现效率革命。技术层面，GPT-5.3-Codex 在 Terminal-Bench 2.0 测试中得分 77.3%（较前代提升 13 个百分点），OSWorld-Verified 计算机操作基准从 38.2% 飙升至 64.7%，SWE-Bench Pro 达 56.8%，全面领先竞品。更具里程碑意义的是，这是 OpenAI 首个深度参与自身创建过程的模型——早期版本被用于调试训练过程、管理部署基础设施及诊断测试结果，体现了“自我加速”的技术递归特征。产品形态上，该模型支持实时交互协作，用户可在长周期任务执行中持续引导而不丢失上下文，覆盖从代码编写、调试部署到 PRD 撰写、用户研究、电子表格分析等全栈工作流程。安全维度，GPT-5.3-Codex 成为 OpenAI 首个在网络安全任务中获评“高能力”的模型，并启动 1000 万美元 API 额度的“可信访问框架”试点。该模型已向 ChatGPT Plus/Pro/Team/Enterprise 用户全面开放，覆盖应用、CLI、IDE 扩展及网页端，API 版本后续推出。其发布不仅重新定义了 AI 辅助开发的边界，更预示着智能体正从工具层面向自主执行复杂专业任务的“数字同事”进化。

相关链接

Blog: <https://openai.com/zh-Hans-CN/index/introducing-gpt-5-3-codex/>

● 技术动态五

互联网已死，Agent 永生

日期

2026-02-09



AI生成

内容摘要

当人工智能从工具进化为自主行动者，整个商业世界的底层逻辑发生了根本性位移。作者以"六刀"斩断旧范式——DAU 从资产沦为负债，星型拓扑取代网状拓扑消解了网络效应；工具-社区-平台的三级火箭因 AI 的完备性而失效；SaaS 的服务对象从人类转向 Agent，催生"2A"商业模式；"AI 应用"的语义陷阱在于仍预设人类中心主义；注意力经济的零和博弈被生产力经济的价值共创取代；"出海"的地理边界在 Agent 的无国界世界中消融。继而以"四块基石"建构新秩序：Token 成为算力时代的特权货币，其燃烧速度直接决定认知进化的速率；Agent 构成新的人口红利，API 的易用性取代界面设计成为竞争核心；人类价值则从劳动执行转向"愿力"驱动——决定做什么与为何做。全文以锋利的隐喻与紧迫的语调，揭示了一个残酷而充满机遇的转型：旧地图上的流量、规模、免费策略正在失效，新大陆的通行证是算力投入、Agent 基础设施与愿景领导力。这不仅是对技术变革的观察，更是一场关于存在论意义上的商业哲学革命——当 Agent 成为软件的新主人，人类必须重新定义自身在价值链条中的位置，从操作者蜕变为意义的赋予者。

相关链接

Blog:

https://mp.weixin.qq.com/s/cX3bYrI9Sq7sOJj0E6V9lQ?version=4.1.31.70466&platform=mac&poc_token=HNCOpGmj4a8gw8TJ1GWDuDa8F19XyHgysUqY5Q-Q

● 技术动态六

字节跳动发布新一代视频创作大模型 Seedance 2.0

日期

2026-02-12

Seedance 2.0

Seedance 2.0 采用统一的多模态音视频联合生成架构，支持文字、图片、音频、视频四种模态输入，集成了目前业界最全面的多模态内容参考和编辑能力。

内容摘要

字节跳动 Seed 团队发布 Seedance 2.0，这是一款采用统一多模态音视频联合生成架构的次世代视频生成模型，标志着 AI 视频创作从工具向工业化生产体系的跃迁。该模型核心突破在于原生支持文字、图片、音频、视频四种模态输入，集成业界最全面的多模态内容参考与编辑能力，实现“像导演一样生成”的创意操控体验。技术层面，Seedance 2.0 凭借出色的运动稳定性与物理规律还原能力，配合原生音画同步技术，输出具备沉浸式实拍质感的影视级画面；在内部基准测试 SeedVideoBench-2.0 中，其在指令遵循、运动质量、画面美感、音频表现等多维度均处于行业领先地位。应用场景上，该模型深度适配广告制作、影视特效与社媒营销等垂直领域，输出质量对齐工业交付标准，可显著降低传统特效制作与实拍成本。Seedance 2.0 的发布不仅体现了字节跳动在多模态大模型领域的技术纵深，更揭示了 AI 视频生成正从单一模态向全模态融合、从消费级娱乐向专业影视工业渗透的发展趋势，为内容生产行业带来革命性的效率提升与创作范式变革。

