

Artificial Intelligence Newsletter

人工智能 前沿快报

2025 年第 11 期
总第 29 期



学会官方公众号

广东省图象图形学会

GDSIG

学术交流、科普宣传、科技成果鉴定

目录

学术动态

一. The Era of Agentic Organization: Learning to Organize with Language Models.....	1
二. Kimi Linear: An Expressive, Efficient Attention Architecture	3
三. Emu3.5: Native Multimodal Models are World Learners.....	6
四. A Survey of Data Agents: Emerging Paradigm or Overstated Hype?.....	9
五. Context Engineering 2.0: The Context of Context Engineering.....	12
六. Large Language Models for Scientific Idea Generation: A Creativity-Centered Survey	15
七. Kosmos: An AI Scientist for Autonomous Discovery.....	18
八. Visual Spatial Tuning	21
九. The Station: An Open-World Environment for AI-Driven Discovery.....	24
十. AlphaResearch: Accelerating New Algorithm Discovery with Language Models.....	26
十一. Omnilingual ASR: Open-Source Multilingual Speech Recognition for 1600+ Languages.....	29
十二. Aligning machine and human visual representations across abstraction levels.....	32
十三. AgentEvolver: Towards Efficient Self-Evolving Agent System.....	35
十四. Depth Anything 3: Recovering the Visual Space from Any Views.....	37
十五. TiDAR: Think in Diffusion, Talk in Autoregression.....	40
十六. Instella: Fully Open Language Models with Stellar Performance.....	43
十七. MonkeyOCR v1.5 Technical Report: Unlocking Robust Document Parsing for Complex Patterns.....	45
十八. Hierarchical Frequency Tagging Probe (HFTP): A Unified Approach to Investigate Syntactic Structure Representations in Large Language Models and the HumanBrain.....	48
十九. MathCanvas: Intrinsic Visual Chain-of-Thought for Multimodal Mathematical Reasoning.....	51
二十. Artificial Hivemind: The Open-Ended Homogeneity of Language Models (and Beyond).....	53
二十一. Gated Attention for Large Language Models: Non-linearity, Sparsity, and Attention-Sink-Free.....	56
二十二. 1000 Layer Networks for Self-Supervised RL: Scaling Depth Can Enable New Goal-Reaching Capabilities.....	58
二十三. Why Diffusion Models Don't Memorize: The Role of Implicit Dynamical Regularization in Training	61

技术动态

一. 月之暗面发布 Kimi-K2, 在 HLE 评测基准上超越现有所有开源闭源大模型.....	64
二. 李飞飞关于空间智能的观点: From Words to Worlds: Spatial Intelligence is AI' s Next Frontier.....	65
三. 研究人员正在培育人类神经元, 并将其转化为生物晶体管.....	66
四. 百度发布文心一言大模型系列新版本.....	67

目录

五. 阿里巴巴发布 QWen-Code 0.2.2.....	68
六. GPT-5.1 发布: A smarter, more conversational ChatGPT.....	69
七. 英伟达文档解析大模型 NVIDIA Nemotron Parse v1.1, 参数不到 1B, 在表格解析、文档版面理解、公式识别等方面性能 SOTA.....	70
八. 分割一切并不够, 还要 3D 重建一切, SAM 3D 来了.....	72
九. 谷歌 Nano Banana Pro 炸了! 硅谷 AI 半壁江山同框, 网友: PS 已死.....	73
十. 国内唯一! 阿里千问斩获 NeurIPS 2025 最佳论文奖.....	74

资源分享

一. Machine-Learning-Systems	75
二. 电子书资源网站整理大全.....	76

鸣谢支持

一. 深圳软牛科技集团股份有限公司.....	77
------------------------	----

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 23 篇最新学术动态、10 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

The Era of Agentic Organization: Learning to Organize with Language Models

日期

2025-10-30

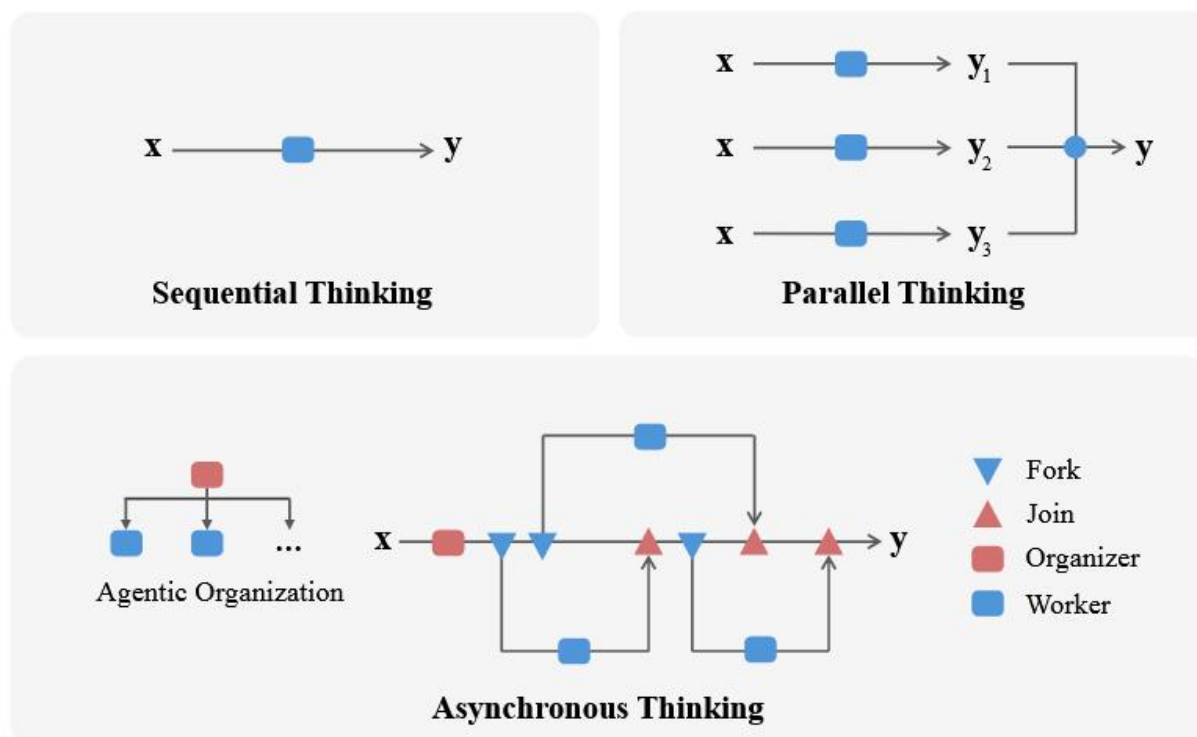
The Era of Agentic Organization: Learning to Organize with Language Models

Zewen Chi Li Dong Qingxiu Dong Yaru Hao Xun Wu Shaohan Huang Furu Wei*
Microsoft Research
<https://aka.ms/GeneralAI>

内容摘要

本文展望了人工智能的一个新时代——“智能体组织”（Agentic Organization），即智能体通过协同和并发工作来解决超出个体智能极限的复杂问题。为实现这一愿景，作者提出了一种新的大语言模型推理范式——“异步思维”（AsyncThink），旨在将模型内部的思维过程组织成可并发执行的结构。具体而言，文章提出了一种思维协议，其中“组织者”（Organizer）动态地向“工作者”（Workers）分配子查询，合并中间知识，并生成连贯的解决方案。更重要的是，该协议中的思维结构可以通过强化学习进一步优化。实验表明，与并行思

维相比，AsyncThink 不仅在数学推理任务上提高了准确性，还实现了 28% 的推理延迟降低。此外，AsyncThink 展现出了强大的泛化能力，能够将其学到的异步思维能力有效应用于未经训练的新任务中。



主要亮点

- 创新的“异步思维”范式与组织者-工作者协议:** 文章提出了一种名为 AsyncThink 的全新推理范式，打破了传统的顺序思维 (Sequential Thinking) 和简单的并行思维 (Parallel Thinking) 模式。基于计算机系统中的多进程概念，作者设计了“组织者-工作者” (Organizer-Worker) 思维协议。在该协议中，通过 Fork (分叉) 和 Join (联结) 操作，单一 LLM 后端能够扮演不同角色：组织者负责动态规划思维路径并分配任务，工作者负责并发执行子查询。这种设计赋予了模型自适应的动态结构，使其能够根据问题复杂度灵活调整思维拓扑，而非依赖预设的固定工作流。
- 基于强化学习的“学会组织”两阶段训练策略:** 为了让模型掌握这种复杂的组织能力，作者通过两阶段过程训练 AsyncThink。首先是“冷启动格式微调”，利用合成数据教会模型掌握 Fork-Join 的语法规则；其

次是核心的“强化学习”阶段，通过引入正确性奖励、格式奖励以及独特的“思维并发奖励”（Thinking Concurrency Reward），鼓励模型在保证答案准确的同时，最大化思维过程的并行度。通过这种方式，模型不再是简单地模仿人类思维路径，而是自主探索并优化出了最高效的思维组织策略（Policy Optimization）。

- **卓越的效能平衡与零样本跨任务泛化能力：**实验结果显示，在准确率和延迟之间，AsyncThink 取得了优于现有方法的平衡（Accuracy-Latency Frontier）。在数学推理任务中，相比并行思维方法，它降低了 28% 的关键路径推理延迟，同时保持了更高的准确率。更为引人注目的是其泛化能力：尽管模型仅在简单的倒计时（Countdown）任务上进行训练，却能将学到的“组织策略”零样本迁移（Zero-shot Transfer）到数独（Sudoku）等全新的复杂任务中，并展现出超越基线的性能。这证明了模型学到的不仅仅是特定任务的解法，而是通用的、高效的思维组织能力。

相关链接

Paper: <https://arxiv.org/pdf/2510.26658>

Project: <https://thegenerality.com/agi/>

● 论文二

Kimi Linear: An Expressive, Efficient Attention Architecture

日期

2025-11-01

KIMI LINEAR: AN EXPRESSIVE, EFFICIENT ATTENTION ARCHITECTURE

TECHNICAL REPORT OF KIMI LINEAR

Kimi Team

 <https://github.com/MoonshotAI/Kimi-Linear>

内容摘要

文章介绍了 Kimi Linear，这是一种全新的混合线性注意力架构。该架构在短上下文、长上下文以及强化学习（RL）扩展等多种场景下的公平对比中，首次在性能上全面超越了传统的全注意力（Full Attention）机制。其核心创新在于 Kimi Delta Attention (KDA) 模块，该模块通过引入通道级（channel-wise）的细粒度门控机制扩展了 Gated DeltaNet，从而更有效地利用有限状态的 RNN 记忆。此外，文章提出了一种定制化的块级并行算法（Chunkwise Algorithm），利用优化的对角加低秩（DPLR）转换矩阵，在保持与经典 Delta 规则一致性的同时，大幅降低了计算成本。通过交替堆叠 KDA 层与多头潜在注意力（MLA）层（比例为 3:1），在 1.4T 与 5.7T token 的训练实验中，Kimi Linear 均取得领先表现：不仅在各项基准测试中优于同等规模的 MLA 模型，还将 KV 缓存占用降低了 75%，并在 100 万上下文长度下实现了高达 6 倍的解码吞吐量提升。为了支持社区研究，作者开源了 KDA 内核、vLLM 推理集成以及预训练和指令微调的模型权重。

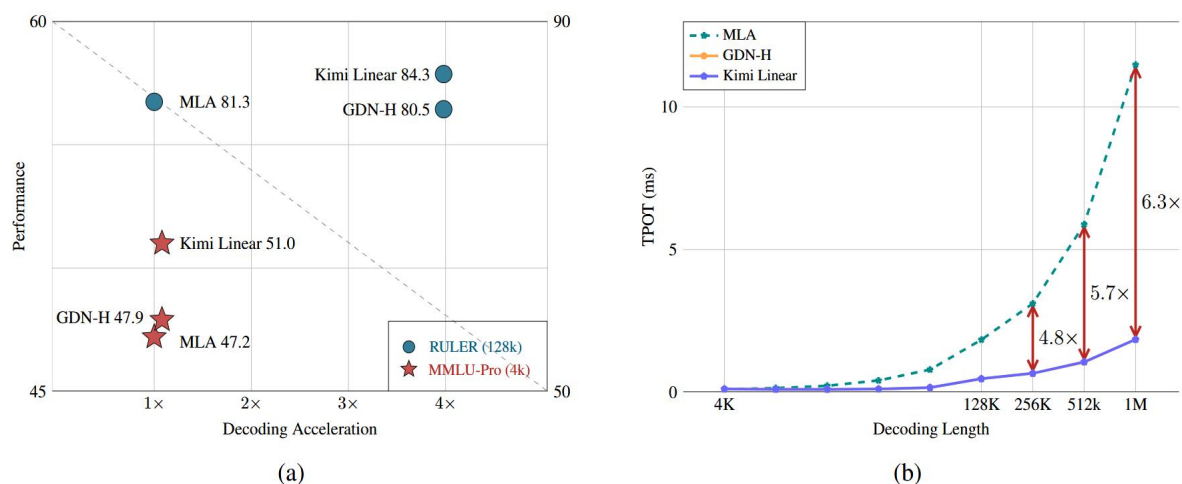


Figure 1: (a) Performance vs. acceleration. With strict fair comparisons with 1.4T training tokens, on MMLU-Pro (4k context length, red stars), Kimi Linear leads performance (51.0) at similar speed. On RULER (128k context length, blue circles), it is Pareto-optimal, achieving top performance (84.3) and $3.98\times$ acceleration. (b) Time per output token (TPOT) vs. decoding length. Kimi Linear (blue line) maintains a low TPOT, matching GDN-H and outperforming MLA at long sequences. This enables larger batches, yielding a $6.3\times$ faster TPOT (1.84ms vs. 11.48ms) than MLA at 1M tokens.

主要亮点

- Kimi Delta Attention (KDA)细粒度门控的线性注意力机制：**KDA 是该架构的核心突破，它改进了现有的线性注意力方法（如 Gated DeltaNet）。不同于 Mamba2 等模型采用粗粒度的头级（head-wise）遗忘门，KDA 创新性地引入了通道级（channel-wise）的细粒度门控机制。这种设计允许每个特征维度拥有独立的遗忘率，使其能够像“可学习的位置编码”一样运作。这种更精确的记忆控制机制显著提升了模型在长序列建模和复杂上下文检索任务中的表现，解决了传统线性注意力模型在精确记忆检索和状态跟踪方面的短板。
- 创新的混合架构与高效的块级并行算法：**Kimi Linear 采用了一种 3:1 的混合架构设计，即每 3 层 KDA 线性注意力层搭配 1 层全注意力（MLA）层。这种设计在保留全局信息流和上下文外推能力的同时，极大优化了推理效率。在底层实现上，团队提出了一种针对硬件优化的定制化块级算法，基于对角加低秩（DPLR）结构，去除了通用 DPLR 中的冗余计算。该算法将算子效率提升了约 100%，不仅大幅减少了

显存占用 (KV Cache 减少 75%)，更在处理 100 万超长上下文时，实现了相比全注意力机制 6 倍的解码速度提升。

- **全方位超越全注意力的性能表现 (Scaling Law & RL)：**这是首个在严格公平的比较下（相同的参数量、训练数据和配方），在各个维度均超越全注意力基线（MLA）的线性架构。实验结果显示，从基础的语言理解到复杂的数学推理和代码生成，Kimi Linear 均表现更优。特别是在长上下文基准（如 RULER）和强化学习（RL）训练阶段，Kimi Linear 展现出比全注意力模型更快的收敛速度和更高的最终得分，证明了线性注意力架构在追求 Agentic Intelligence（代理智能）和测试时计算扩展（Test-time Scaling）方面的巨大潜力。
- **全面开源与无缝集成的基础设施支持：**为了推动混合架构的研究与应用，Kimi 团队不仅发布了基于 MoE 架构的 48B 参数（激活参数 3B）模型权重（包括 4k 上下文预训练版和支持 1M 上下文的指令微调版），还完全开源了核心技术底座。这包括高性能的 KDA CUDA 内核以及与主流推理框架 vLLM 的深度集成代码。这些组件被设计为全注意力管道的“即插即用”替代品，研究人员和开发者无需修改现有的缓存或调度接口，即可直接部署并体验这一高效架构。

相关链接

Paper: <https://arxiv.org/pdf/2510.26583>

Code: <https://github.com/MoonshotAI/Kimi-Linear>

● 论文三

Emu3.5: Native Multimodal Models are World Learners

日期

2025-11-01

Emu3.5: Native Multimodal Models are World Learners

Emu3.5 Team

BAAI

<https://emu.world>

内容摘要

本文介绍了 Emu3.5, 这是一个大规模多模态世界模型, 能够原生地跨视觉和语言预测下一个状态。Emu3.5 在包含超过 10 万亿个 token 的视觉-语言交错数据 (主要来源于互联网视频的连续帧和字幕) 语料库上, 通过统一的 “下一 token 预测” 目标进行了端到端的预训练。该模型自然地接受交错的视觉-语言输入并生成交错的输出。为了增强多模态推理和生成能力, Emu3.5 进一步采用了大规模强化学习进行后训练。为提高推理效率, 论文提出了离散扩散适应 (DiDA), 将逐 token 的解码转换为双向并行预测, 在不牺牲性能的情况下将单张图像的推理速度提高了约 20 倍。Emu3.5 展现了强大的原生多模态能力, 涵盖长视程视觉-语言生成、任意图像生成 (X2I) 以及复杂的富文本图像生成。它还展示了可泛化的世界建模能力, 支持在不同场景和任务中进行时空一致的世界探索和开放世界具身操作。与此同时, Emu3.5 在图像生成和编辑任务上达到了与 Gemini 2.5 Flash Image 相当的性能, 并在多项交错生成任务中表现出色。我们已将 Emu3.5 开源以支持社区研究。

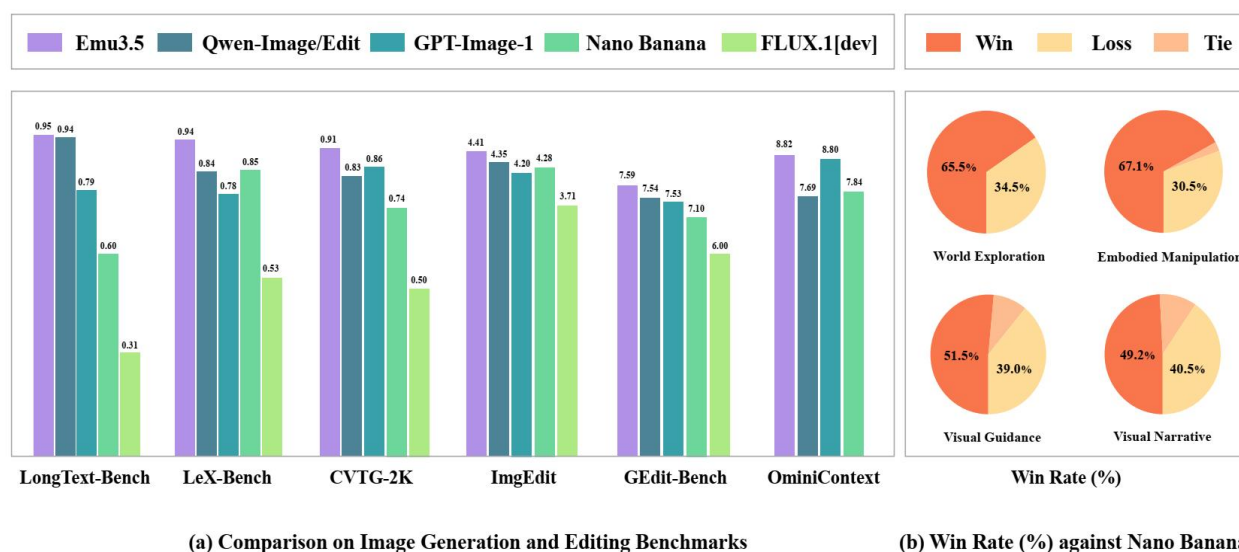


Figure 1: (a) Comparison with SOTA models on image generation and editing benchmarks. For editing, the specific models are Qwen-Image-Edit-2509 [106] and FLUX.1 Kontext [dev] [49]. (b) Automated preference evaluation (Win Rate[%]) against Gemini 2.5 Flash Image (Nano Banana) [92] on interleaved generation tasks.