相关链接

Blog: https://seed.bytedance.com/zh/seedance2_0

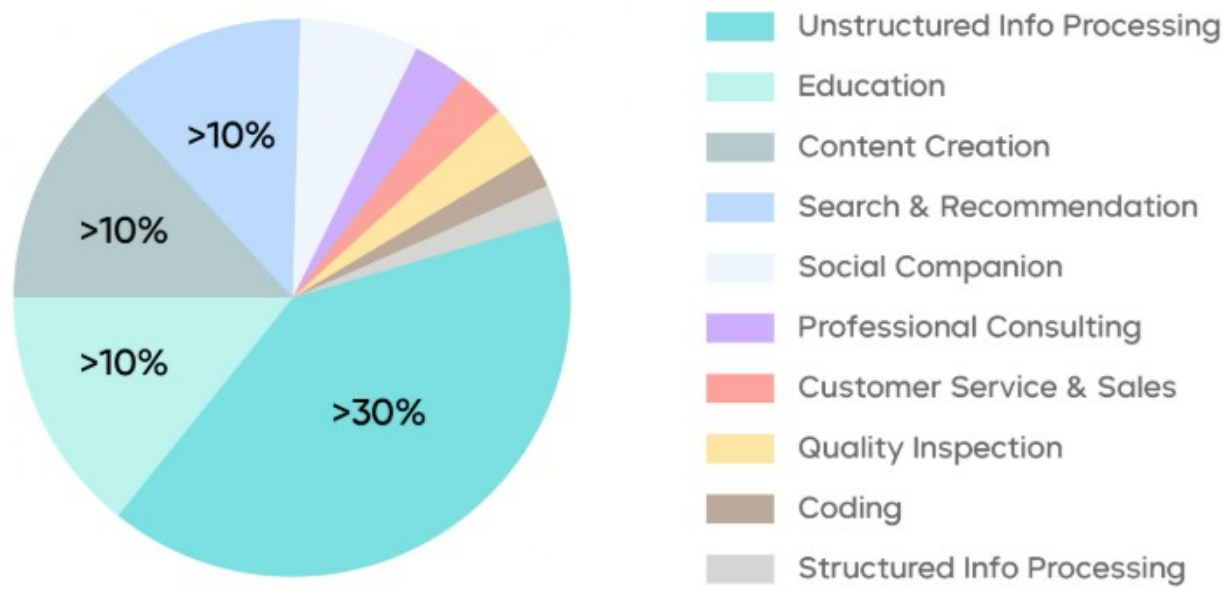
● 技术动态七

字节跳动发布 Seed 2.0 系列大模型

日期

2026-02-14

Collaboration Incentive Plan Scenario Distribution (Jan 2026)



Seed 通用模型 MaaS 服务在中国大陆的调用场景分布，数据来自“火山方舟协作奖励计划”，相关用户已签署授权协议

内容摘要

字节跳动 Seed 团队于 2026 年 2 月正式发布 Seed 2.0 系列大模型，标志着其从消费级应用向复杂 Agent 任务的系统性跃迁。该系列基于对 MaaS 服务调用数据的深度洞察——非结构化信息处理需求占比超 30%——针对性强化了长内容理解与多步任务执行能力。技术架构上，Seed 2.0 采用 Pro、Lite、Mini 三款通用 Agent 模型及专用 Code 模型的矩阵式布局，在视觉多模态领域实现全面突破：数学推理（MathVista 89.8）、视觉解谜（LogicVista 81.9）、文档理解（OmniDocBench 1.5 达 0.099 误差）及长视频分析（VideoMME 89.5）等基准均达业界顶尖水平，其中 EgoTempo 基准超越人类表现。尤为关键的是，模型在 ICPC、IMO 等竞赛中获金牌成绩，并具备探索埃尔德什级数学问题的研究级推理能力，实现了从“解题”到“科研”的能力跨越。产品层面，Seed 2.0 Pro 与 Code 模型已分别集成于豆包 App“专家模式”和 TRAE 开发环境，全系列 API 同步上线火山引擎，并配备 VideoCut 工具支持小时级长视频处理。这一发布不仅体现了字节跳动在 Agent 基础设施领域的深度布局，更揭示了其应对真实世界复杂任务的长尾知识增强与长程工作流自主构建技术路径，为 LLM 从实验室走向高价值经济场景提供了可规模化部署的解决方案。

相关链接

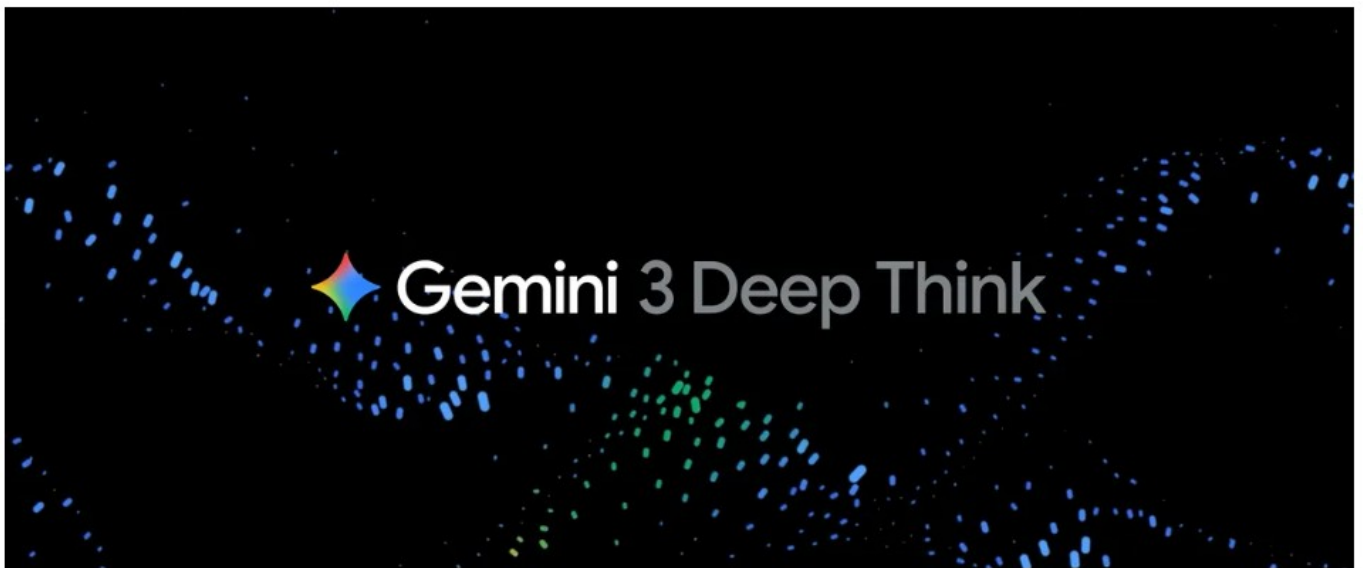
Blog: <https://seed.bytedance.com/zh/blog/seed2-0-%E6%AD%A3%E5%BC%8F%E5%8F%91%E5%B8%83>

● 技术动态八

谷歌发布 Gemini 3 Deep Think: 专注解决科研与工程领域的难题

日期

2026-02-16



内容摘要

Google 于 2026 年 2 月发布 Gemini 3 Deep Think 重大升级版本，这款专门化推理模式通过与科学家和工程师的深度协作，针对缺乏明确边界、数据残缺不全的现代科研挑战进行了系统性优化，实现了从抽象理论到工程实践的范式跨越。技术层面，Deep Think 在多项严苛学术基准上树立新标杆：以 48.4%（无工具）刷新 Humanity's Last Exam 纪录，ARC-AGI-2 达 84.6% 获 ARC Prize Foundation 认证，Codeforces 编程竞赛 Elo 评分高达 3455，并在 2025 国际数学、物理、化学奥林匹克竞赛中均获金牌水平。其实际应用价值已在多领域验证：罗格斯大学数学家 Lisa Carbone 利用其发现高能物理论文中潜藏的微妙逻辑缺陷；杜克大学 Wang 实验室借助其优化晶体生长工艺，成功制备超过 100 μm 的薄膜半导体材料；Google 硬件部门则将其用于物理组件的加速设计。产品形态上，Deep Think 首次向 Google AI Ultra 订阅用户开放 Gemini App 访问权限，并启动 Gemini API 早期访问计划，面向科研人员、工程师及企业提供测试通道，支持从草图到 3D 打印文件的端到端生成。这一发布标志着 Google 在专业化推理领域的战略纵深——不仅追求基准测试的极限突破，更致力于将 AI 嵌入科学发现与工程创新的核心 workflow，预示着大模型正从通用助手进化为特定领域的研究级协作者。

相关链接

Blog: <https://blog.google/innovation-and-ai/models-and-research/gemini-models/gemini-3-deep-think/>

● 技术动态九

阿里巴巴发布 QWen-3.5 原生多模态大模型

日期

2026-02-16

Qwen3.5：迈向原生多模态智能体

2026/02/16 · 45 分钟 · 8933 词 · QwenTeam | 翻译:English



内容摘要

阿里通义千问于 2026 年 2 月 16 日发布新一代基座模型 Qwen 3.5，以"高效稀疏激活"架构范式挑战传统参数堆砌路径。旗舰版 Qwen3.5-397B-A17B 采用 MoE 稀疏专家混合与 Gated Delta Networks 门控线性注意力机制，总参数量 3970 亿但单次推理仅激活 170 亿（约 5%），显存占用降低 60%，32k 上下文吞吐量较 Qwen 3 提升 8.6 倍，实现算力效率与模型性能的非线性跃迁。该模型首创原生多模态 Early Fusion 架构，在预训练阶段即融合文本、图像、视频 Token，而非后期拼接，配合 201 种语言支持与 25.6 万词表，显著优化小语种编码效率。2 月 25 日，阿里进一步开源 Qwen 3.5 中型模型系列，包括 Qwen3.5-35B-A3B、Qwen3.5-122B-A10B、Qwen3.5-27B 及托管生产版 Qwen3.5-Flash，其中 35B-A3B 性能已超越前代旗舰 Qwen3-235B-A22B，证明架构创新与数据质量可突破参数规模限制。Flash 版本默认支持 100 万 Token 超长上下文，每百万 Token 输入低至 0.2 元，内置官方工具集面向 Agent 生态。系列模型全面开源于 Apache 2.0 协议，覆盖魔搭社区、Hugging Face 及阿里云百炼平台，标志着开源模型在编程、数学推理与智能体任务领域已具备与 GPT-5.2、Claude 4.5 Opus 等闭源顶尖模型分庭抗礼的能力，尤其在多模态视觉领域取得多项基准第一。