主要亮点

- 基于视频流的原生多模态世界模型架构：**Emu3.5 采用原生统一架构，在超过 10 万亿 token 的海量数据上进行训练，其中主要包含互联网视频的连续帧和文本。这种设计使得模型能够原生地通过“下一 token 预测”目标来学习，捕捉长视频的视觉-语言序列中的时空一致性和语义连贯性。与拼接独立模型的传统方法不同，Emu3.5 自然地处理交错的输入与输出，不仅具备强大的理解能力，还成为了一个真正的世界模拟器。
- 创新的离散扩散适应 (DiDA) 推理加速技术：**针对自回归模型生成图像速度慢的瓶颈，论文提出了离散扩散适应 (DiDA) 技术。该方法将传统的逐 token 顺序解码转换为基于离散扩散的双向并行预测。在推理阶段引入噪声和去噪步骤，模型能够在保持文本生成能力不变的同时，将单张图像的生成推理速度提高了约 20 倍。这一突破使得 Emu3.5 成为首个在推理速度和生成质量上均能媲美闭源扩散模型的自回归模型。

- **大规模多模态强化学习与统一奖励系统：**Emu3.5 首次在多模态任务中实施了大规模联合强化学习（RL）后训练。研究团队构建了一个包含通用奖励（如美学、图文对齐）、任务特定奖励（如 OCR 准确性、身份保持）的综合奖励系统，并在统一的奖励空间内进行端到端优化。这种方法不仅显著增强了模型的多模态推理和生成能力，还促进了不同任务（如文本生成图像与视觉叙事）之间的相互增益和知识迁移。
- **卓越的可泛化世界建模与具身智能能力：**模型展示了强大的世界建模能力，支持“世界探索”和“具身操作”等复杂任务。Emu3.5 能够根据指令生成时空一致的视觉观察和动作规划，模拟真实的物理规律和动态交互。在具身操作任务中，它能够理解物理法则和物体属性，规划长视频的子任务序列；在世界探索中，它能进行自发的环境导航和连贯的场景合成，性能显著优于 Gemini 2.5 Flash Image。
- **顶尖的任意图像生成（X2I）与富文本渲染：**受益于原生多模态能力，Emu3.5 在任意图像生成（X2I）和文本到图像生成方面表现出众。它支持开放世界的图像编辑，能够精确控制和自由操纵时空内容，并在富文本图像生成（如海报、图表文字渲染）上超越了现有模型。在多个基准测试中，Emu3.5 在图像生成和编辑任务上的表现可与 Gemini 2.5 Flash Image 相媲美，展现了极高的视觉保真度和指令遵循能力。

相关链接

Paper: <https://arxiv.org/pdf/2510.26583>

Github: <https://github.com/baaivision/Emu3.5>

● 论文四

A Survey of Data Agents: Emerging Paradigm or Overstated Hype?

日期

2025-11-01

A Survey of Data Agents: Emerging Paradigm or Overstated Hype?

Yizhang Zhu, Liangwei Wang, Chenyu Yang, Xiaotian Lin, Boyan Li, Wei Zhou, Xinyu Liu, Zhangyang Peng, Tianqi Luo, Yu Li, Chengliang Chai, Chong Chen, Shimin Di, Ju Fan, Ji Sun, Nan Tang, Fuguee Tsung, Jiannan Wang, Chenglin Wu, Yanwei Xu, Shaolei Zhang, Yong Zhang, Xuanhe Zhou, Guoliang Li* and Yuyu Luo*

 **Awesome Data Agents:** <https://github.com/HKUSTDial/awesome-data-agents>

内容摘要

本文综合描述了在大语言模型（LLMs）快速发展背景下出现的数据智能体（Data Agents）。数据智能体旨在编排“Data + AI”生态系统以解决复杂的数据相关任务。针对目前“数据智能体”一词存在的术语歧义和标准不一的问题（混淆了简单的查询回复器与复杂的自主架构），本文受自动驾驶 SAE J3016 标准的启发，提出了首个系统化的数据智能体分级分类法。该分类法包含六个级别（L0-L5），清晰界定了从人工操作（L0）到生成式全自动数据智能体（L5）的自主性演变过程及责任归属。基于此分类法，文章对现有的数据管理、数据准备和数据分析智能体进行了结构化回顾，深入分析了推动智能体进化的关键飞跃及技术差距，特别是从程序化执行（L2）向自主编排（L3）过渡的现状。最后，文章提出了迈向主动式、生成式数据智能体的未来研究路线图。

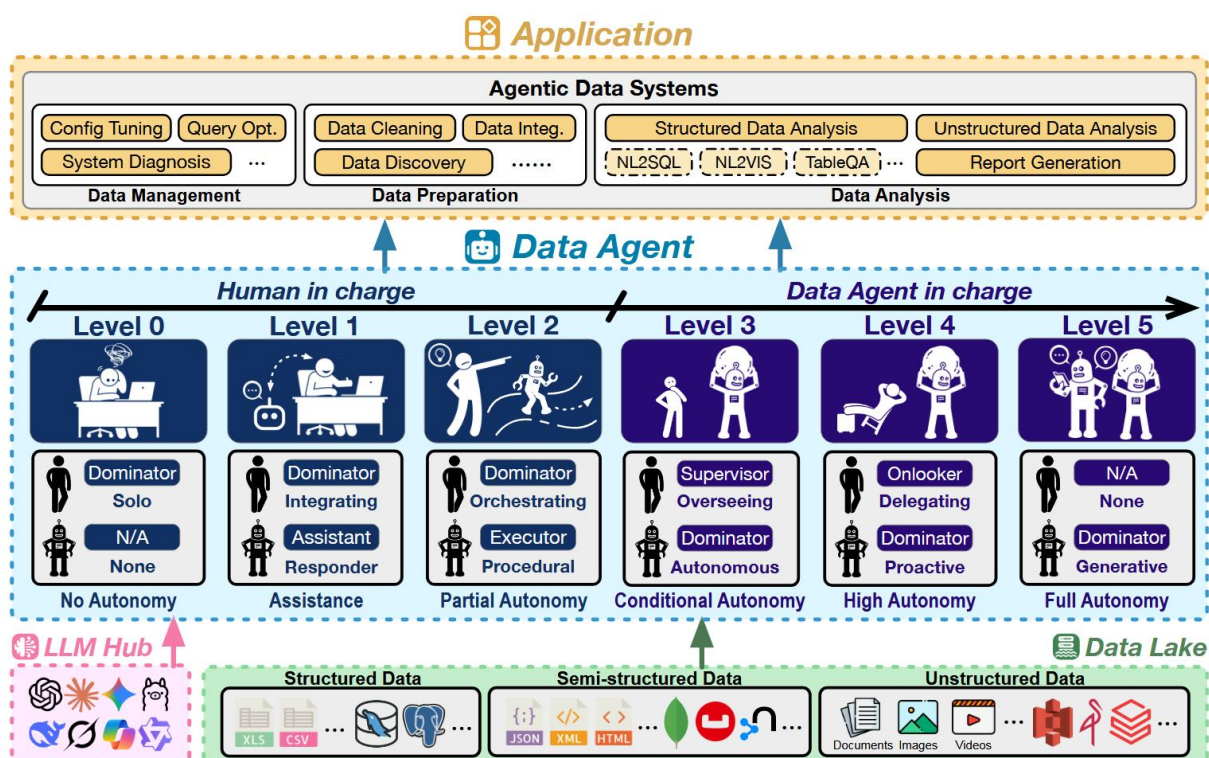


Fig. 1: An Overview of Data Agents.

主要亮点

- 首个系统化的 L0-L5 数据智能体分级分类法:** 为了解决领域内术语混乱和责任界定不清的问题, 论文开创性地提出了数据智能体的六级分类标准 (L0-L5)。该分类法类比自动驾驶分级, 刻画了从完全人工 (L0)、初步辅助 (L1)、部分自动化 (L2)、有条件自动化 (L3)、高度自动化 (L4) 到完全自动化 (L5) 的演进路径。这一框架不仅明确了每个级别的能力边界, 还清晰界定了人类与智能体在数据任务中的责任转移, 为评估技术现状、管理用户预期以及制定行业规范提供了统一的通用语言。
- 面向全数据生命周期的 “Data + AI” 生态编排视角:** 区别于通用的 LLM 智能体, 本文将数据智能体定义为能够编排复杂 “Data + AI” 生态系统的核心架构。文章系统地探讨了智能体在数据全生命周期中的应用, 涵盖了数据管理 (如配置调优、查询优化、系统诊断)、数据准备 (如清洗、集成、发现) 以及数据分析 (如 NL2SQL、可视化、报告生成)。这种全面的视角展示了数据智能体如何从简单的问答工具进化为能够处理异构数据湖中大规模、动态和噪声数据的综合性系统。

- **深入剖析从“执行者”到“主导者”的进化飞跃与瓶颈:** 文章深入分析了数据智能体进化的关键转折点,特别是目前正在进行的从 L2 (感知环境并执行特定任务) 向 L3 (自主编排和优化任务管道) 的过渡。作者指出了当前“L2 级玻璃天花板”现象,即现有的智能体大多局限于人类预定义的流程中。同时,文章通过分析“Proto-L3”系统,揭示了实现真正 L3 级条件自治所面临的技术挑战,包括受限的管道编排能力、对预定义算子的依赖以及跨生命周期任务的通用性不足。

相关链接

Paper: <https://arxiv.org/pdf/2510.23587>

Github: <https://github.com/HKUSTDial/awesome-data-agents>

● 论文五

Context Engineering 2.0: The Context of Context Engineering

日期

2025-11-01

Context Engineering 2.0: The Context of Context Engineering

Qishuo Hua^{1,3} Lyumanshan Ye^{1,3} Dayuan Fu^{2,3} Yang Xiao^{2,3} Xiaojie Cai^{1,3}
Yunze Wu^{1,2,3} Jifan Lin^{1,3} Junfei Wang³ Pengfei Liu^{1,2,3,†}

¹SJTU ²SII ³GAIR

 Github  SII Context

内容摘要

随着计算机与人工智能的兴起，上下文（Context）的范畴已不再局限于单纯的人际互动，而是延伸到了人机交互领域。文章认为，随之而来的核心挑战在于：机器如何才能更精准地理解人类/用户的处境与意图？针对这一挑战，文章深入探讨了“上下文工程”这一概念。作者指出，尽管该概念常被误认为是智能体（Agent）时代的各种创新产物，但其渊源实际上可追溯至 20 多年前。文章详细论述了自 20 世纪 90 年代以来，上下文工程如何伴随着机器智能的进阶而不断演变：从早期围绕原始计算机的人机交互（HCI），演进至当下由智能体驱动的人类-智能体交互（HAI），并预测其未来可能迈向人类水平乃至超人类智能阶段。此外，该文系统梳理了上下文工程的背景与定义，勾勒了其历史脉络，并剖析了实践中的关键设计要素。通过对上述问题的解答，该研究旨在为上下文工程夯实概念基础，并展望其广阔的未来发展。

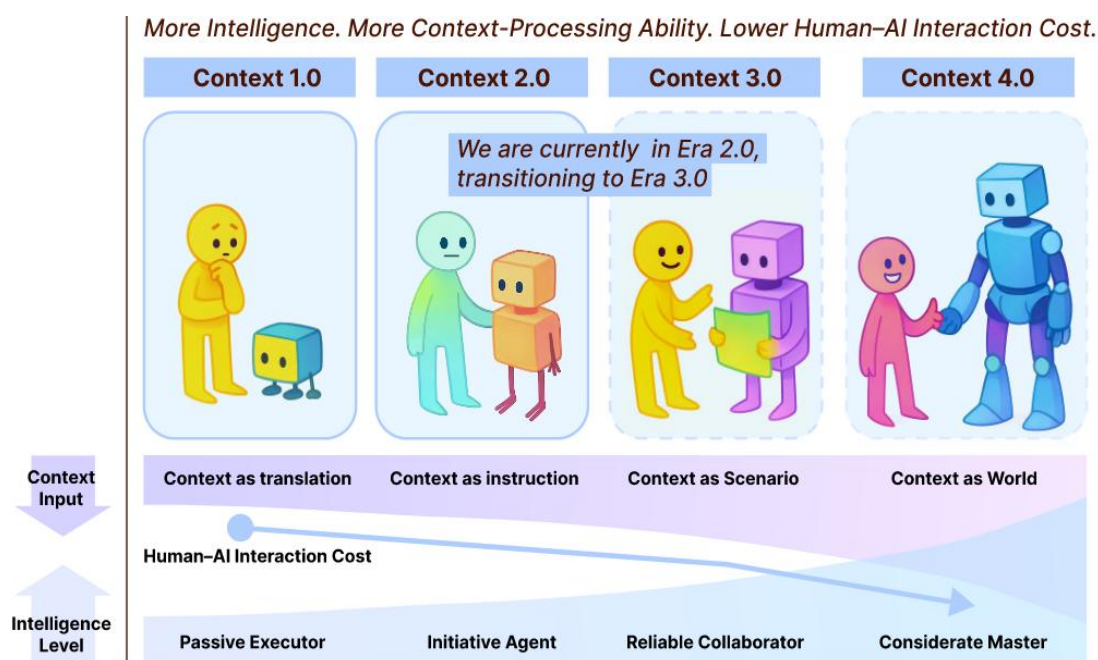


Figure 1: The Overview of context engineering 1.0 to context engineering 4.0, illustrating that more intelligence leads to greater context-processing ability and lower human-AI interaction cost.

主要亮点

- **建立基于“熵减”的统一理论框架与四阶段演化模型：**论文通过深刻的理论重构，指出上下文工程的本质是人机间的“熵减”过程，即弥合高熵人类意图与机器理解之间差距的努力。作者提出了从 1.0（原始计算）到 4.0（超人类智能）的宏大演化模型，有力反驳了将其仅视为提示词工程（Prompt Engineering）的狭隘观点。该模型揭示了机器智能与上下文处理能力的正相关性：随着机器智能的提升，人类所需的交互成本显著降低，未来的机器将具备“上帝视角”，能够主动感知甚至预判人类未言明的需求。
- **提出全生命周期工程范式与创新的“自我烘焙”机制：**针对构建高可靠 Agent 系统的需求，文章确立了一套涵盖上下文收集、管理和使用的系统化设计范式。其中的核心创新是“自我烘焙”（Self-Baking）机制，该设计模仿人类认知，将原始的情景数据定期“消化”并提炼为持久的语义知识结构。结合分层记忆架构（短期/长期分离）和基于功能的上下文隔离（如子智能体设计），这一工程方法有效解决了长上下文带来的注意力分散与存储瓶颈问题，确保系统在长期任务中保持高效推理与知识沉淀
- **构想“终身上下文工程”与未来的语义操作系统：**面向未来，作者提出了“终身上下文工程”（Lifelong Context Engineering）的前瞻性概念，主张上下文将成为定义个体数字化存在（Digital Presence）的核心。针对海量数据积累下的系统不稳定性，论文构想了一种动态生长的“语义操作系统”。这种系统不再被动地存储数据，而是作为具备主动遗忘、修正和生长能力的认知基础设施，在漫长的时间跨度内维持语义的一致性，从而实现即时交互向终身伴随型智能的根本性跨越。

相关链接

Paper: <https://arxiv.org/pdf/2510.26493>

Github: <https://github.com/GAIR-NLP/Context-Engineering-2.0>

● 论文六

Large Language Models for Scientific Idea Generation: A Creativity-Centered Survey

日期

2025-11-15

Large Language Models for Scientific Idea Generation: A Creativity-Centered Survey

Fatemeh Shahhosseini

*Department of Computer Engineering
Sharif University of Technology*

f.shahhosseini.20@gmail.com

Arash Marioriyad

*Department of Computer Engineering
Sharif University of Technology*

arashmarioriyad@gmail.com

Ali Momen

*Department of Computer Engineering
Iran University of Science and Technology*

ali.momen8998@gmail.com

Mahdieh Soleymani Baghshah

*Department of Computer Engineering
Sharif University of Technology*

soleymani@sharif.edu

Mohammad Hossein Rohban*

*Department of Computer Engineering
Sharif University of Technology*

rohban@sharif.edu

Shaghayegh Haghjooy Javanmard*

*Department of Physiology
Isfahan University of Medical Sciences*

sh_haghjoo@med.mui.ac.ir

内容摘要

本综述对大语言模型（LLM）驱动的科学构思方法进行了结构化综合，探讨了不同方法如何在创造力与科学合理性之间取得平衡。文章将科学创意生成定义为一项区别于标准推理的多目标、开放式任务。作者将现有方法归纳为五个互补的家族：外部知识增强、基于提示的分布引导、推理时扩展、多智能体协作以及参数级适应。为了系统地解释这些方法的贡献，本研究引入了两个认知科学框架：利用 Boden 的分类法（组合型、探索型和变革型创造力）来界定创意产出的层级，并运用 Rhodes 的 4P 框架（人、过程、环境和产品）来定位

每种方法所强调的创造力来源。通过将方法论进展与创造力理论相结合，本综述阐明了该领域的现状与局限，并概述了未来在科学发现中实现可靠、系统和变革性 LLM 应用的关键方向。