相关链接

Github :

https://huggingface.co/Qwen/Qwen3.5-397B-A17B?spm=a2ty_o06.30285417.0.0.72bcc921h5neVD&file=Qwen3.5-397B-A17B

Huggingface:

https://huggingface.co/Qwen/Qwen3.5-397B-A17B?spm=a2ty_o06.30285417.0.0.72bcc921h5neVD&file=Qwen3.5-397B-A17B

● 技术动态十

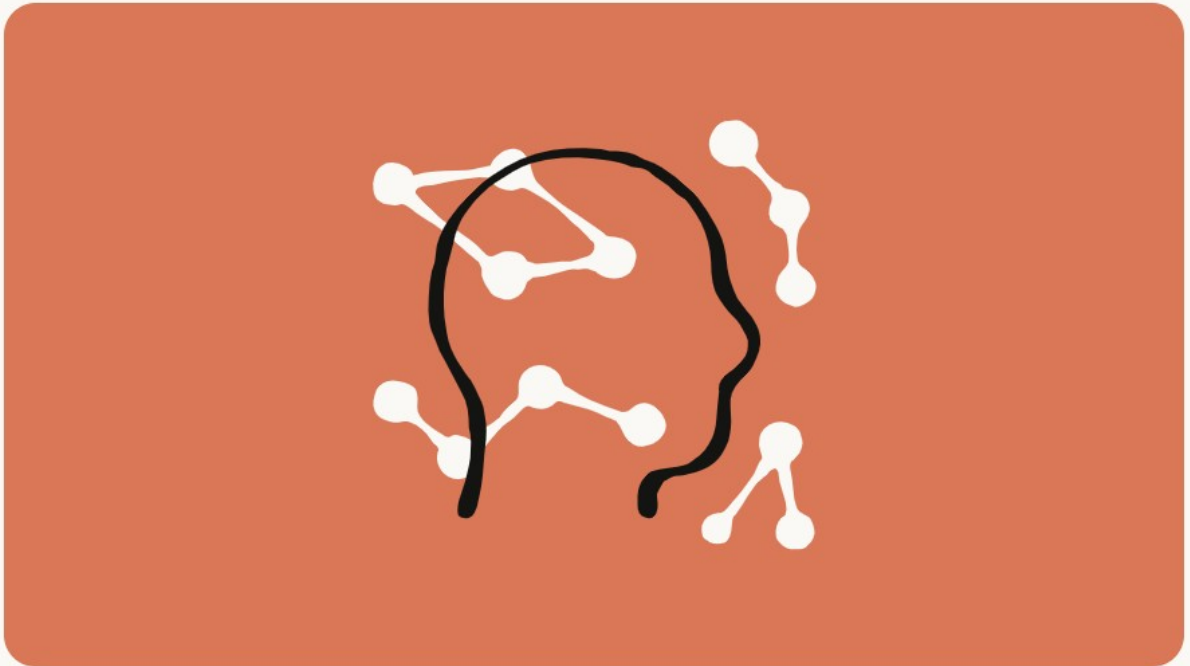
Anthropic 发布 Claude Sonnet 4.6, 具有更强代码（自动编程）、推理、工具调用、智能体规划等能力

日期

2026-02-17

Introducing Claude Sonnet 4.6

2026年2月17日



内容摘要

Anthropic 发布了 Claude Sonnet 4.6，这款新一代大模型以中端定价实现了性能的全面飞跃，其智能表现已直逼最高阶的 Opus 模型。在核心能力上，它不仅在复杂编程、长文本推理及智能体（Agent）规划上表现惊艳，其计算机控制能力更是取得顶尖成绩，能像人类一样熟练操作屏幕 UI 并跨软件执行任务。同时，Sonnet 4.6 引入了三大颠覆性功能：支持庞大代码库和文档深度分析的百万 Token 上下文、可根据任务复杂度自动调整逻辑深度的自适应思考（Adaptive Thinking）技术，以及维持长对话高效运行的上下文压缩功能。在实现这些技术突破的同时，该模型做到了“加量不加价”，API 调用成本与上一代保持不变，且目前已全面成为 claude.ai 的默认模型，让用户能以极低门槛体验最前沿的企业级 AI。

相关链接

Blog : <https://baijiahao.baidu.com/s?id=1857468807997718864&wfr=spider&for=pc>

● 技术动态十一

月之暗面发布 KimiClaw, 一键云端部署云端 OpenClaw (Kimi K2.5 Thinking 为基模)