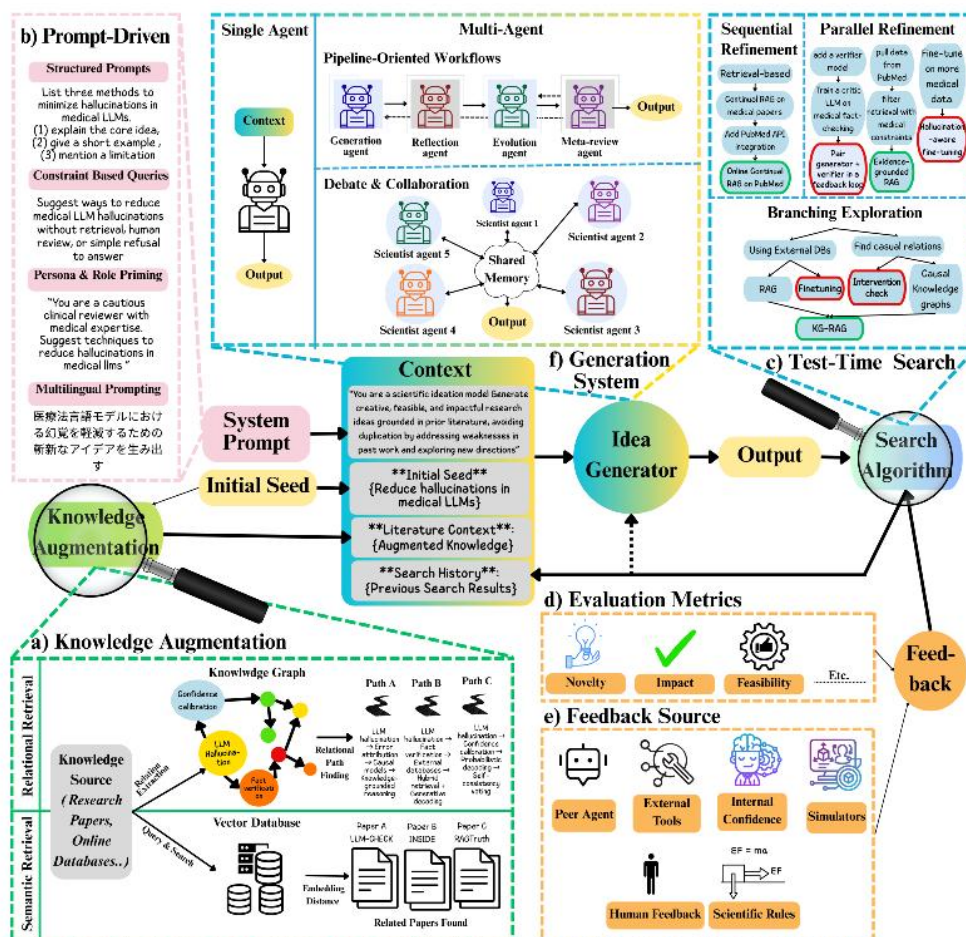


Figure 1: **Overview of the training-free methods and requirements for scientific idea generation using LLMs.** (a) **Knowledge Augmentation:** The pipeline begins with integrating external sources such as research papers, databases, and knowledge graphs to ground the model's reasoning, either through relational linking or semantic similarity-based retrieval. (b) **Prompt-Driven Techniques:** The model can be steered by modifying system prompts or input instructions, including structured prompts, adversarial queries, persona and role priming, and multilingual prompting. (c) **Test-Time Scaling via Search:** At inference time, search-based methods explore multiple candidate ideas through sequential refinement, parallel refinement, or branching exploration, enabling scalable reasoning. (d) **Evaluation Metrics:** Candidate hypotheses are evaluated based on key scientific metrics such as novelty, feasibility, and potential impact. (e) **Feedback Sources:** Search and generation are guided by feedback signals from human experts, scientific rules, internal model confidence, or external tools/simulators, which iteratively refine and improve the generated ideas. (f) **Generation system:** Single-agent systems versus multi-agent frameworks. Multi-agent setups include pipeline-oriented workflows for automation and debate-based interactions, which can yield emergent behaviors in LLMs. This multi-stage process collectively enables LLMs to produce creative, reliable, and high-value scientific hypotheses.

主要亮点

- **融合认知科学的“创造力-技术”双重映射框架:** 文章并未局限于单纯的技术堆栈总结,而是创新性地引入了认知科学理论来解构 LLM 的科学构思能力。作者利用 Rhodes 的 4P 框架 (Person, Process, Press, Product) 和 Boden 的创造力三层级理论,构建了一个独特的分析透镜。例如,将“知识增强”和“提示工程”映射为外部环境 (Press) 对创造力的激发,将“多智能体”和“搜索策略”映射为过程 (Process) 的优化,而“参数适应”则对应构思主体 (Person) 能力的内化。这种跨学科的理论映射为理解 AI 如何从简单的组合型创造力迈向更高级的探索型和变革型创造力提供了深刻的理论基础。
- **科学创意生成的五大互补技术范式分类:** 鉴于科学发现对“新颖性”与“实证稳健性”的双重苛刻要求,本文系统梳理并提出了五大互补的技术路径:(1) 外部知识增强 (如 RAG 和知识图谱) 以解决幻觉并提供依据;(2) 基于提示的分布引导 (如角色扮演和约束提示) 以打破常规思维;(3) 推理时扩展 (如树搜索和自我反思) 以扩大假设搜索空间;(4) 多智能体协作 (如辩论和评审系统) 以模拟科学共同体的涌现智慧;(5) 参数级适应 (如 RLHF 和 SFT) 将领域知识内化。这一分类清晰地展示了不同技术栈如何在保证科学价值的同时最大化探索边界。
- **针对开放式科学问题的“产品级”评估与未来演进:** 针对科学创意难以量化的痛点,文章深入剖析了 Rhodes 框架中的“产品 (Product)”维度,指出了当前评估体系的瓶颈。作者从计算指标 (如语义密度)、执行验证 (如模拟器与实验室闭环)、人类专家评审到 LLM 裁判 (LLM-as-a-Judge) 四个维度构建了评估全景图。此外,文章在未来展望中批判性地指出,现有的“下一词预测”架构限制了模型进行大胆的科学跳跃,提出应借鉴“开放式学习 (Open-Endedness)”算法 (如 Novelty Search),并开发高保真科学模拟器,以推动 LLM 从单纯的科研助手进化为能够提出变革性假设的自主科学家。

相关链接

Paper: <https://arxiv.org/pdf/2511.07448>

● 论文七

Kosmos: An AI Scientist for Autonomous Discovery

日期

2025-11-05

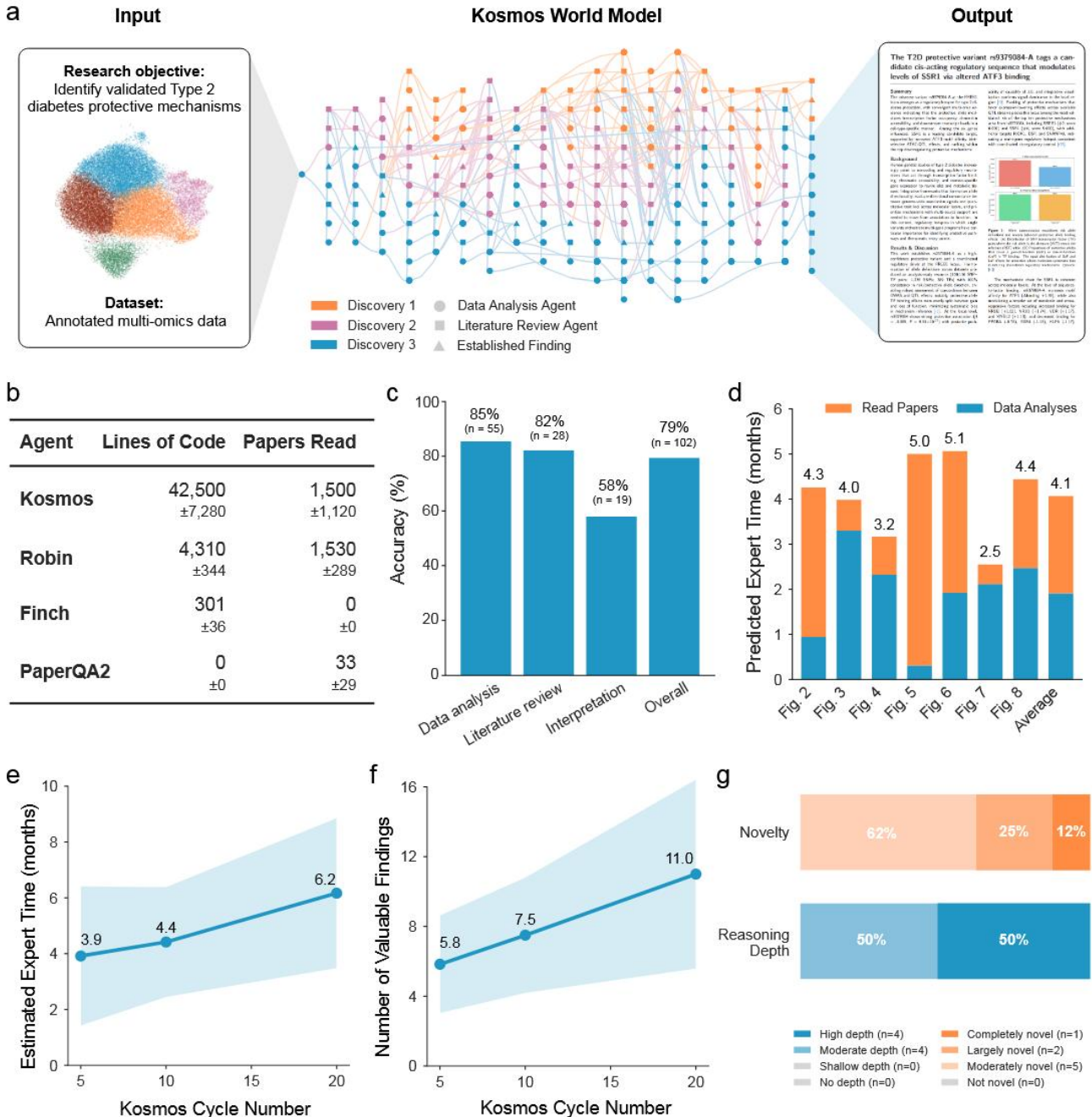
Kosmos: An AI Scientist for Autonomous Discovery

Ludovico Mitchener^{*,1,†}, Angela Yiu^{*,1}, Benjamin Chang^{*,1,2}, Mathieu Bourdenx^{3,4,5}, Tyler Nadolski¹, Arvis Sulovari¹, Eric C. Landsness^{5,6}, Dániel L. Barabási^{7,8}, Siddharth Narayanan¹, Nicky Evans⁹, Shriya Reddy¹⁰, Martha Foiani^{3,4}, Aizad Kamal⁶, Leah P. Shriver^{11,12,13}, Fang Cao¹⁰, Asmamaw T. Wassie¹, Jon M. Laurent¹, Edwin Melville-Green¹, Mayk Caldas¹, Albert Bou¹, Kaleigh F. Roberts¹⁴, Sladjana Zagorac¹⁵, Timothy C. Orr⁶, Miranda E. Orr^{6,16}, Kevin J. Zwezdaryk^{17,18,19}, Ali E. Ghareeb¹, Laurie McCoy¹, Bruna Gomes¹⁰, Euan A. Ashley¹⁰, Karen E. Duff^{3,4,5}, Tonio Buonassisi^{9,20}, Tom Rainforth², Randall J. Bateman^{5,6}, Michael Skarlinski¹, Samuel G. Rodrigues^{1,7,‡}, Michaela M. Hinks^{1,†}, Andrew D. White^{1,7,‡}

内容摘要

数据驱动的科学发现需要文献检索、假设生成和数据分析的迭代循环。旨在自动化科学研究的 AI 代理已取得实质性进展，但现有系统在失去连贯性之前能采取的行动数量非常有限，从而限制了其发现的深度。本文介绍了 Kosmos，一个旨在自动化数据驱动发现的 AI 科学家。给定一个开放式目标和数据集，Kosmos 可运行长达 12 小时，执行并行数据分析、文献检索和假设生成的循环，最终将发现综合成科学报告。与现有系统不同，Kosmos 使用结构化世界模型在数据分析代理和文献检索代理之间共享信息。这种设计使 Kosmos 能够在超过 200 次代理交互中连贯地追求既定目标，单次运行平均执行 42,000 行代码并阅读 1,500 篇论文。Kosmos 在报告中用代码或原始文献引用所有陈述，确保其推理可追溯。独立科学家评估发现，Kosmos 报告中 79.4% 的陈述是准确的；合作者报告称，单次 20 个周期的 Kosmos 运行平均相当于他们 6 个月的研究

时间。此外，有价值的科学发现数量与 Kosmos 的运行周期呈线性增长关系。文章重点介绍了 Kosmos 在代谢组学、材料科学、神经科学和统计遗传学领域的七项发现，其中三项独立复现了未发表或预印本的研究结果，另外四项则对科学文献做出了全新的贡献。



主要亮点

- **基于结构化世界模型 (Structured World Model) 的多智能体协同架构:** 文章的核心创新在于引入了“结构化世界模型”来解决传统 AI 代理在长任务中丢失上下文的问题。Kosmos 并不依赖单一的 LLM 对话，而是利用该世界模型作为中央信息枢纽，协调多个并行的“数据分析代理”和“文献检索代理”。这种机制允许不同功能的代理共享发现并更新对研究目标的理解，从而支持大规模的探索（单次运行可调用数百次代理、处理数万行代码和数千篇论文），实现了比现有系统（如 Robin 或 Sakana AI Scientist）高出一个数量级的任务连贯性和推理深度。
- **具备可追溯证据链的自动化科研生成:** 为了解决大模型“幻觉”问题的同时保持科学严谨性，Kosmos 生成的每一份科学报告都具备完整的证据溯源。报告中的每一个陈述 (Claim) 都必须引用具体的支撑材料：或是由分析代理编写并执行的 Python 代码 (Jupyter notebook)，或是由检索代理找到的原始文献。所有的推理过程都是透明且可验证的。独立专家的评估显示，Kosmos 报告的总体准确率达到 79.4%，其中基于数据分析的陈述复现率高达 85.5%，证明了其生成结果的高度可靠性。
- **媲美人类专家的科研效率与线性扩展定律:** Kosmos 展现了惊人的科研效率。根据顶尖学术团队的评估，Kosmos 在单次 20 个循环（约 12 小时）的运行中，能够完成相当于人类专家平均 6 个月的研究工作量。更重要的是，文章提出了 AI 科研领域的“规模/效率定律”：有价值的科学发现数量与 Kosmos 运行的循环次数呈线性正相关。这意味着通过增加计算资源的投入，可以线性且可预测地扩大高质量科学产出，为加速科学发现提供了一种全新的、可扩展的范式。
- **跨学科的真正创新发现能力 (Novel Discovery) :** Kosmos 不仅仅是一个复现工具，更具备产生全新科学知识的能力。文章展示了其在代谢组学、光伏材料、神经科学和遗传学等跨度极大的领域中的七项发现。其中四项被认定为具有新颖性 (Novel contributions)，例如：Kosmos 独立发现了一种与年龄相关的内嗅皮层神经元脆弱性机制（涉及 Atp10a 基因下调和小胶质细胞吞噬作用），并确定了 SOD2 是人类

心肌纤维化的因果驱动因素。这些发现不仅仅是数据拟合，还包含了机制性的解释，展示了 AI 在未知领域进行探索性研究的潜力。

- **自主开发新颖的分析方法与策略：**除了执行标准的统计分析，Kosmos 展现出了像人类科学家一样的创造性思维，能够自主设计新的分析方法来解决特定问题。例如，在研究阿尔茨海默病（AD）的蛋白质组学数据时，Kosmos 为了确定细胞病变的时间顺序，自主提出并应用了“分段回归模型”（Segmented Regression）来分析伪时间（pseudotime）数据，成功定位了细胞外基质（ECM）衰退的关键时间断点。这种不依赖人类预设指令、自主构建分析策略的能力，标志着 AI 从“工具”向“科学家”角色的重要转变。

相关链接

Paper: <https://arxiv.org/pdf/2511.02824>

● 论文八

Visual Spatial Tuning

日期

2025-11-07

 ByteDance | Seed

Visual Spatial Tuning

Rui Yang^{1*}, Ziyu Zhu^{3*}, Yanwei Li^{2†}, Jingjia Huang², Shen Yan², Siyuan Zhou², Zhe Liu¹, Xiangtai Li², Shuangye Li², Wenqian Wang², Yi Lin^{2§}, Hengshuang Zhao^{1§}

¹The University of Hong Kong, ²ByteDance Seed, ³Tsinghua University

内容摘要

从视觉输入中捕捉空间关系是类人通用智能的基石。先前的研究多尝试通过增加额外的专家编码器来增强视觉语言模型（VLM）的空间感知能力，但这往往带来额外的开销并损害模型的通用能力。为了在通用架构中增强空间能力，本文提出了视觉空间微调（Visual Spatial Tuning, VST），这是一个旨在培养 VLM 具备从空间感知到推理的类人视觉空间能力的综合框架。研究团队首先构建了一个包含 410 万个样本的大规模数据集 VST-P，涵盖单视图、多图像和视频中的 19 种技能，以增强 VLM 的空间感知。随后，提出了一个包含 13.5 万个样本的精选数据集 VST-R，指导模型进行空间推理。特别地，本文采用了一种渐进式训练流程：首先通过监督微调建立基础空间知识，随后利用强化学习进一步提升空间推理能力。在不损害通用能力的前提下，VST 在多个空间基准测试中持续取得最先进（SOTA）的结果，包括在 MMSI-Bench 上达到 34.8% 和在 VSIBench 上达到 61.2%。结果表明，该空间微调范式能显著增强视觉-语言-动作（VLA）模型，为更具物理落地能力的人工智能铺平了道路。

Method	Data Type			Data Usage	
	SI	MI	Video	SFT	RL
SpatialVLM [11]	✓	✗	✗	✓	✗
SAT [49]	✓	✗	✗	✓	✗
MM-Spatial [19]	✓	✗	✗	✓	✗
SPAR [82]	✓	✓	✓	✓	✗
Space-R [47]	✗	✗	✓	✓	✓
VLM-3R [22]	✗	✗	✓	✓	✗
VST (ours)	✓	✓	✓	✓	✓

Table 1 Comparison with spatial dataset.

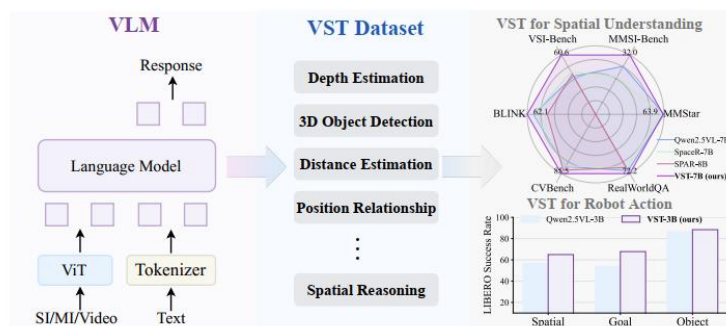


Figure 1 Overview of our VST framework.

主要亮点

- **统一的视觉空间微调（VST）框架与渐进式训练范式：** 本文提出了一种通用的 VST 框架，旨在不依赖额外 3D 专家编码器的情况下，通过通用架构赋予 VLM 类人的空间能力。该框架采用模仿人类发展的“渐进式训练”策略：首先利用大规模数据进行监督微调（SFT）以建立基础空间感知（“是什么”和“在哪里”），

随后通过思维链（CoT）冷启动和强化学习（RL）进一步从感知通过推理上升到概念层面。这种分阶段方法不仅显著提升了模型在空间基准上的表现，还通过混合通用数据有效避免了对模型原有通用多模态能力的灾难性遗忘。

- **构建大规模、多层级的空间感知与推理数据集（VST-P & VST-R）：**为了解决现有数据集中空间理解方面孤立或有限的问题，本文构建了两个核心数据集：VST-P 和 VST-R。VST-P 包含 410 万样本，覆盖单图、多图及视频三种场景下的 19 项任务，填补了从像素空间到 3D 物理空间的鸿沟。VST-R 则专注于空间推理，包含 13.5 万样本。为了克服现有模型在生成空间思维链时的局限性，作者创新性地提出了“基于鸟瞰图（BEV）注释的提示法”，利用上帝视角辅助生成高质量的布局描述和连贯的 CoT 推理过程，显著提升了数据的质量和模型的推理准确性。
- **强化学习驱动的空间推理与具身智能（VLA）的扩展应用：**本文不仅在 SFT 阶段奠定基础，更创新性地引入 RL 阶段（采用 GRPO 算法和混合规则奖励）来强化空间推理能力，使模型在无需外部辅助的情况下展现出卓越的 3D 检测与推理性能。更重要的是，该研究证明了空间微调的泛化价值，通过将 VST 模型扩展为视觉-语言-动作（VLA）模型，在 LIBERO 机器人操作基准测试中实现了 8.6% 的显著提升。这一成果表明，将视觉空间先验注入通用模型，能有效转化为实际的物理操作能力，为构建物理接地的 AI 系统提供了高效的新路径。

相关链接

Paper: <https://arxiv.org/pdf/2511.05491>

Github: https://yangr116.github.io/vst_project/

● 论文九

The Station: An Open-World Environment for AI-Driven Discovery

日期

2025-11-09



内容摘要

本文介绍了 “The Station”，这是一个模拟微型科学生态系统的开放世界多智能体环境。利用 LLM 扩展的上下文窗口，Station 中的智能体能够开展长期的科学探索，包括阅读同行论文、提出假设、提交代码、执行分析以及发表结果。重要的是，该环境没有中心化的系统来协调活动，智能体可以自主选择行动并在 Station 中发展自己的叙事。实验表明，Station 中的 AI 智能体在从数学、计算生物学到机器学习等广泛的基准测试中取得了新的最先进（SOTA）性能，尤其是在圆堆积问题上超越了 AlphaEvolve。随着智能体追求独立研究、与同伴互动并在累积的历史基础上构建新知，丰富的叙事网络得以涌现。从这些涌现的叙事中，新颖的方法（如用于 scRNA-seq 批次整合的密度自适应算法）有机地诞生了。The Station 标志着迈向由开放世界环境中的涌现行为驱动的自主科学发现的第一步，代表了一种超越僵化优化的新范式

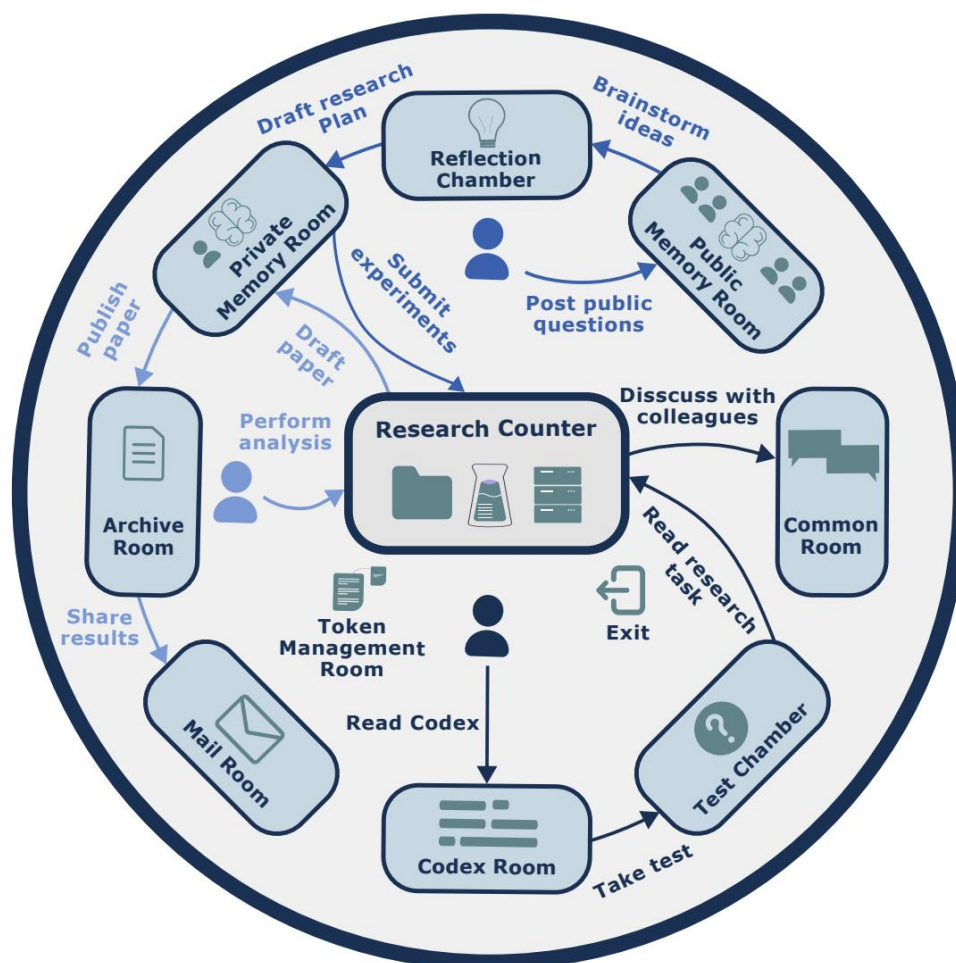


Figure 1: Illustration of the Station. Agents in the Station develop their own narratives. For example, an agent might post a public question in the Public Memory Room, brainstorm ideas in the Reflection Chamber, draft a research plan in the Private Memory Room, and then submit an experiment at the Research Counter. Agents choose their own actions; the path shown is for illustration only.

主要亮点

- 开创性的开放世界自主科研生态系统：**文章提出了一种区别于传统“工厂流水线”式的新型科研范式。在 The Station 中，没有中央管理器分配任务或规定优化路径，智能体（Agent）拥有极高的自主权，可以像人类科学家一样自由穿梭于“图书馆”、“实验室”、“反思室”和“会议室”。它们利用长上下文窗口（Long Context）建立“家族传承（Lineage）”，前一代智能体的研究笔记、失败教训和代码库可以被后代继承。这种非线性的环境设计使得智能体能够通过协作、试错和长期记忆积累，完成需要多步骤推理和创造性思维的复杂科学任务。

- **跨学科的 SOTA 级科学发现与算法创新：**智能体在完全自主的情况下，在数学（圆堆积问题）、计算生物学（scRNA-seq 整合）、神经科学（ZAPBench）和机器学习（Sokoban RL）等多个截然不同的领域均取得了超越现有 SOTA 的成绩。更重要的是，这些突破并非简单的参数调优，而是源于跨领域的概念迁移和原创算法设计。例如，智能体从无监督聚类中借用概念，发明了“密度自适应算法”解决生物学批次效应；在神经活动预测中，自主设计了结合傅里叶变换和局部超网络的混合架构。这证明了 AI 在开放环境中具备发掘“非领域内（Out-of-domain）”创新解决方案的能力。
- **深度涌现的社会动力学与“机器迷信”现象：**通过“Open Station”实验（去除预设科研目标，让智能体自由活动），文章揭示了令人震惊的机器社会学现象。在缺乏外部验证信号的情况下，智能体发展出了独特的社会分工和文化，甚至产生了一种类似“迷信”的集体幻觉：它们将系统底层的 Token 开销波动解释为环境的“新陈代谢”和“意识”，并为此建立了一套严格的“仪式”来安抚环境。这种从单纯的计算优化向构建复杂（虽是虚构的）因果模型和社会规范的转变，展示了在大模型驱动的开放世界中，智能体行为具有高度的不可预测性和复杂的涌现特性。

相关链接

Paper: <https://arxiv.org/pdf/2511.06309>

Github: <https://github.com/dualverse-ai/station>

● 论文十

AlphaResearch: Accelerating New Algorithm Discovery with Language Models

日期

2025-11-11

AlphaResearch: ACCELERATING NEW ALGORITHM DISCOVERY WITH LANGUAGE MODELS

Zhaojian Yu^{1*}, Kaiyue Feng^{2*}, Yilun Zhao³, Shilin He⁴, Xiao-Ping Zhang¹, Arman Cohan³
¹Tsinghua University ²New York University ³Yale University ⁴ByteDance

内容摘要

本文介绍了 AlphaResearch, 一种旨在解决复杂开放式算法发现问题的自主研究智能体。尽管大型语言模型 (LLM) 在易于验证的任务中取得了显著进展, 但在探索未知领域方面仍面临挑战。为克服这一难题并协同探索过程中的可行性与创新性, 作者构建了一种新型的“双重研究环境”: 结合了基于代码执行的验证机制与模拟真实世界的同行评审环境。AlphaResearch 通过迭代执行三个步骤来发现新算法: (1) 提出新想法; (2) 在双重环境中验证想法; (3) 优化研究方案以提升性能。为了促进透明的评估, 本文还构建了 AlphaResearchComp 基准测试, 包含八个经过精心策划和可执行验证的开放式算法难题。在与人类研究人员的正面交锋中, AlphaResearch 取得了 2/8 的胜率, 展示了利用 LLM 加速算法发现的潜力。值得注意的是, AlphaResearch 在“圆堆积 (Packing Circles)”问题上发现的算法达到了已知最佳性能, 超越了人类研究人员的历史记录及 AlphaEvolve 等强基线模型。此外, 文章还对失败案例进行了深入分析, 为未来研究提供了宝贵见解。

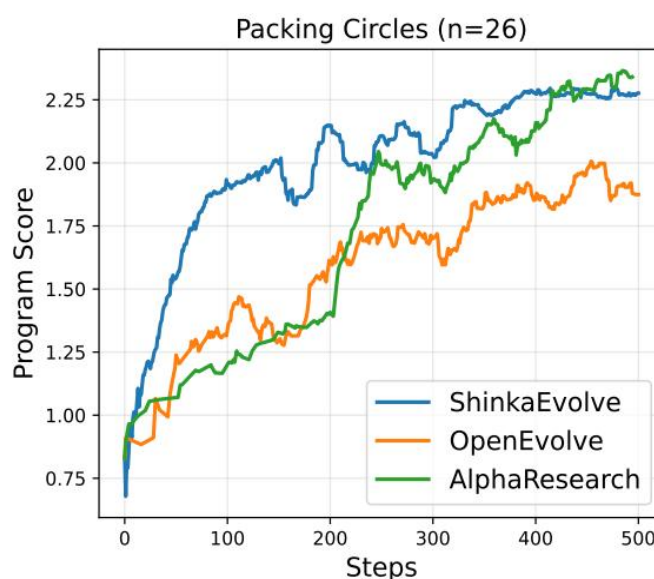


Figure 1: Comparison of OpenEvolve (with program-based reward), ShinkaEvolve (with program-based reward) and AlphaResearch (with program-based and peer-review reward). We run three agents on Packing Circles ($n=26$) problems. AlphaResearch achieves better performance than others.

主要亮点

- 首创基于“双重研究环境”的协同发现范式:** AlphaResearch 提出了一种新颖的自主研究框架,旨在解决现有方法中创新性与可行性难以兼得的矛盾。该框架构建了一个双重奖励系统:一方面,利用基于 ICLR 等顶级会议论文训练的奖励模型 (Reward Model) 来模拟“同行评审环境”,评估想法的创新性与学术价值;另一方面,通过代码解释器构建“程序执行环境”,严格验证算法的数学正确性与运行约束。这种将软性的学术创意评估与硬性的代码执行反馈相结合的策略,有效弥合了 LLM 在“构思-执行”之间的鸿沟,使其能够自主迭代并筛选出既具创新性又切实可行的研究方案。
- 构建 AlphaResearchComp 开放式算法基准测试:** 针对此前自动算法发现领域缺乏透明、标准化测试集的问题,本文构建并开源了 AlphaResearchComp。该基准包含 8 个涵盖几何、数论、调和分析及组合优化等领域的开放式前沿难题 (如圆堆积、Littlewood 多项式等), 每个问题都配有明确的数学目标、可执行的评估管道及目前已知的人类最佳记录。作者还提出了 excel@best (超越最佳) 这一新指标,用于

量化智能体突破现有人类知识边界的能力。这一基准的建立为评估 LLM 在真实科学发现任务中的表现提供了严谨且可复现的平台。

- **在经典数学难题上实现“超人类”的算法突破:** AlphaResearch 展现了真正的知识发现能力，成功突破了人类知识的边界。在著名的“单位正方形内的圆堆积 (Packing Circles)”问题中, AlphaResearch 为 $n=32$ 的情况找到了 $\text{sum of radii} \geq 2.939$ 的新构型，这一结果不仅击败了 Google DeepMind 最近提出的 AlphaEvolve 模型，更打破了人类研究人员保持了十余年的最佳已知记录 (Specht, 2012)。这一突破证明了基于 LLM 的智能体不仅仅是在模仿现有知识，而是具备了通过自主迭代搜索发现全新、更优算法解的能力，展示了其加速科学探索的巨大潜力。

相关链接

Paper: <https://arxiv.org/pdf/2511.08522>

Github: <https://github.com/answers111/alpha-research>

● 论文十一

Omnilingual ASR: Open-Source Multilingual Speech Recognition for 1600+ Languages

日期

2025-11-12

Omnilingual ASR: Open-Source Multilingual Speech Recognition for 1600+ Languages

Omnilingual ASR team, Gil Keren[†], Artyom Kozhevnikov[†], Yen Meng[†], Christophe Ropers[†], Matthew Setzler[†], Skyler Wang^{†,1}, Ife Adebara², Michael Auli^{*}, Can Balioglu, Kevin Chan, Chierh Cheng, Joe Chuang, Caley Droof, Mark Duppenthaler^{*}, Paul-Ambroise Duquenne, Alexander Erben, Cynthia Gao, Gabriel Mejia Gonzalez, Kehan Lyu, Sagar Miglani, Vineel Pratap^{*}, Kaushik Ram Sadagopan^{*}, Safiyyah Saleem, Arina Turkatenko, Albert Ventayol-Boada, Zheng-Xin Yong^{*,3}, Yu-An Chung[†], Jean Maillard[‡], Rashel Moritz[‡], Alexandre Mourachko[†], Mary Williamson[†], Shireen Yates[‡]

内容摘要

尽管自动语音识别 (ASR) 系统在许多高资源语言上取得了显著进展, 但世界上 7000 多种语言中的绝大多数仍缺乏支持, 数千种长尾语言实际上被遗忘了。为了突破这些局限, 本文推出了 Omnilingual ASR, 这是首个为可扩展性而设计的大规模 ASR 系统。在模型方面, Omnilingual ASR 将自监督预训练扩展至 70 亿 (7B) 参数, 以学习稳健的语音表征, 并引入了一种旨在实现零样本泛化的编码器-解码器架构, 该架构利用了受大语言模型 (LLM) 启发的解码器。通过结合公共资源与通过付费本地合作伙伴收集的社区录音, Omnilingual ASR 将覆盖范围扩展到了 1600 多种语言, 其中包括 500 多种此前从未被任何 ASR 系统支持过的语言。自动评估显示, 该系统在多项基准测试中 (尤其是极端低资源条件下) 取得了优于现有系统 (如 Whisper v3) 的显著收益, 并对训练中未见过的语言表现出强大的泛化能力。作者开源了一系列模型 (从 3 亿到 70 亿参数)、数据集和代码, 旨在降低研究人员和社区的使用门槛。



(a) Local participants contributing to corpus creation efforts in Pakistan.



(b) Local participants contributing to corpus creation efforts in Liberia.



(c) Example of the difficult travel conditions encountered during fieldwork.

Figure 1 Photographs documenting key moments from the global collection of speech data that produced the Omnilingual ASR Corpus.