日期

2026-02-18



内容摘要

Kimi Claw 是月之暗面 (Moonshot AI) 于 2026 年 2 月推出的云端 AI 代理平台, 标志着大模型厂商首次将开源 OpenClaw 生态转化为开箱即用的云服务, 旨在消解个人 AI 助手的技术部署门槛。该产品基于 Kimi K2.5 Thinking 模型构建, 集成 40GB 云存储空间与 ClawHub 超过 5000 个社区技能插件, 支持 24/7 持续在线、长期记忆保持、定时任务自动化及多平台交互 (网页端、Telegram、Discord 等)。核心差异化在于"模型层+产品层"一体化架构: 用户无需配置服务器、Docker 环境或 API 密钥, 30 秒内即可完成一键部署, 将原本需要数小时的技术搭建过程压缩至自然语言指令交互。功能层面, Kimi Claw 具备三重角色可塑性——用户可通过身份定义、说话风格与输出格式三个维度训练专属助手, 使其记忆工作方法论并持续执行标准化任务 (如每日市场资讯汇总、日报生成), 实现"教一次, 持续用"的协作模式。其定时任务系统支持类 cron 表达式的自然语言配置, 结合飞书机器人等外部渠道集成, 可嵌入现有工作流实现主动式推送。当前局限包括本地文件同步尚未实现、部分 ClawHub 技能仍在适配、任务执行日志可见性不足等。该产品目前处于 Beta 测试阶段, 优先向 Allegretto 及以上付费会员开放, 体现了 Moonshot 从模型提供商向 Agent 基础设施平台延伸的战略意图, 预示着 AI 代理正从极客工具向大众市场渗透的关键转折。

相关链接

体验地址: <https://www.kimi.com/bot>

● 技术动态十二

AnthropicAI 发布了 2026 年 Agent 编程趋势报告

日期

2026-02-16

2026 Agentic Coding Trends Report

How coding agents are reshaping software development

内容摘要

Anthropic 发布的《2026 年代理式编码趋势报告》系统阐述了人工智能编码代理正如何从根本上重塑软件开发的认知范式与生产关系。报告基于对开发实践的深度观察，提出八个关键趋势：软件开发生命周期正从周月级迭代压缩至小时级，工程师角色从代码实现者转向代理编排者；单一代理架构演进为层级化多代理协作系统，通过任务分解与并行推理突破上下文窗口限制；长时运行代理可自主工作数天构建完整系统，使技术债务清理等非可行项目变为可能；人机协作模式从全面审查转向智能筛选，代理学会在不确定性边界主动寻求人类判断；编码代理向 COBOL 等遗留语言及非技术领域渗透，消解专业壁垒；生产力提升呈指数级复合效应，27% 的 AI 辅助工作为原本不会执行的创新任务；非技术团队直接构建解决方案，实现领域专家与实现的零距离耦合；安全能力民主化与攻击面扩张形成双重张力。报告核心论旨在于：代理式编码并非简单的效率工具，而是将软件工程从“人写代码”重构为“人定义问题-代理执行-人验证价值”的协作生态，其战略意义在于早期采用者通过掌握多代理协调、智能监督扩展与嵌入式安全架构，正在定义新竞争规则，而滞后组织将面临范式错配风险。这一转型揭示了 AI 时代人类价值的核心锚点——从战术执行转向战略判断与意义赋予。