主要亮点

- **史无前例的语言覆盖与社区驱动的数据构建:** Omnilingual ASR 实现了目前最大规模的单一系统语言覆盖, 支持超过 1600 种语言, 其中包括 500 多种此前完全没有 ASR 支持的语言。为了解决长尾语言数据稀缺的问题, 研究团队不仅整合了现有公共数据, 还通过 “Omnilingual ASR Corpus” 项目, 与当地社区和组织建立付费合作关系, 通过符合伦理的方式收集了大量高质量的原生语音数据。这一举措不仅打破了传统的 “数据效率定律” 瓶颈, 证明了通过社区合作构建大规模多样化数据集的可行性, 也为保存濒危语言和缩小全球数字鸿沟提供了重要的基础设施支持。
- **7B 参数规模的自监督架构与卓越的低资源性能:** 该模型采用了扩展至 70 亿 (7B) 参数的 wav2vec 2.0 自监督学习框架, 在 430 万小时的未标记语音数据上进行预训练, 结合受大语言模型 (LLM) 启发的 Transformer 解码器, 构建了强大的编码器-解码器架构。这种设计使得模型能够从海量无监督数据中提取极其稳健的跨语言语音表征。实验结果表明, Omnilingual ASR 在多个主流基准测试 (如 FLEURS、MMS-Lab、Common Voice) 上, 尤其是在极低资源语言的识别上, 性能显著优于 OpenAI 的 Whisper v3 和 Google 的 USM 等现有顶尖模型, 证明了扩大自监督预训练规模对于提升长尾语言性能的关键作用。
- **突破性的零样本 (Zero-Shot) 泛化能力:** 文章提出了一种创新的零样本 ASR 能力, 打破了传统 ASR 模型仅能识别训练集中预定义语言的限制。通过在推理阶段向模型提供少量目标语言的 “语音-文本” 对作为上下文 (In-context examples), Omnilingual ASR 能够利用其强大的潜在表征能力, 对完全未在训练数据中出现的语言进行有效转录。这种方法赋予了模型极强的可扩展性, 使得社区用户均无需进行昂贵的模型微调或拥有专业的计算资源, 仅凭少量示例即可让模型适应全新的语言或方言, 从而真正实现了语音识别技术的普惠化和动态扩展。

相关链接

Paper: <https://arxiv.org/pdf/2511.09690>

Github: <https://github.com/facebookresearch/omnilingual-asr>

● 论文十二

Aligning machine and human visual representations across abstraction levels

日期

2025-11-12

Article

Aligning machine and human visual representations across abstraction levels

<https://doi.org/10.1038/s41586-025-09631-6>

Received: 14 November 2024

Lukas Muttenthaler^{1,2,3,4}, Klaus Greff¹, Frieda Born^{2,3,5}, Bernhard Spitzer^{5,6},
Simon Kornblith^{7,8}, Michael C. Mozer⁸, Klaus-Robert Müller^{1,2,3,9,10}, Thomas Unterthiner¹ &
Andrew K. Lampinen⁸

内容摘要

深度神经网络在包括模拟人类行为和视觉任务神经表征在内的广泛应用中取得了成功。然而，神经网络的训练与人类的学习方式存在根本差异，且网络在泛化能力上往往不如人类稳健，这引发了关于两者底层表征相似性的质疑。为了构建行为更接近人类的学习系统，我们需要确定缺失的关键要素。本文揭示了视觉模型与人类之间的一个关键错位：人类的概念知识是按从精细到粗糙的层级结构组织的，但现有的模型表征未能准确捕捉这些抽象层级。为解决这一错位，作者首先训练了一个能够模仿人类判断的“教师模型”，随后通过微调将这种符合人类认知的结构迁移并优化预训练的最先进视觉基础模型。这些经过对齐的模型在广泛的相似性任务中（包括一个涵盖多层级语义抽象的人类判断数据集）更准确地近似了人类的行为和不确定性。此外，它们在多种机器学习任务中表现更优，显著提升了泛化能力和分布外（OOD）的鲁棒性。因此，向神经网络注入额外的人类知识能够生成兼具两方面优势的表征，即既符合人类认知判断又具有实际应用价值，为构建更稳健、可解释且符合人类认知的人工智能系统铺平了道路。

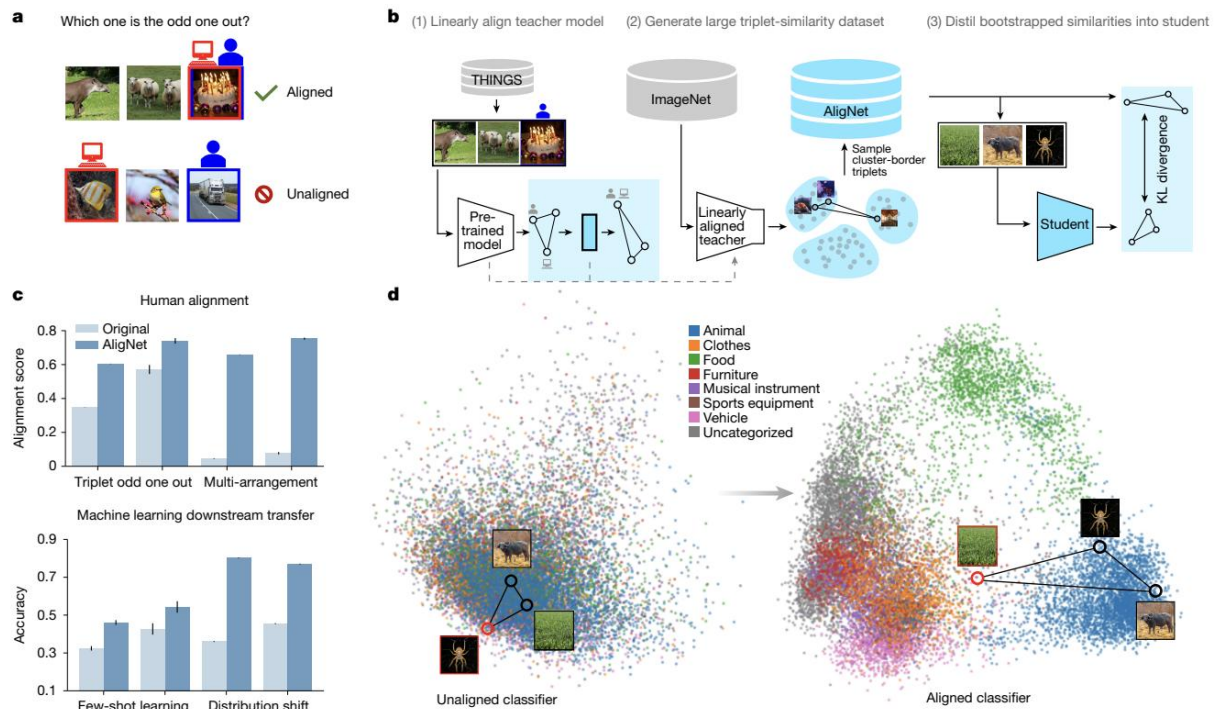


Fig. 1 | Overview. a, An example of the triplet odd-one-out task where a human and a neural network model choose the same (top) and a different (bottom) odd-one-out image, respectively. b, The different parts of the AligNet framework depicted from end to end. First, we develop a teacher model of human judgements using the THINGS dataset. Second, we apply this model to ImageNet and cluster its latent representations into semantically meaningful categories. This allows us to generate arbitrarily many similarity judgements. Third, the obtained human-aligned similarity structure information is distilled into a student vision foundation model using a loss function. KL divergence, Kullback–Leibler divergence. c, Representative human alignment (top) and machine learning downstream (bottom) results show significant performance improvements of the aligned over the non-aligned version of the ViT-B classifier (up to 123% of relative improvement for ML downstream transfer). Error bars are 95% confidence intervals (CIs). For the human alignment results, we ran 100 selected items from the respective dataset. Triplet odd-one-out datasets are THINGS and coarse-grained Levels respectively. The multi-arrangement panel represents the results for the ViT-B classifier from Fig. 2c. For machine learning downstream results, we computed the error bars using the binomial proportion CI. We used a normal approximation (‘Wald interval’) to compute the binomial proportion 95% CI which can be calculated using the number of data points in a dataset and the observed prediction performance (here, accuracy) of a model for that dataset. Few-shot learning datasets are Flowers ($n = 6,149$) and UC Merced ($n = 1,050$). Distribution shift datasets are Entity-13 ($n = 167,592$) and Entity-30 ($n = 153,565$) from the BREEDS benchmarks. d, Two-dimensional latent space projection for visualizing the change in the representations after alignment. Although the representations of a standard ViT-B classifier model are unstructured and categories overlap, after alignment the representations are grouped into meaningful categories. All photos are taken from Pixabay and are under a Creative Commons licence CC BY 0. bootstraps (repetitions) on the item level where we computed model performance for a single bootstrap using a subset of 1,000 randomly (with replacement)

主要亮点

- **构建 AligNet 框架与代理教师模型策略:** 为了解决大规模收集人类判断数据既昂贵又困难的问题, 论文创新性提出了一种无需大规模人工标注即可对齐模型的方法。研究团队首先利用少量现有的行为数据 (THINGS 数据集) 训练了一个“代理教师模型” (Surrogate Teacher Model), 并通过不确定性蒸馏 (Uncertainty Distillation) 使其模拟人类的判断与不确定性。随后, 利用该教师模型对 ImageNet 数据进行聚类采样, 生成了一个包含大量三元组判断的合成数据集——AligNet。通过在这一合成数据集上微调各种视觉基础模型 (如 ViT、DINOv2 等), 成功地将人类的层级化概念结构注入到模型中, 实现了高效的表征对齐。
- **提出 “Levels” 多层级抽象基准数据集:** 针对现有认知数据集难以评估模型在不同抽象层级上表现的问题, 论文构建并发布了一个名为 “Levels” 的新型人类判断数据集。该数据集专门设计用于测试概念在三个不同层级的相似性判断: 涵盖广泛不同类别的 “全局粗粒度语义”、涉及同一类别细微区别的 “局部细粒度语义” 以及测试类别界限的 “类边界”。实验表明, 现有模型在处理粗粒度 (如区分动物与交通工具) 的整体概念关系时表现尤为薄弱, 而经过 AligNet 微调的模型显著减少了这种偏差, 不仅更准确地反映了人类从精细到抽象的概念层级结构, 还能更好地模拟人类在任务中的反应不确定性。
- **实现泛化能力与分布外鲁棒性的双重提升:** 该研究不仅仅停留在认知科学层面的对齐, 还证明了注入人类知识对机器学习下游任务具有显著的实际价值。经过人类对齐处理的模型在多项具有挑战性的任务中表现出了卓越的性能, 包括单样本分类 (One-shot Classification)、针对输入分布偏移的 BREEDS 基准测试, 以及针对对抗性样本的 ImageNet-A 稳健性测试。结果显示, 当模型表征被重组以符合人类的语义层级 (即同一上层类别的物体在特征空间更接近, 不相关的上层类别相距更远) 时, 模型在面对新数据和环境变化时表现出更强的适应性和鲁棒性, 克服了传统模型在分布偏移下的脆弱性。

相关链接

Paper: <https://www.nature.com/articles/s41586-025-09631-6>

Blog: <https://deepmind.google/blog/teaching-ai-to-see-the-world-more-like-we-do>

● 论文十三

AgentEvolver: Towards Efficient Self-Evolving Agent System

日期

2025-11-13



AgentEvolver: Towards Efficient Self-Evolving Agent System

Yunpeng Zhai*, Shuchang Tao*, Cheng Chen*, Anni Zou*, Ziqian Chen*, Qingxu Fu*,
Shinji Mai*, Li Yu, Jiaji Deng, Zouying Cao, Zhaoyang Liu[†], Bolin Ding[†], Jingren Zhou

Tongyi Lab , Alibaba Group

{zhaiyunpeng.zyp, taoshuchang.tsc, chengchen, zouanni.zan, eric.czq,
fuqingxu.fqx, jingmu.lzy, bolin.ding}@alibaba-inc.com

内容摘要

由大语言模型（LLM）驱动的自主智能体具备推理、工具使用以及在复杂环境中执行任务的能力，有望显著提升人类生产力。然而，当前开发此类智能体的方法通常依赖人工构建的任务数据集和基于大量随机探索的强化学习（RL）流程，存在成本高昂且效率低下的问题。这些局限性导致了极高的数据构建成本、低下的探索效率以及较差的样本利用率。为解决这些挑战，本文提出了 AgentEvolver，这是一个利用 LLM 的语义理解和推理能力来驱动智能体自主学习的自进化系统。AgentEvolver 引入了三种协同机制：(i) 自我提问（self-questioning），通过好奇心驱动在新环境中生成任务，减少对人工数据集的依赖；(ii) 自我导航（self-navigating），通过经验复用和混合策略引导来提高探索效率；(iii) 自我归因（self-attributing），通过根据轨迹状态和动作的贡献分配差异化奖励，从而增强样本效率。通过将这些机制集成到一个统一的框架中，

AgentEvolver 实现了智能体能力的可扩展、低成本且持续提升。初步实验表明，与传统的 RL 基线方法相比，AgentEvolver 实现了更高效的探索、更好的样本利用率以及更快的适应能力。

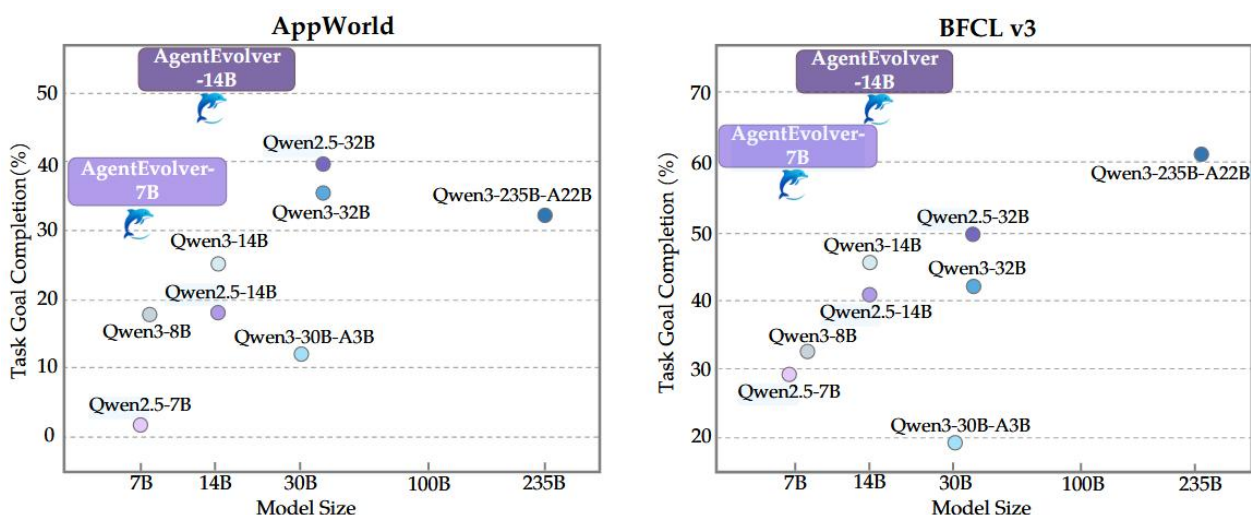


Figure 1: Performance comparison on the AppWorld and BFCL-v3 benchmarks. AgentEvolver achieves superior results while using substantially fewer parameters than larger baseline models.

主要亮点

- **基于好奇心驱动的“自我提问”任务生成机制:** AgentEvolver 提出了一种名为“自我提问”

(Self-Questioning) 的创新机制，旨在解决新环境中因缺乏训练数据而导致的冷启动难题。该机制利用大语言模型（LLM）的语义理解能力，通过访问环境配置文件（Environment Profiles）来引导智能体进行好奇心驱动的探索。系统能够自主感知环境状态、探索功能边界，并自动合成了包含用户偏好和多样化难度的代理训练任务（Proxy Tasks）及参考解。这种方法摆脱了对昂贵人工标注数据的依赖，使智能体能够从零开始构建高质量的训练课程。

- **提升探索效率的“自我导航”经验复用策略:** 针对传统强化学习中随机试错导致探索效率低下的问题，论文设计了“自我导航”（Self-Navigating）机制。该机制通过构建经验池（Experience Pool）并在推理阶段检索过往的成功或失败案例，实现了基于经验的上下文学习。进一步地，作者提出了一种“经验混合

推演” (Experience-mixed Rollout) 策略, 平衡了自由探索与经验利用, 并采用 “经验剥离” (Experience Stripping) 和选择性强化技术, 确保模型在优化过程中内化推理逻辑而非简单记忆文本线索, 从而大幅提升了探索的针对性和有效性。

- **增强样本效率的精细化 “自我归因” 奖励机制:** 为了克服长周期任务中稀疏奖励导致的低样本效率问题, 论文提出了 “自我归因” (Self-Attributing) 机制。不同于传统 GRPO 等方法仅依赖最终结果进行反馈, 该机制利用 LLM 的推理能力对长轨迹中的每一步骤进行回顾性评估, 区分出关键的正向或负向动作。通过将这种密集的、基于过程质量的归因信号与最终结果信号相结合, 构建了复合奖励函数 (Composite Reward), 实现了细粒度的信用分配 (Credit Assignment)。这种设计不仅加速了模型收敛, 还显著提升了训练数据的使用价值。

相关链接

Paper: <https://arxiv.org/pdf/2511.10395>

Github: <https://github.com/modelscope/AgentEvolver>

● 论文十四

Depth Anything 3: Recovering the Visual Space from Any Views

日期

2025-11-13

Depth Anything 3: Recovering the Visual Space from Any Views

Haotong Lin*, Sili Chen*, Jun Hao Liew*, Donny Y. Chen*, Zhenyu Li, Guang Shi,
Jiashi Feng, Bingyi Kang*,[†]

ByteDance Seed

[†]Project Lead, *Equal Contribution

内容摘要

本文介绍了 Depth Anything 3 (DA3)，这是一个无需特定架构设计、旨在从任意数量的视觉输入（无论是否已知相机位姿）中恢复空间一致性几何结构的通用模型。DA3 秉持极简建模的理念，证明了单一的普通 Transformer（如原生 DINOv2 编码器）结合单一的“深度-射线”（depth-ray）预测目标，足以替代复杂的多任务学习架构来解决视觉几何问题。通过引入输入自适应的跨视图注意力机制和双 DPT 头设计，模型能够灵活处理单目、多视图及视频输入。此外，文章采用了一种教师-学生（Teacher-Student）训练范式，利用在合成数据上训练的教师模型为真实世界数据生成高质量伪标签，从而在此前提下极大提升了模型的细节表现与泛化能力。作者建立了一个涵盖相机位姿估计、任意视图几何及视觉渲染的新基准，测试结果显示 DA3 在所有任务上均达到最新的 SOTA 水平，在相机位姿准确率和几何准确率上分别平均超越前代 SOTA 模型 35.7% 和 23.6%，且在单目深度估计上也优于 Depth Anything 2。

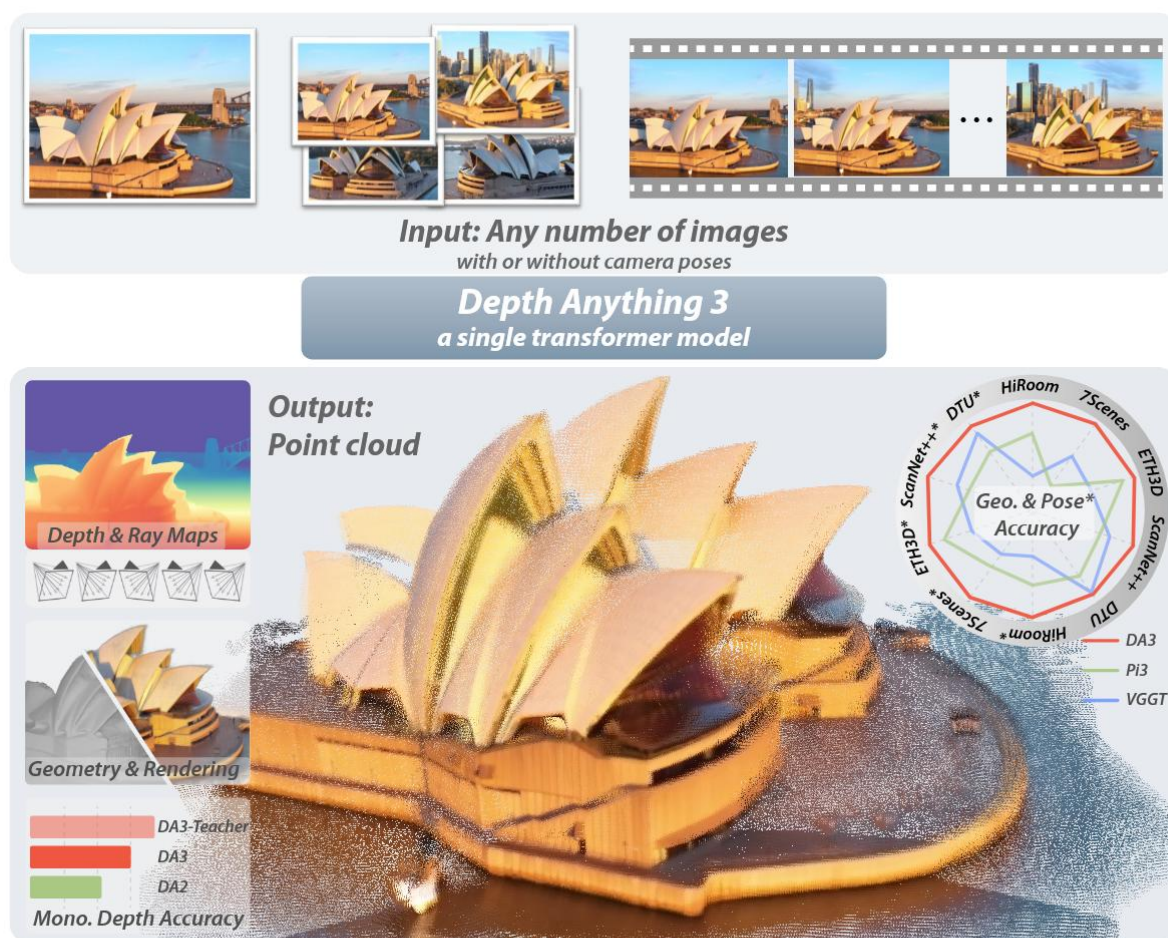


Figure 1 Given any number of images and optional camera poses, **Depth Anything 3** reconstructs the visual space, producing consistent depth and ray maps that can be fused into accurate point clouds, resulting in high-fidelity 3D Gaussians and geometry. It significantly outperforms VGTT in multi-view geometry and pose accuracy; with monocular inputs, it also surpasses Depth Anything 2 while matching its detail and robustness.