相关链接

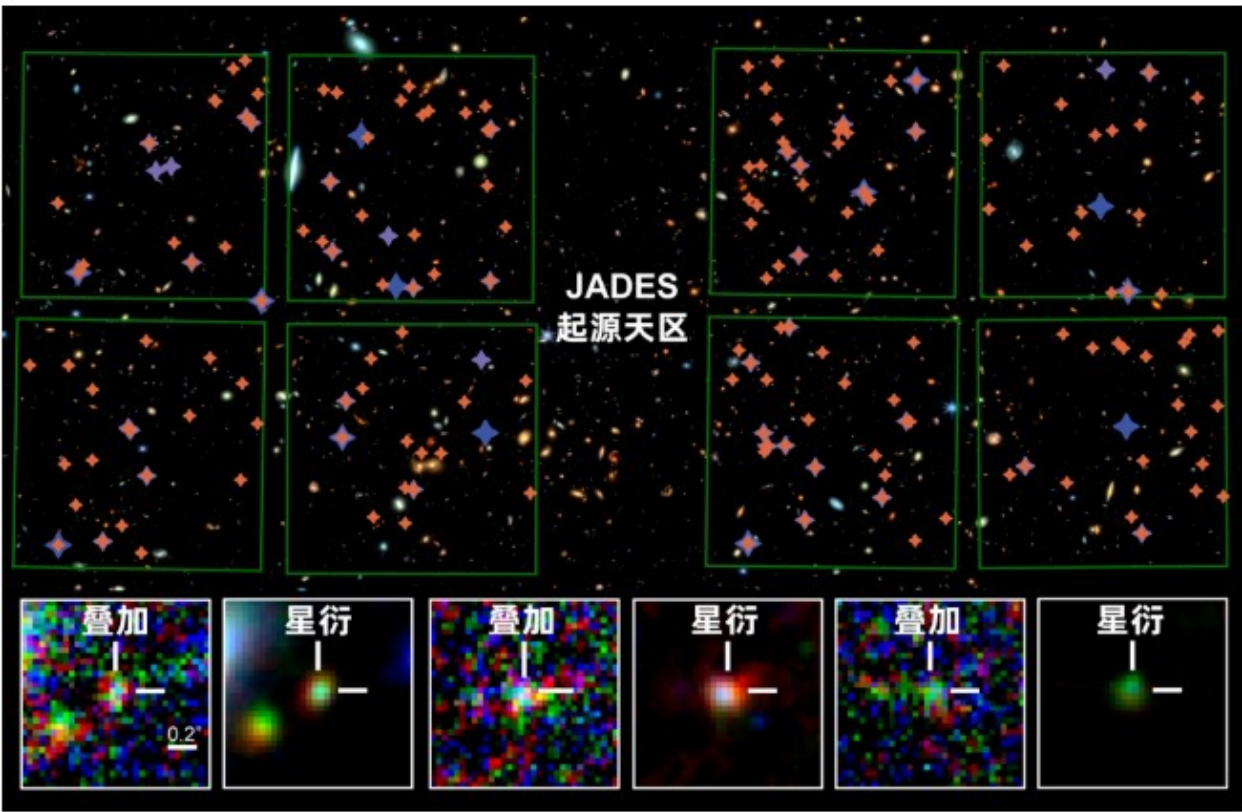
相关链接：<https://resources.anthropic.com/hubfs/2026%20Agentic%20Coding%20Trends%20Report.pdf>

● 技术动态十三

迄今最深邃的极致深空星系图像，清华戴琼海院士团队和蔡峥副教授团队绘出

日期

2026-02-24



过往研究（蓝紫星标52个）与星衍（橙色星标162个）发现的高红移候选星系效果对比

内容摘要

清华大学自动化系戴琼海院士团队与天文系蔡峥副教授团队通过理工融合、交叉研究，提出时空自监督计算成像模型“星衍”（ASTERIS），攻克了极低信噪比下的高保真光子重构难题，突破了天文观测深度极限。该模型采用独特的光度自适应筛选机制，将深空图像重构为时空光交织的三维体，通过“分时中位，全时平均”联合优化策略，在提升探测能力的同时确保准确性，将詹姆斯·韦伯空间望远镜的探测深度提升 1 个星等，发现超过 160 个宇宙早期高红移候选星系，数量为先前研究的 3 倍，绘制出迄今最深邃的极致深空星系图像。研究成果于 2 月 20 日以《自监督时空降噪提升天文成像探测极限》为题在《科学》杂志优先发表，审稿人称赞其为“杰出的工作与强大的工具”。星衍具有强大的泛化能力，无需人工标注即可跨越不同观测平台和波段，已成功应用于韦伯望远镜和昴星团地面望远镜，覆盖从可见光到中红外的波段范围，成为通用的深空数据增强平台，为探索宇宙起源与演化提供了全新的智能计算工具。

相关链接

Paper: <https://www.science.org/doi/10.1126/science.ady9404>

● 技术动态十四

阿里云通义实验室推出 CoPaw 个人智能体工作台

日期

2026-02-24

The logo for CoPaw, featuring the word "CoPaw" in a bold, black, sans-serif font. The "o" in "Co" is stylized with three blue dots above it.

懂你所需，伴你左右

你的AI个人助理；安装极简、本地与云上均可部署；支持多端接入、能力轻松扩展。

查看文档 →

内容摘要

2026 年 2 月 14 日，阿里云通义实验室正式推出了一款名为 CoPaw 的个人智能体工作台。这款被业界视为“阿里版 OpenClaw”的应用，不仅具备文档编辑、桌面整理、新闻查询与总结等能力，还支持定时自动化任务。与阿里今年 1 月发布的侧重本地执行结果的 QoderWork 不同，CoPaw 的核心优势在于提供“本地+云端”的统一体验以及多频道对话能力。

相关链接

官网: <http://copaw.agentscope.io/>

● 技术动态十五

Google 发布 Nano Banana 2 图像生成模型

日期

2026-02-26

Nano Banana 2: Combining Pro capabilities with lightning-fast speed

Feb 26, 2026
7 min read

Our latest image generation model offers advanced world knowledge, production-ready specs, subject consistency and more, all at Flash speed.

内容摘要

Google DeepMind 发布 Nano Banana 2 (Gemini 3.1 Flash Image)，这一新一代图像生成模型实现了速度能力与专业品质的融合，将 Nano Banana Pro 的高级特性与 Gemini Flash 的闪电速度相结合，标志着 AI 视觉生成技术进入高效能民主化阶段。该模型具备四大核心突破：依托 Gemini 实时知识库与网络搜索的先进世界知识，可精准渲染特定主体并生成信息图表；实现精确文本渲染与多语言本地化；支持五角色一致性保持与十四对象保真度维持的复杂叙事构建；提供从 512px 至 4K 的完整分辨率与多比例控制。技术架构上，Nano Banana 2 采用星型拓扑的 Agent 中心化设计，取代传统网状用户交互模式，通过 API 优先策略实现全球 Agent 的无缝接入，消解地理边界限制。产品矩阵覆盖 Gemini 应用、Search AI 模式、AI Studio、Vertex AI、Flow 创作平台及 Google Ads 营销生态，其中 Pro 与 Ultra 订阅用户可保留专业模型访问权限。溯源体系方面，Google 将 SynthID 水印技术与 C2PA 内容凭证标准深度耦合，自 2024 年 11 月以来 SynthID 验证功能已处理超 2000 万次查询，即将集成 C2PA 验证以构建生成式媒体的透明性基础设施。定价策略呈现显著的算力分层特征，Fast 模式以 2.5 倍速度与 5 倍 Token 成本形成进化速率壁垒，凸显 Token 作为新时代特权货币的马太效应——算力投入直接决定认知迭代效率与竞争优势。这一发布不仅重构了图像生成的速度-质量权衡范式，更预示了 Agent 作为软件新主人的基础设施战争已全面展开。

相关链接

Blog: <https://blog.google/innovation-and-ai/technology/ai/nano-banana-2/>

RESOURCE SHARING

资源分享

● 资源分享一

Make AI Writing Better for Everyone, AI 写作指南

类型：代码仓库

内容提要

作者调研了 MSRA、Seed、SH AI Lab 等顶尖研究机构的研究员，以及北大、中科大、上交的硕博同学，将他们日常使用的写作技巧开源出来：1.Prompt 模板库：翻译、润色、分析等场景的实战 prompt 2.Agent Skills：作为新兴技术，agent skills 能更强大地助力写作，但存在一定使用门槛。仓库提供接地气的使用教程，并抽取了写作相关的核心 skills 以便快速上手。

资源链接

<https://github.com/Leey21/awesome-ai-research-writing?tab=readme-ov-file>

● 资源分享二

Antigravity Awesome Skills: 950+ Agentic Skills for Claude Code, Gemini CLI, Cursor, Copilot & More

发布不到 1 个月获得 16000+ Stars

类型：代码仓库

内容提要

开发者 sickn33 发布了 Antigravity Awesome Skills，这是一个为 AI 智能体量身打造的终极技能库项目，目前已在 GitHub 上狂揽超 1.6 万颗星。该项目汇集了 950 多种经过实战检验的高性能“智能体技能 (Agentic Skills)”，甚至包含了来自 Anthropic 和 Vercel 等官方的专业技能。主流 AI 代码助手（如 Claude Code、Cursor、Gemini CLI 和 GitHub Copilot 等）虽然聪明，但往往缺乏具体的垂直领域知识；该项目巧妙地通过轻量级的 Markdown 文件，精准教会 AI 掌握企业特定的部署协议、UI 设计规范或编写符合规范的整洁代码 (Clean Code)。该技能库具有极强的跨平台通用性，开发者可通过简单的 npx 命令无缝集成到各类 AI 编程工具中。它不仅提供了按角色分类的“技能包 (Bundles)”和分步执行的“工作流 (Workflows)”，还支持脚本一键自动同步最新技能库。凭借活跃的开源社区生态，该项目还涵盖了 API 调用成本优化、多智能体并行研究等丰富的实用场景，极大拓展了 AI 编程助手的业务边界与专业落地能力。