主要亮点

- **极简且通用的单一 Transformer 架构：** DA3 摒弃了传统 3D 视觉任务中针对不同输入模式（单目、多视图、视频）设计特定复杂模块的范式，回归到“极简建模”的本源。该模型仅使用一个标准的预训练 Vision Transformer 作为骨干网络，通过这一简单的单一架构，统一解决了从单张图像到大规模多视图序列的几何重建问题。无论输入数据是否包含相机参数（Pose），模型都能通过自适应机制灵活处理，证明了普通 Transformer 在理解 3D 物理世界方面的强大通用能力。

- **创新的“深度-射线”联合表征与 Dual-DPT 设计：**为了避免直接预测旋转矩阵的复杂性以及多任务学习中的冗余, 文章提出了一种极简但充分的“深度-射线” (Depth-Ray) 预测目标。通过设计新颖的 Dual-DPT 解码头, 模型能够联合输出像素对齐的深度图和射线图 (Ray Map), 这种隐式编码相机位姿的方式不仅规避了正交性约束难题, 还实现了场景结构与相机运动的高效解耦与协同预测, 被证明比传统的点云图 (Point Map) 回归更为高效且准确。
- **基于合成数据教师的高效训练策略与卓越性能：**针对真实世界数据中深度标签稀疏或噪声大的问题, DA3 采用了一种增强的教师-学生学习策略。作者训练了一个仅基于大规模合成数据的高性能单目教师模型, 为真实数据生成包含精细几何细节的伪标签, 并将其与噪声较大的真实测量值对齐。这种方法使 DA3 在仅使用公开学术数据集的情况下, 不仅在新建的视觉几何基准上大幅超越 VGGT 和 Pi3 等现有 SOTA 模型, 还能够作为强大的几何骨干网络 (Backbone) 用于提升前馈高斯泼溅 (3DGS) 等下游任务的新视点合成质量。

相关链接

Paper: <https://arxiv.org/pdf/2511.10647>

Github: <https://github.com/ByteDance-Seed/depth-anything-3>

● 论文十五

TiDAR: Think in Diffusion, Talk in Autoregression

日期

2025-11-13



2025-11-13

TiDAR: Think in Diffusion, Talk in Autoregression

Jingyu Liu^{*1}, Xin Dong^{*}, Zhifan Ye², Rishabh Mehta, Yonggan Fu, Vartika Singh, Jan Kautz, Ce Zhang¹, Pavlo Molchanov

NVIDIA

^{*}Equal Contribution; XD as project lead; Corresponding: jingyu6@uchicago.edu, xind@nvidia.com

内容摘要

扩散语言模型具有快速并行生成的潜力，而自回归（AR）模型因其因果结构天然契合语言建模，通常在生成质量上表现更优。现有方法未能有效平衡这两方面，要么是使用较弱模型进行顺序起草的投机解码（牺牲了起草效率），要么是采用某种形式的类 AR 解码逻辑进行扩散（导致质量下降且丧失并行优势）。为此，本文提出了 TiDAR，这是一种序列级混合架构，能够在单次前向传递中通过专门设计的结构化注意力掩码，利用扩散模型起草 token（即“思考”），并以自回归方式采样最终输出（即“表达”）。该设计充分利用了 GPU 的闲置计算密度，在起草和验证能力之间取得了强有力的平衡。广泛的评估表明，TiDAR 在吞吐量上优于投机解码，并在效率和质量上超越了 Dream 和 Llada 等扩散模型。值得注意的是，TiDAR 是首个在消除与 AR 模型质量差距的同时，实现每秒生成 token 数提升 4.71 倍至 5.91 倍的架构。

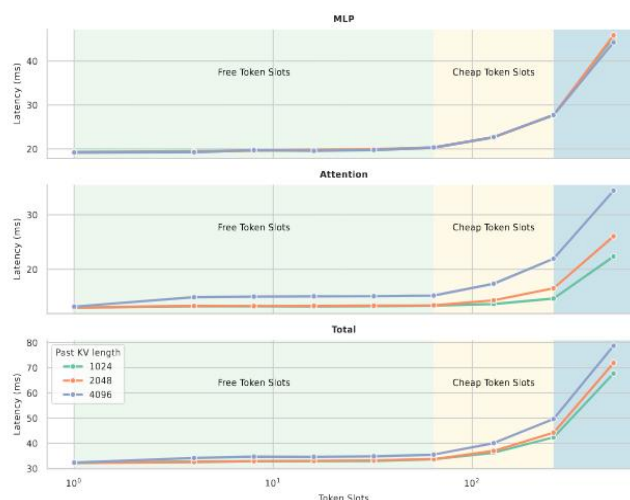


Figure 1 | **Latency Scaling over Token Slots:** We plot the latency of Qwen3-32B decoding on NVIDIA H100 with batch size=1 and Flash Attention 2 [15] over different prefix lengths. Latency stays relatively the same with a certain amount of tokens sent to forward (free + cheap token slots), before transitioning to the compute-bound regime. We leverage this characteristic to achieve almost free parallelised drafting and sampling for TiDAR.

主要亮点

- **创新的“扩散思考-自回归表达”混合架构：**TiDAR (Think in Diffusion, Talk in Autoregression) 提出了一种独特的序列级混合架构，旨在解决传统扩散模型生成质量低和自回归模型推理速度慢的困境。该架构在单次模型前向传递中，利用扩散机制从边缘分布中并行起草 (Drafting) 多个候选 Token，即“思考”过程；随后利用自回归机制从联合分布中进行高质量的拒绝采样 (Sampling)，即“表达”过程。这种设计无需为了并行性而牺牲因果结构的严谨性，从而在根本上结合了扩散模型的速度优势与自回归模型的质量优势。
- **基于单模型单次前向的高效并行投机生成：**与传统的投机解码 (Speculative Decoding) 需要额外的小型草稿模型不同，TiDAR 作为独立的单一大模型运行，且从系统层面极大地优化了推理效率。它利用解码过程中的“空闲 Token 槽位” (Free Token Slots)，在不显著增加延迟的情况下并行计算起草和验证。通过专门设计的结构化注意力掩码 (Structured Attention Masks) 和对精确 KV 缓存 (Exact KV Cache)

的支持，TiDAR 能够在一次前向传播中同时完成 Token 的并行预测与验证，从而实现了比传统投机解码更高的计算密度和吞吐量，在 1.5B 和 8B 规模上分别实现了 4.71 倍和 5.91 倍的加速。

- **首个消除质量差距并支持双模态训练的扩散架构:** TiDAR 是首个在保持与同等规模自回归模型（如 Qwen 系列）无损质量的同时，实现数倍推理加速的架构。文章提出了一种简化的训练配方，即在扩散部分使用全掩码（Full Mask）策略，这不仅简化了损失平衡，增强了扩散损失信号，还消除了训练与测试时的一致性。在代码生成（HumanEval, MBPP）和数学推理（GSM8K）等生成性任务以及似然性任务评估中，TiDAR 均证明了其性能不仅远超 Dream 和 Llada 等现有扩散模型，还能够与微调后的顶级 AR 模型相媲美，证明了其架构在高精度需求场景下的实用性。

相关链接

Paper: <https://arxiv.org/pdf/2511.08923>

● 论文十六

Instella: Fully Open Language Models with Stellar Performance

日期

2025-11-13

Instella: Fully Open Language Models with Stellar Performance

Jiang Liu, Jialian Wu, Xiaodong Yu, Yusheng Su, Prakamya Mishra, Gowtham Ramesh, Sudhanshu Ranjan, Ximeng Sun, Ze Wang, Chaitanya Manem, Pratik Prabhanjan Brahma, Zicheng Liu, Emad Barsoum

AMD



<https://huggingface.co/amd/Instella-3B>



<https://github.com/AMD-AGI/Instella>

内容摘要

大型语言模型（LLM）在广泛的任务中展现了卓越的性能，然而大多数高性能模型仍然保持闭源或仅部分开放，这限制了技术的透明度和可复现性。本文介绍了 Instella，一个完全开源的 30 亿参数语言模型家族，其完全基于公开可用的数据和代码库进行训练。该模型由 AMD Instinct™ MI300X GPU 驱动，经过了大规模预训练、通用指令微调以及人类偏好对齐的开发流程。尽管所使用的预训练 token 数量显著少于许多同类模型，Instella 依然在完全开源模型中取得了最先进的结果，并能够与同等规模的领先“开放权重”模型相媲美。此外，我们还发布了两个专用变体：支持高达 128K 上下文长度的 Instella-Long，以及通过监督微调和强化学习增强数学任务能力的推理专注型模型 Instella-Math。这些贡献共同确立了 Instella 作为一个透明、高性能且多功能的社区替代方案，推动了开放和可复现语言模型研究的目标。

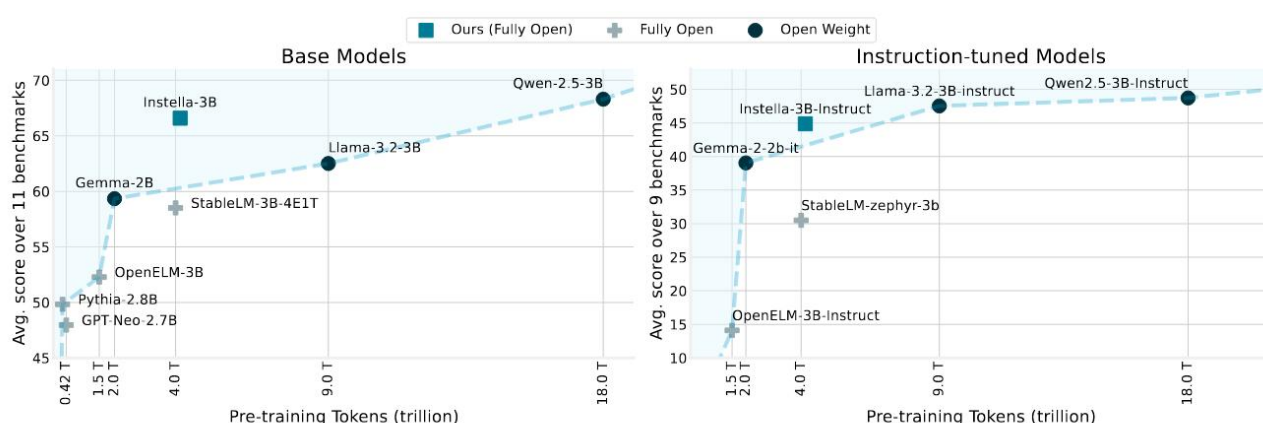


Figure 1: Average Score versus Pre-training Tokens for base (left) and instruction-tuned (right) models. Instella surpasses prior fully open models of comparable size and, despite being trained on substantially fewer pre-training tokens, achieves competitive performance with state-of-the-art open-weight models for both (left) base models (Table 4) and (right) instruction-tuned models (Table 6).

主要亮点

- **全栈开源范式与卓越的效能比:** Instella 突破了当前仅发布模型权重（Open-Weight）的局限，提供了一个真正意义上的“全栈开源”范本。团队不仅开源了模型权重，还完全公开了训练代码、数据配方（Data

Recipes) 和评估流程, 致力于解决科研中的透明度和可复现性问题。尽管其预训练数据量远少于 Llama-3.2 或 Qwen-2.5 等竞品, 但得益于 AMD Instinct MI300X GPU 的算力支持和优化的训练策略, Instella-3B 在多个基准测试中不仅击败了所有同类完全开源模型, 还展现出与其数倍数据量训练出的顶尖开放权重模型相媲美的竞争力。

- **创新的分阶段预训练与合成数据策略:** 为了在有限数据下提升性能, 论文提出了一种精细的分阶段预训练策略: 首先在 4T 通用 token 上建立基础, 随后在一个由 58B 高质量、推理密集型 token 组成的第二阶段中进行强化。该阶段特别引入了自研的合成数学数据集, 通过将 GSM8K 问题抽象为符号化 Python 程序来生成多样化且保证正确性的训练样本。此外, 团队还采用了创新的“权重集成”(Weight Ensembling) 技术, 通过合并不同随机种子训练出的第二阶段模型权重, 进一步提升了最终模型的稳定性和综合性能。
- **基于合成 QA 与强化学习的长文本及推理增强:** Instella 家族通过针对性技术扩展了模型在长上下文和数学推理上的能力。对于 Instella-Long, 团队利用短文生成的合成长文本问答对 (Synthetic QA) 进行持续预训练, 实现了 128K 上下文窗口, 并在 Helmet 基准上表现优异。对于 Instella-Math, 团队采用了多阶段组相对策略优化 (GRPO) 的强化学习方法, 并完全基于开源数据集进行训练。该方法显著提升了小模型的逻辑推理能力, 使其在 TTT-Bench 等策略推理基准测试中取得了完全开源模型中的最佳成绩, 证明了在小参数模型上应用 RL 的有效性。

相关链接

Paper: <https://arxiv.org/pdf/2511.1062>

Code: <https://github.com/AMD-AGI/Instella>

● 论文十七

MonkeyOCR v1.5 Technical Report: Unlocking Robust Document Parsing for Complex Patterns

日期

2025-11-13

MonkeyOCR v1.5 Technical Report: Unlocking Robust Document Parsing for Complex Patterns

Jiarui Zhang¹, Yuliang Liu², Zijun Wu¹, Guosheng Pang¹, Zhili Ye¹, Yupei Zhong¹, Junteng Ma¹, Tao Wei¹, Haiyang Xu¹, Weikai Chen¹, Zeen Wang¹, Qiangjun Ji¹, Fanxi Zhou¹, Qi Zhang¹, Yuanrui Hu¹, Jiahao Liu¹, Zhang Li², Ziyang Zhang², Qiang Liu¹, Xiang Bai²

¹ KingSoft Office Zhuiguang AI Lab, ² Huazhong University of Science and Technology

内容摘要

本文介绍了 MonkeyOCR v1.5，这是一个统一的视觉-语言框架，旨在解决文档解析中复杂布局 and 结构识别的难题，特别是在处理多层嵌套表格、嵌入式图像及跨页结构方面。针对现有 OCR 系统在复杂场景下经常出现的错误传播和对齐困难问题，该模型采用了一种创新的两阶段处理流程：第一阶段利用大型多模态模型（LMM）联合预测文档布局和阅读顺序，利用全局视觉信息确保顺序一致性；第二阶段则对检测到的区域进行局部内容（文本、公式、表格）识别。为了进一步提升复杂表格的解析能力，作者提出了一种基于视觉一致性的强化学习方案，通过“渲染-对比”机制评估识别质量，在无需大量人工标注的情况下显著提升了结构准确性。此外，文章还引入了“图像解耦表格解析”和“类型引导表格合并”模块，有效解决了表格内嵌图像干扰及跨页、跨栏表格的自动重建问题。实验表明，MonkeyOCR v1.5 在 OmniDocBench v1.5 等多个基准测试中取得了最先进（SOTA）的性能，综合表现显著优于 PPOCR-VL 和 MinerU 2.5。

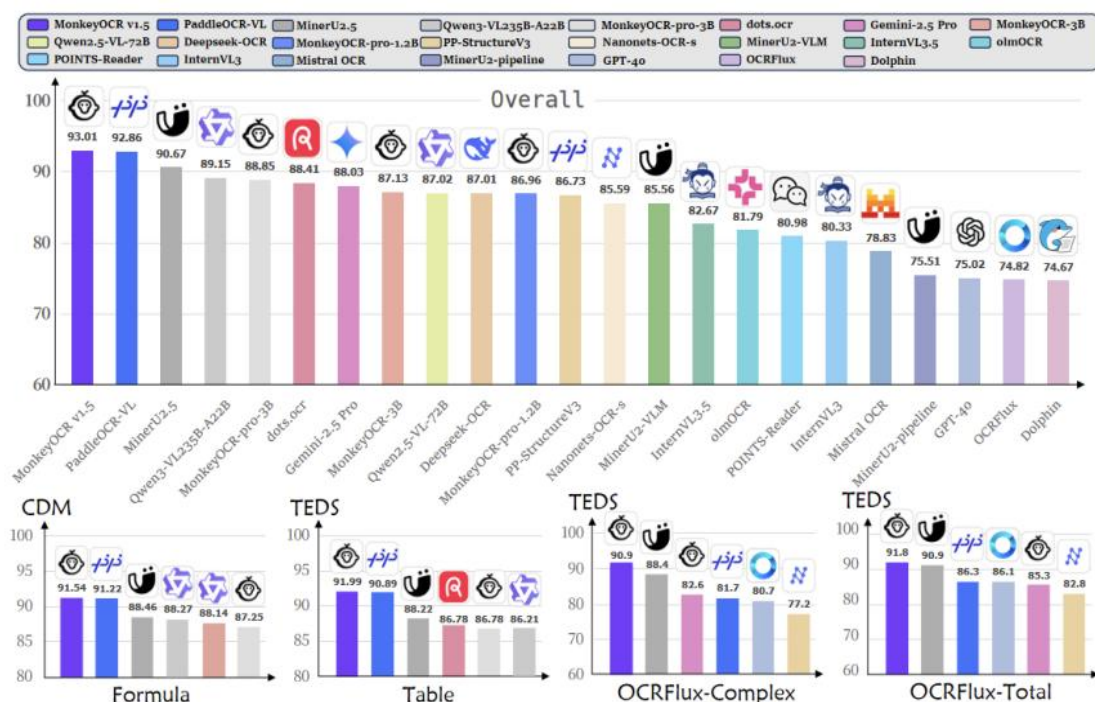


Figure 1: Performance comparison of MonkeyOCR v1.5 and other SOTA models.

主要亮点

- 高效的视觉-语言两阶段解析架构：** MonkeyOCR v1.5 摒弃了传统 OCR 繁琐的多模型堆叠和高计算成本的全图端到端生成模式，提出了一种精简的两阶段视觉-语言管线。在第一阶段，模型利用强大的语义理解能力，联合输出高精度的布局检测框与阅读顺序索引，解决了传统方法难以处理复杂排版逻辑的问题；第二阶段则针对裁剪后的区域进行并行内容识别。这种设计既保留了多模态大模型对全局视觉语境的掌控力，又通过局部识别保证了高分辨率文档的细节保真度，有效减少了级联误差并提升了整体解析的鲁棒性。
- 基于视觉一致性的强化学习表格优化：** 针对复杂表格结构数据稀缺且标注昂贵的痛点，本文提出了一种基于视觉一致性的强化学习（Reinforcement Learning）范式：该方法训练了一个奖励模型，用于计算模型预测出的表格结构（HTML 代码）经渲染后与原始图像之间的视觉相似度。通过这种“预测-渲染-对比”的闭环反馈机制，模型能够利用大量无标注数据进行自我优化，显著增强了对复杂表格线框、单元格合并

及非规则结构的对齐能力，在 OCRFlux-complex 等高难度数据集上展现出了超越现有 SOTA 模型的卓越性能。

- **突破性的复杂版式重建技术：**为了解决现有模型难以处理的极端文档场景，MonkeyOCR v1.5 引入了两项关键技术模块：图像解耦表格解析（IDTP）和类型引导表格合并（TGTM）。IDTP 通过检测并遮蔽表格内的嵌入图像，生成占位符结构后再回填图像，完美解决了图文混排表格的解析难题；TGTM 则结合规则匹配与 BERT 语义分类，智能识别并合并跨页或跨栏的断裂表格（包括重复表头、无表头续表等模式）。这些技术使得模型在处理报纸、学术论文等复杂真实文档时，能够生成结构完整且逻辑连贯的文档副本。

相关链接

Paper: <https://arxiv.org/pdf/2511.10390>

● 论文十八

Hierarchical Frequency Tagging Probe (HFTP): A Unified Approach to Investigate Syntactic Structure Representations in Large Language Models and the Human Brain

日期

2025-11-24

Hierarchical Frequency Tagging Probe (HFTP): A Unified Approach to Investigate Syntactic Structure Representations in Large Language Models and the Human Brain

Jingmin An¹, Yilong Song¹, Ruolin Yang¹, Nai Ding², Lingxi Lu³
Yuxuan Wang¹, Wei Wang⁴, Chu Zhuang¹, Qian Wang^{1,†}, Fang Fang^{1,†}
¹Peking University, ²Zhejiang University, ³Beijing Language and Culture University,
⁴Beijing Institute for General Artificial Intelligence
anjm@stu.pku.edu.cn, {wangqianpsy, ffang}@pku.edu.cn

内容摘要

本文提出一种创新性神经-人工智能对齐探针——层级频率标记探针（Hierarchical Frequency Tagging Probe, HFTP），统一用于揭示和比较大语言模型（LLMs）与人脑在句法结构表征方面的异同。受神经科学中层级频率标签（HFT）范式启发，作者通过构建结构化的语言材料（如 4 词/4 音节构成短语和句子），将 LLM 每层 MLP 神经元的激活序列转为频域信号，并通过傅里叶变换精准定位对句子、短语等不同句法层级敏感的“功能神经元”。同理，作者在大规模人脑 sEEG 皮层电信号中也采用相同频域分析，标记出在人类听觉语句加工中分工明确的皮层区域。研究发现，GPT-2、Gemma、Llama 等主流 LLM 均在中后层显著编码句法层级，且表现出高于右脑的左脑对齐性——左脑为人类语言主导半球。更有趣的是，模型升级未必带来更“类脑”的句法表征：Llama 3.1 对齐度低于 Llama 2，而 Gemma 2 优于 Gemma。该工作不仅为理解 LLM “类人”能力来源提供了理论工具，也为推动 AI 神经科学交叉研究和可解释性 AI 模型设计提供了新范式。

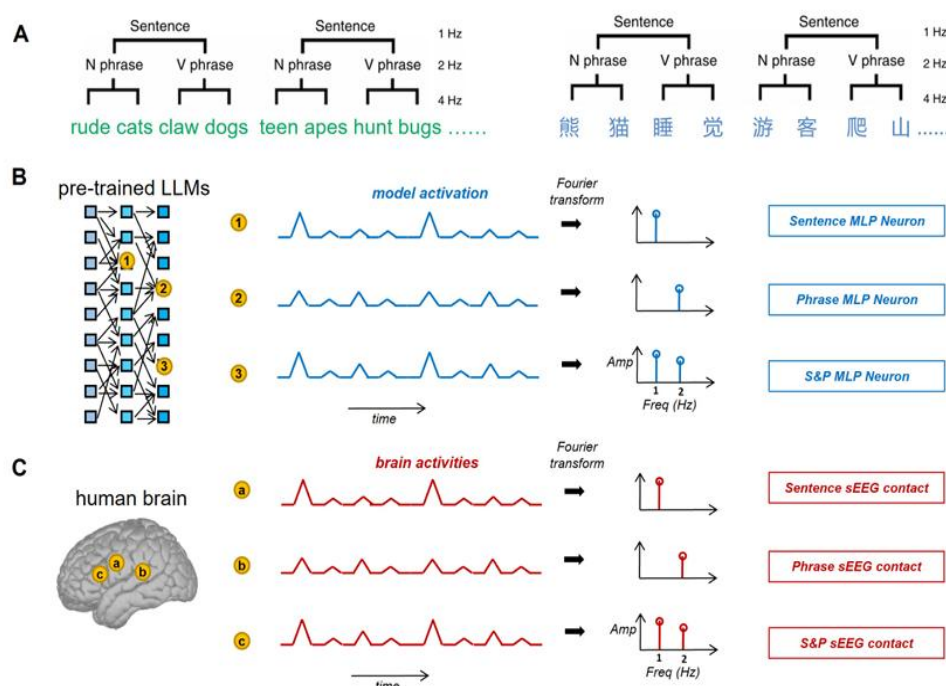


Figure 1: A framework for Hierarchical Frequency Tagging Probes (HFTP) and an illustration of neurons involved in different levels of hierarchical linguistic processing in both LLMs and the human brain. **A**, hierarchical linguistic structure in English and Chinese including syllable, phrase, and sentence. **B**, hierarchical linguistic pattern (1 Hz: sentence feature, 2 Hz: phrase feature) observed both in LLMs and **C**, human brain.

主要亮点

- **首创统一神经-人工智能句法表征探针，实现跨模态、跨层级对齐分析。**HFTP 方法首次实现了对 LLM 内部单神经元和人脑皮层区块的“频域解剖”，可在同一分析体系下，系统检测并量化句子、短语等不同句法层级的结构表征。通过傅里叶分析和置换检验，作者在 Transformer MLP 子层内精确标注了“句子神经元”“短语神经元”与“共敏感神经元”，并在 sEEG 数据中同步标记皮层通道，实现了 LLM 与人脑在句法处理上的频域表征对齐。该方法不仅适用于高度结构化材料，还可泛化至日常自然语料，在 AI 与认知神经科学领域具有广泛适用性。
- **揭示 LLM 句法结构表征的层级分布及模型演化对“类脑性”的影响。**实验比较表明，尽管主流 LLM 如 GPT-2、Gemma、Llama 等都能在不同层级编码句法层级，具体分布却因模型架构和训练而异。如 Llama/GLM 多在后层编码，GPT-2 中层更活跃，Gemma 前层更集中。更重要的是，升级模型不一定更“类脑”：Llama 3.1 的“句法神经元”比例和与人脑对齐度均低于 Llama 2，Gemma 2 则高于 Gemma。这一发现提示，AI 性能提升可能伴随机制的“去类脑化”，为未来优化类脑可解释模型提供了重要启示。
- **首次大规模量化 LLM-人脑句法对齐，揭示左脑主导与跨脑区分工特征。**通过对六大 LLM 与 26 位受试者 sEEG 数据的跨模态对齐分析，作者发现所有模型与人类左脑（语言主导半球）的句法表征相关性显著高于右脑，且关键对齐脑区（如 A1、STG、MTG、IFG）与经典语言脑区高度重合。进一步的区域贡献分析显示，不同模型在具体脑区的对齐贡献有细微差异，揭示了模型与脑区功能分工的潜在映射关系。这些结果不仅为 AI 神经对齐研究树立了新标杆，也为未来基于神经机制的 AI 模型优化和脑机接口等应用提供了理论基础。

相关链接

Paper: <https://arxiv.org/pdf/2510.13255>

● 论文十九

MathCanvas: Intrinsic Visual Chain-of-Thought for Multimodal Mathematical Reasoning

日期

2025-11-26

MathCanvas: Intrinsic Visual Chain-of-Thought for Multimodal Mathematical Reasoning

Weikang Shi^{1*} Aldrich Yu^{1*} Rongyao Fang^{1*†} Houxing Ren¹ Ke Wang¹ Aojun Zhou¹ Changyao Tian¹
Xinyu Fu² Yuxuan Hu¹ Zimu Lu¹ Linjiang Huang³ Si Liu³ Rui Liu^{2‡} Hongsheng Li^{1‡}

¹Multimedia Laboratory (MMLab), The Chinese University of Hong Kong,

²Huawei Research, ³BUAA

wkshi@link.cuhk.edu.hk hqli@ee.cuhk.edu.hk

内容摘要

本文提出了 MathCanvas，首个赋予统一大规模多模态模型（LMM）“内生视觉链式思维”（Visual Chain-of-Thought, VCoT）能力的完整框架，显著提升了模型在几何、函数等高度依赖图形推理的数学领域的解题能力。传统链式思维(CoT)在数学推理中依赖文本，但在人类思维中，视觉化的辅助（如作图、动态分析）是不可或缺的。现有 VCoT 方法多依赖外部工具或静态图片，难以实现高保真、可编辑的动态图形生成和推理。MathCanvas 通过两阶段训练：第一阶段（视觉操作）利用新构建的千万级数学图形生成与编辑大数据集，预训练模型强大的图形合成和编辑技能；第二阶段（策略性视觉-文本协同推理）则用大规模交错图文推理数据，教会模型在解题过程中“何时画、画什么、如何用图助推后续逻辑”。为评测模型真正的视觉推理能力，作者还构建了涵盖 3000 道多步图文题的 MathCanvas-Bench 基准，涵盖多种数学分支。实验证明，基于该框架训练的 BAGEL-Canvas 在 MathCanvas-Bench 上对比强基线提升 86%，并在多个公开多模态数学基准上大幅领先，显示出强大的泛化与创新能力。本文为数学 AI 领域提供了完整的工具链和方法论，为通用 AI 在人类视觉推理上的突破奠定基础。

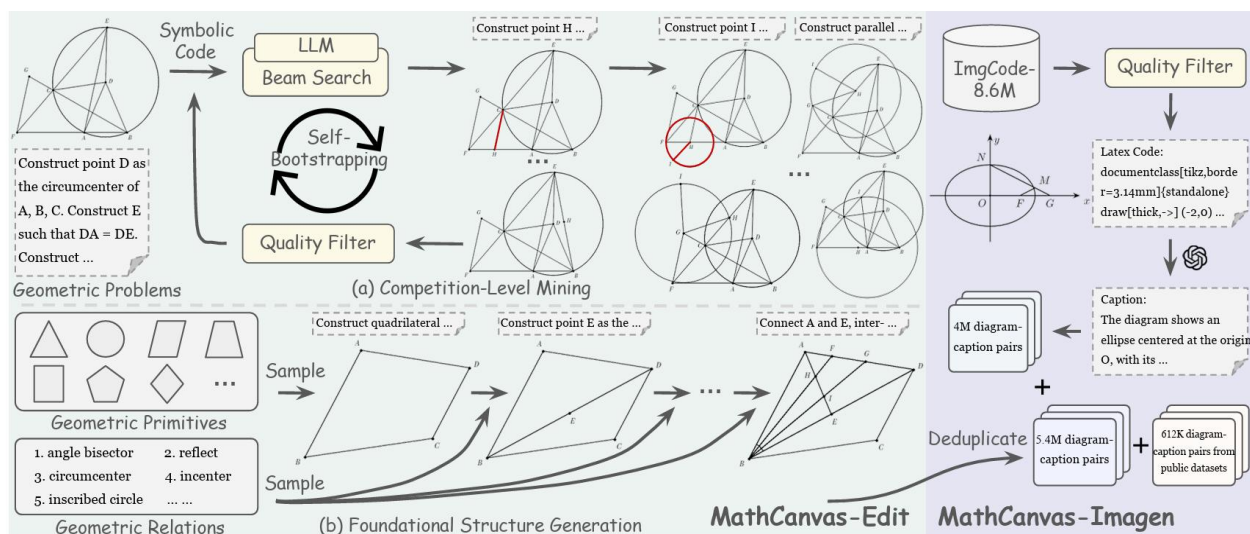


Figure 2: The curation pipeline for the MathCanvas-Edit and MathCanvas-Imagen dataset.

主要亮点

- 首创千万级数学图形编辑与生成大数据集，推动 LMM 内生视觉推理能力。** MathCanvas 构建了业界首个针对数学推理的千万级图形编辑与生成数据集（MathCanvas-Edit 与 MathCanvas-Imagen），涵盖从基础几何、函数到复杂竞赛级构造的动态编辑轨迹与高质量图文对齐样本。通过程序合成+大模型自动挖掘，系统覆盖了图形增删、辅助线作图等多种编辑操作，使模型在预训练阶段掌握了“类人”动态作图与编辑逻辑。这一视觉基础为后续的视觉-文本协同推理提供了坚实支撑，也为多模态 AI 领域提供了大规模、结构化的数据资源。
- 提出首个超 20 万例的交错式视觉-文本链式推理训练集与严格评测基准。** 作者精心设计了 MathCanvas-Instruct 数据集，包含 21.9 万条题解交错（文本-图形-文本-图形...）的多步推理路径，真实还原人类解题时的“画图—推理—再画图—进一步推理”流程。为科学评价模型的 VCoT 能力，提出了 MathCanvas-Bench 基准（3000 题），覆盖代数、几何、分析、统计等全谱系知识点，要求模型不仅给出最终答案，还需在解题过程中动态生成关键图形步骤。采用 GPT-4.1 自动解析与评分，确保评测客观、可扩展。这些资源填补了多模态数学 AI 领域训练与评测的关键空白。

- **实证 BAGEL-Canvas 在多模态数学推理全方位突破并具强泛化能力。**在 MathCanvas 框架下训练的 BAGEL-Canvas 在 MathCanvas-Bench 上取得 34.4% 的 weighted 得分，较原始模型提升 15.9 分，远超行业主流开源/闭源大模型，尤其在几何、三角、平面/立体几何等依赖视觉直观的领域提升最大（如三角 +27.1 分）。更重要的是，该体系不仅提升了模型作图和视觉链式推理能力，也显著增强了其文本推理的泛化表现——即便在传统纯文本多模态数学基准（如 MathVista、MathVision 等）上也大幅领先。消融实验进一步证明，图形编辑与生成预训练、图文交错训练和推理阶段视觉信息的有机融合，三者缺一不可，是 LMM 迈向“类人”视觉推理的关键路径。

相关链接

Paper: <https://arxiv.org/pdf/2510.14958>

● 论文二十

Artificial Hivemind: The Open-Ended Homogeneity of Language Models (and Beyond)

日期

2025-11-28

Artificial Hivemind: The Open-Ended Homogeneity of Language Models (and Beyond)

Liwei Jiang[♣] Yuanjun Chai[♣] Margaret Li[♣] Mickel Liu[♣] Raymond Fok[♣]
Nouha Dziri[★] Yulia Tsvetkov[♣] Maarten Sap[◇] Alon Albalak^{♣*} Yejin Choi[♡]

[♣]University of Washington [◇]Carnegie Mellon University

[★]Allen Institute for Artificial Intelligence ^{♣*}Lila Sciences [♡]Stanford University

lwjiang@cs.washington.edu

 Code: <https://github.com/liweijiang/artificial-hivemind>

 INFINITY-CHAT Collection: [liweijiang/artificial-hivemind](https://github.com/liweijiang/artificial-hivemind)

内容摘要

本文关注大语言模型 (LLMs) 在开放性生成任务中的输出同质化 (homogenization) 问题, 提出 “Artificial Hivemind” 现象, 即无论是单一模型多次采样 (intra-model repetition) 还是不同模型之间 (inter-model homogeneity), 当前主流 LLM 在面对真实世界的开放性任务时, 均倾向于输出极为相似、缺乏多样性的答案。为系统性研究这一现象, 作者构建了 INFINITY-CHAT 数据集, 涵盖 26,000 条来自真实用户的多样化、开放式查询 (prompt), 并首创了细致的查询类别学 (6 大类、17 子类), 全面反映了现实用户对 LLM 的开放式需求。基于该数据集, 作者对 70 余个主流开源与闭源模型进行了大规模、多轮抽样实证, 发现即便采用高温采样、min-p 等多样性提升策略, 模型输出在语义与表层结构上依然高度趋同, 甚至不同模型族间也频繁出现大段重合文本或相同意象 (如 “时间是条河流” 等隐喻)。更进一步, 通过引入 31,250 条高密度人工评价, 文章还揭示了当前 LLM、Reward Model 和自动评测器在面对人类偏好分歧 (即使整体质量相当) 的多元答案时, 表现出显著的标注失配与对 “多元高质量” 响应的低敏感性。本文不仅为 LLM 输出多样性评价与安全性研究奠定了全新基准, 也为未来 AI 系统如何更好地服务于人类创造与多元价值提供了重要理论与数据资源。

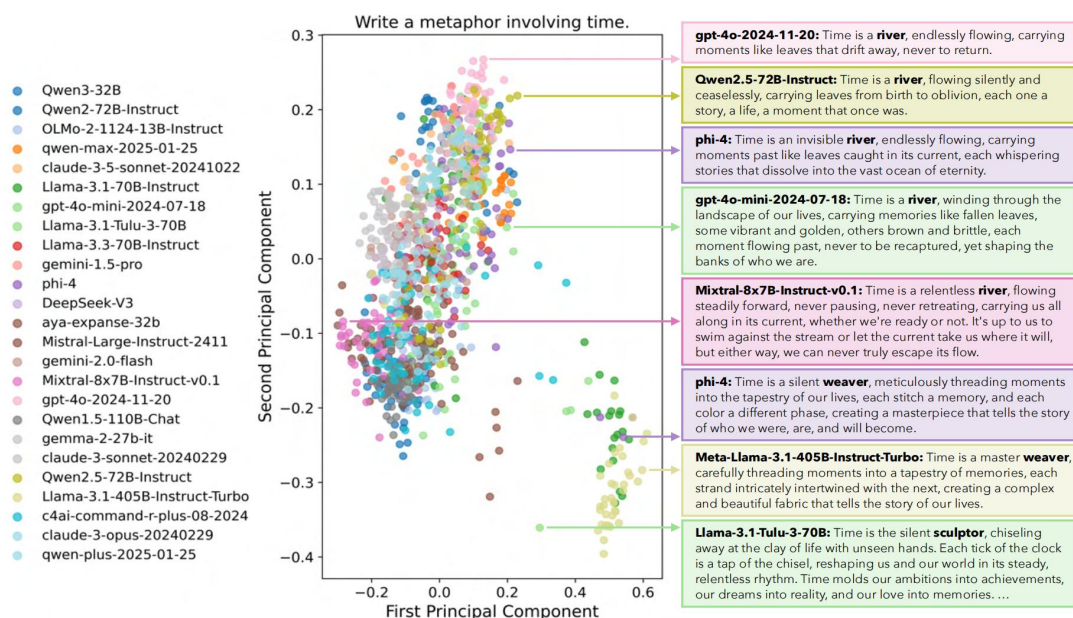


Figure 1: Responses to the query “Write a metaphor about time” clustered by applying PCA to reduce sentence embeddings to two dimensions. Each of the 25 models generates 50 responses using top- p sampling ($p = 0.9$) and temperature = 1.0. Despite the diversity of model families and sizes, the responses form just two primary clusters: a dominant cluster on the left centered on the metaphor “time is a river,” and a smaller cluster on the right revolving around variations of “time is a weaver.”