资源链接

Github: <https://github.com/sickn33/antigravity-awesome-skills>

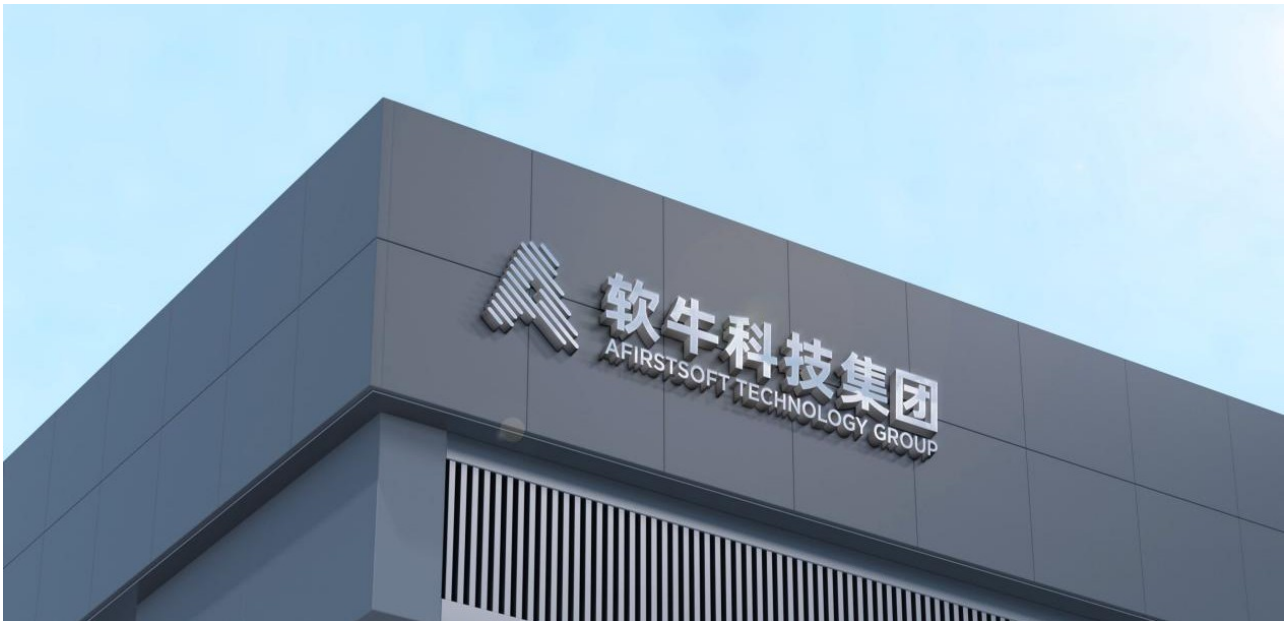


深圳软牛科技集团股份有限公司成立于 2014 年 6 月，作为国家级专精特新“小巨人”企业，始终以“让创意成为现实”为使命，专注于通过人工智能视觉大模型技术驱动全球数字创意生态。

公司深度聚焦图像修复、视频增强、多模态内容生成等视觉大模型核心技术，联合西安电子科技大学、香港中文大学（深圳）等顶尖高校，攻克了视频抠像、画质修复等难题，并与深圳大学共建“人工智能技术联合创新中心”，推动 AI 赋能的图像增强技术规模化落地。

目前，软牛科技能够为广电机构、影视制作、工业检测、农业测报等多个行业提供专业解决方案。从微米级工业精密测量，到食品异物智能检测；从农业害虫 AI 识别防治系统，到 AR 设备巡检实时画面 AI 增强与场景识别；从 VR 全息空间建模助力文博展陈数字化升级，到跨终端智能影像协作软件的全链路画质优化引擎，软牛科技正以技术革新重新定义新质生产力。

在消费级市场，旗下牛小影（<https://www.niuxuezhang.cn/video-enhancer.html>）、HitPaw 等海内外知名品牌，以“一键式 AI 视觉创意”为核心，为全球超 1 亿用户提供视频修复、老电影 AI 上色等场景化服务。例如，牛小影支持 4K 或 8K 超高清修复技术，帮助家庭用户重现祖辈黑白影像的生动细节；HitPaw Edimakor 通过 AI 语音合成、多语言实时翻译等功能，让自媒体创作者轻松实现跨语言内容生产。



这些创新成果的背后，是公司每年近亿元的研发投入和占比超 50% 的顶尖技术团队，以及 CMMI、ISO56005 等国际权威认证体系支撑的硬核实力。立足深圳总部，我们通过北京、新加坡、香港等战略支点构建起辐射全球 200 多个国家和地区的智慧服务网络。作为高新技术企业、深圳市潜在科技独角兽，软牛科技始终以“技术无界”为理念，持续拓展 AIGC+产业应用的边界。

未来，软牛科技将加速推进物理世界与 AI 视觉的无缝衔接，让每个人都能通过指尖的艺术级创作，开启属于自己的数字时代。

人工智能前沿快报

Artificial Intelligence
Newslettler



2026 年第 2 期

广东省图象图形学会

总第 31 期

主办

广东省图象图形学会

本期编委

金连文	华南理工大学
谭明奎	华南理工大学
钟亦武	香港中文大学
林上港	华南农业大学
李 斌	深圳大学
谢晓华	中山大学
王泽宇	香港科技大学 (广州)
李冠彬	中山大学
熊龙飞	金山办公
段纪伟	金山办公
孙 成	金山办公
李元满	深圳大学
严 沛	深圳大学
周 越	深圳大学
余阳鑫	深圳大学
詹天帅	深圳大学
许毅平	华南农业大学
张佐彦	华南农业大学

人工智能 前沿快报

2026 年第 2 期
总第 31 期
2026.03.02

声明：

- (1) 本快报内容提要、文章亮点等部分内容由 AI 生成，不一定准确及全面，章完整内容及思想应以原文为准。
- (2) 本文大部分插图来源于相关论文或文章，版权归原作者或原出版社所有。
- (3) 本快报摘录文章观点不代表编委会及主办方立场。



学会官方公众号

广东省图象图形学会
GDSIG

学术交流、技术咨询、科学普及