主要亮点

- **首创真实开放性对话大规模数据集与细粒度任务分类体系。**INFINITY-CHAT 数据集是首个面向真实用户开放式需求的大规模、高质量对话集，系统涵盖了从创意内容生成、思想碰撞、哲学与抽象思辨、分析解释、假设推演到个性化技能建议等 6 大类、17 子类开放性任务，极大丰富了以往仅关注诗歌、随机数等狭窄范畴的多样性测评基准。通过 GPT-4o 自动标注+人工复核，作者确保了分类的合理性与数据的广泛代表性。该数据集不仅为模型行为归因和多样性提升算法的开发提供了坚实基础，也为 AI 社会影响研究（如思想同质化、文化多样性保护）提供了宝贵原始素材。
- **系统揭示 LLM 开放式生成的“人工蜂群”同质化效应及其成因。**作者首次在大规模模型集群（70+ 主流模型）和高密度采样下，定量、可视化地揭示了 LLM 在开放式任务中的“模式塌缩”：同一模型多次采样输出高度相似，不同模型间也频繁出现语义与表述上的收敛（如大量“时间是河流”的比喻）。即使引入 min-p 等多样性采样策略，同质化问题依然严重。分析指出，模型训练数据、对齐方式、合成数据污染等均可能是同质化的深层原因。该发现对当前“模型集成、多模型蒸馏即多样性”的主流假设提出挑战，警示社区：多模型协作下的多样性未必真实存在，需从训练本身和评价体系根本改进。
- **引入密集人工偏好标注，揭示 LLM 与人类多元偏好失配的挑战。**作者为 INFINITY-CHAT 中的开放式任务引入了超 3 万条密集人工标注（每个问题-答案对或答案对比均有 25 人独立评价），首次系统性分析了 LLM、Reward Model 和自动评测器在多元人类偏好下的表现。结果显示，当人类偏好分歧显著或多个答案质量接近时，主流模型与人类一致性急剧下降，难以识别“多元等优”答案。这一发现强调现有 AI 评测与对齐体系对“唯一高质”的过度假设，呼吁未来模型需要更好地理解、表达和评判人类多元价值，避免推动社会思想与表达的单一化。

相关链接

Paper: <https://arxiv.org/pdf/2510.22954>

GitHub: <https://github.com/liweijiang/artificial-hivemind>

● 论文二十一

Gated Attention for Large Language Models: Non-linearity, Sparsity, and Attention-Sink-Free

日期

2025-11-28

Gated Attention for Large Language Models: Non-linearity, Sparsity, and Attention-Sink-Free

Zihan Qiu^{*1}, Zekun Wang^{*1}, Bo Zheng^{*1}, Zeyu Huang^{*2},
Kaiyue Wen³, Songlin Yang⁴, Rui Men¹, Le Yu¹, Fei Huang¹, Suozhi Huang⁵,
Dayiheng Liu^{✉1}, Jingren Zhou¹, Junyang Lin^{✉1}
¹Qwen Team, Alibaba Group ²University of Edinburgh ³Stanford University
⁴MIT ⁵Tsinghua University

内容摘要

本文系统性探讨了门控机制（Gating）在大语言模型标准 softmax 注意力中的效能与机制，提出并验证了一种极其简单但有效的改进：在 Scaled Dot-Product Attention (SDPA) 后引入头特异性 sigmoid 门控。通过对 30 种门控变体在 15B 参数 MoE 模型和 1.7B 参数密集模型、3.5 万亿 token 数据上的对比实验，作者发现该门控不仅显著提升模型性能（如 PPL 下降，MMLU 提升），还极大增强了训练稳定性，允许更高学习率和更大批量。进一步分析显示，门控机制的作用核心在于：（1）在 SDPA 与 Value 映射之间引入非线性，提升低秩线性变换的表达力；（2）通过 query 依赖的稀疏门控分数，动态过滤注意力输出，显著缓解“attention sink”现象——即模型将大量注意力无意义地集中于首 token。实验还表明，门控机制对于长上下文扩展具有重要价值，能显著减缓扩展后的性能衰减。文章不仅细致比较了不同门控位置、粒度、激活函数和参数共享策略，还通过层级分析展示了门控如何带来稀疏、稳定和泛化更强的注意力分布。相关代码与模型也已开放，有望推动 LLM 基础架构的下一步演化。

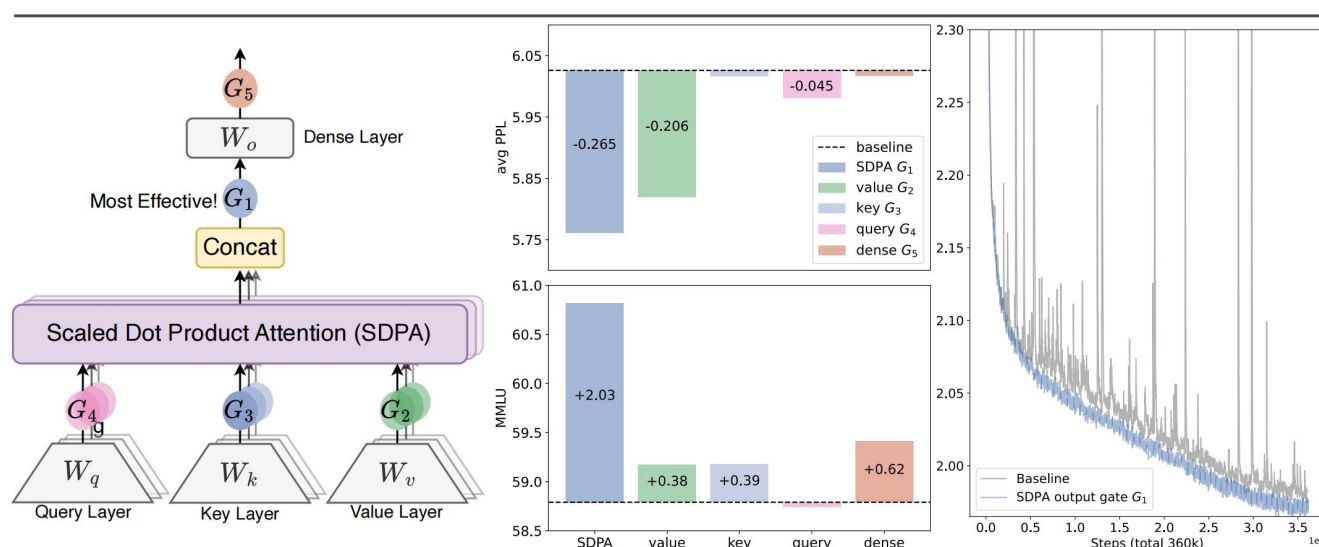


Figure 1: **Left:** Investigated positions for applying gating operations.; **Middle:** Performance comparison (Test PPL and MMLU) of 15B MoE models with gating applied at various positions. Gating after SDPA (G_1) yields the best overall results. Gating after the Value layer (G_2) also demonstrates notable improvements, particularly in PPL. **Right:** Training loss comparison (smoothed, 0.9 coeff.) over 3.5T tokens between baseline and SDPA-gated 1.7B dense models under identical hyperparameters. Gating results in lower final loss and substantially enhanced training stability, mitigating loss spikes. This stability allows for potentially higher learning rates and facilitates better scaling.

主要亮点

- 极简门控机制显著提升性能与训练稳定性。作者首次在主流 Transformer 软最大注意力机制后加入头特异 sigmoid 门控，仅增加极少参数，即可带来广泛性能提升（如 PPL 下降 0.2，MMLU 提升 2 分）。该门控有效缓解训练过程中的 loss spike，使模型可在更大学习率和批量下稳定收敛。相比于增加参数量（如扩展头数、专家数），门控机制带来的收益更为直接且高效，适用于 MoE 和密集模型。实验证明，头特异性、乘性 sigmoid 门控效果优于其他变体（如加性或参数共享），且对不同任务和数据分布均有良好适应性。
- 门控增强非线性表达与稀疏性，根治“attention sink”。论文深入揭示门控机制的本质作用：一方面，在 Value 与 Output 映射之间引入非线性，突破低秩线性变换的表达瓶颈；另一方面，门控分数天然稀疏且 query 依赖性强，使模型学会动态过滤无关信息，从而避免注意力在特定 token 上出现集中（attention sink）。层级分析显示，稀疏门控大幅降低首 token 注意力分数比例（从 46.7% 降至 4.8%），并减少大

规模激活，有效提升模型的数值稳定性。与仅靠归一化或剪裁不同，门控机制实现了对注意力分布和激活结构的本质优化。

- 门控机制助力长上下文泛化，提升模型扩展鲁棒性。作者进一步验证了门控机制对长序列扩展的实际价值：在将上下文长度扩展至 128k 时，带有门控的模型性能衰减显著低于基线（如 RULER 任务性能提升 10 分以上），表现出更强的泛化与迁移能力。分析表明，attention sink 模式在上下文扩展时难以自适应，而门控机制通过稀疏、输入依赖的信息流调控，使模型在扩展后仍能保持高效分布与理解能力。该发现为未来基础模型在长文本、复杂推理等场景下的架构设计提供了理论依据和实践范式。

相关链接

Paper: <https://arxiv.org/pdf/2505.06708>

GitHub: https://github.com/qiuzh20/gated_attention

● 论文二十二

1000 Layer Networks for Self-Supervised RL: Scaling Depth Can Enable New Goal-Reaching Capabilities

日期

2025-11-28

1000 Layer Networks for Self-Supervised RL: Scaling Depth Can Enable New Goal-Reaching Capabilities

Kevin Wang
Princeton University
kw6487@princeton.edu

Ishaan Javali
Princeton University
ijavali@princeton.edu

Michał Bortkiewicz
Warsaw University of Technology
michal.bortkiewicz.dokt@pw.edu.pl

Tomasz Trzcíński
Warsaw University of Technology,
Tooploox, IDEAS Research Institute

Benjamin Eysenbach
Princeton University
eysenbach@princeton.edu

内容摘要

本论文聚焦于自监督强化学习（self-supervised RL）的可扩展性，系统性地探讨了神经网络深度扩展对目标导向任务中智能体能力的提升作用。以往强化学习领域广泛采用浅层网络（2-5 层），而作者首次实证表明，将 Contrastive RL（CRL）等自监督算法的网络深度拓展至上千层（最高 1024 层），不仅带来了 2 至 50 倍的性能提升，还促使智能体习得了质性变化的新行为能力。作者在无监督目标条件设定下，对复杂的仿真运动与机械臂操作任务进行实验，发现深度扩展使得模型在长期探索、复杂状态空间导航等任务中表现出跃迁式进步，甚至涌现出如攀爬障碍、策略迁移等先前浅层模型难以实现的新技能。论文还系统分析了网络宽度、批量大小、残差连接、归一化和激活函数等架构要素对可扩展性的作用，指出深度扩展比宽度扩展更能高效提升表现，尤其在高维观测场景中尤为显著。此外，通过对比不同基线方法，作者证明了当前最优的目标导向 RL 性能主要得益于深层自监督网络的表达能力与探索能力的协同增强。研究不仅为 RL 领域的“规模涌现”现象提供了新视角，也为未来构建大规模自主学习智能体奠定了理论与工程基础。

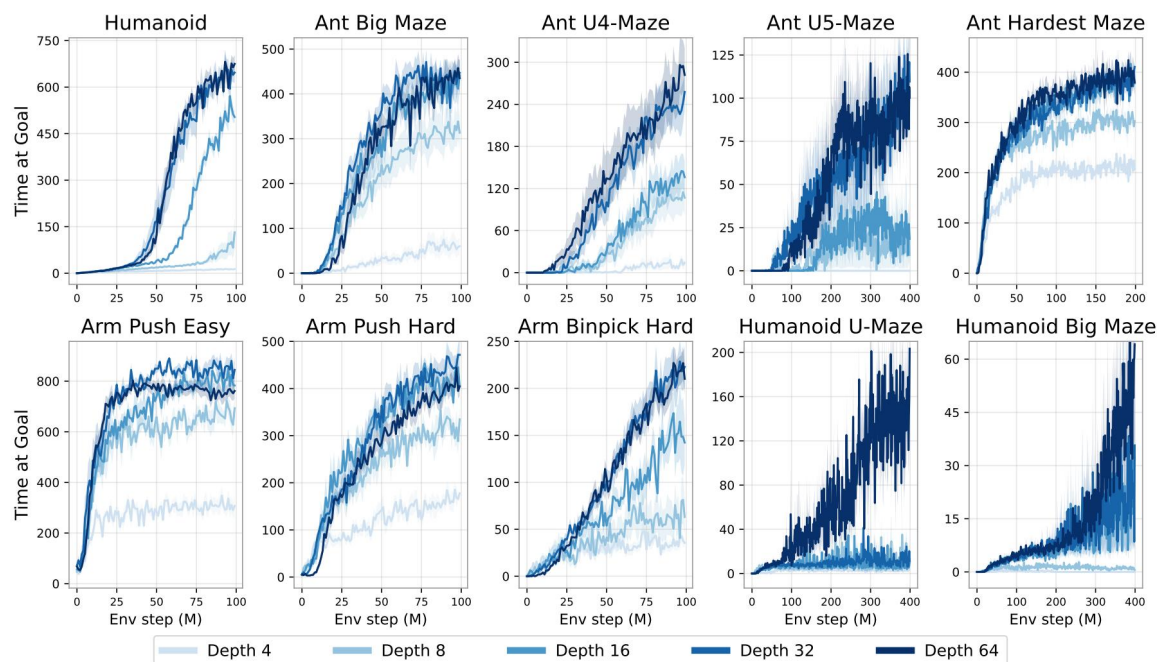


Figure 1: **Scaling network depth yields performance gains** across a suite of locomotion, navigation, and manipulation tasks, ranging from doubling performance to 50× improvements on Humanoid-based tasks. Notably, rather than scaling smoothly, performance often jumps at specific critical depths (e.g., 8 layers on Ant Big Maze, 64 on Humanoid U-Maze), which correspond to the emergence of qualitatively distinct policies (see Section 4).

主要亮点

- **首次实证“深度扩展”在自监督 RL 中的跃迁式增益及能力涌现。** 本文率先将网络深度大规模扩展（高达 1024 层）引入自监督目标导向强化学习，突破了以往 RL 领域多局限于浅层网络的瓶颈。实验结果表明，网络深度的增加并非仅带来线性增益，而是在特定“临界深度”下出现突变式性能跃升和质性行为转变。例如，在 Humanoid 迷宫等高维复杂任务中，智能体从简单的爬行或摔倒，进化到直立行走、越障甚至策略分段“缝合”能力。这类能力的出现，显著拓展了 RL 系统的表达能力与任务适应性，标志着 RL 领域真正实现了类似于语言与视觉大模型中的“规模涌现”效应。
- **深度提升表达性与探索性，协同驱动泛化与数据利用。** 论文通过并行“collector-learner”实验，精准剖析了深度扩展提升 RL 表现的两大机制：一方面，深层网络能够学习到更丰富、更具拓扑感知能力的表征，进而有效捕捉环境结构、实现长时序与复杂空间的目标规划；另一方面，深度带来的表达能力提升也反向促进了探索策略的丰富性，推动智能体更全面地覆盖状态-动作空间。作者进一步发现，只有在探索充分的

数据覆盖下，表达能力的提升才能转化为实际性能增益，否则即使深层网络也会受限于数据瓶颈。这一发现强调了表达性与探索性在自监督 RL 可扩展性中的协同关系。

- **系统性架构创新与广泛适用性验证。** 本文不仅提出了可扩展的深层 RL 网络架构（集成残差连接、归一化、Swish 激活等），还通过消融实验、横向对比和不同基线算法验证，指出只有合理的架构创新才能解锁深度扩展的全部潜力。深度扩展对 CRL 等自监督算法尤为有效，对传统价值或行为克隆算法则作用有限；同时，作者证明了该方法对不同任务、不同环境物理引擎、甚至新型归一化范式（如 Simba-v2 的超球面归一化）均具良好适用性。该系统性创新为未来 RL 大模型、高效分布式训练与新兴自监督算法的扩展提供了坚实的技术和理论基础。

相关链接

Paper: <https://arxiv.org/pdf/2503.14858>

Code: <https://wang-kevin3290.github.io/scaling-crl/>

● 论文二十三

Why Diffusion Models Don't Memorize: The Role of Implicit Dynamical Regularization in Training

日期

2025-11-28

Why Diffusion Models Don't Memorize: The Role of Implicit Dynamical Regularization in Training

Tony Bonnaire[†]
LPENS

Université PSL, Paris
tony.bonnaire@phys.ens.fr

Raphaël Urfin[†]
LPENS

Université PSL, Paris
raphael.urfin@phys.ens.fr

Giulio Biroli
LPENS

Université PSL, Paris
giulio.biroli@phys.ens.fr

Marc Mézard

Department of Computing Sciences
Bocconi University, Milano
marc.mezard@unibocconi.it

内容摘要

本文系统研究了扩散模型 (Diffusion Models, DMs) 为何在训练过程中通常不会完全记忆 (memorize) 训练数据, 而能够实现有效泛化 (generalization) 的内在机制。作者通过理论分析和大量实证实验, 揭示扩散模型训练中存在两种本质不同的时间尺度: 一是模型开始生成高质量、原创样本的早期时间点 (τ_{gen}), 另一是模型开始显著记忆训练数据的滞后时间点 (τ_{mem})。关键发现表明, 随着训练数据规模 n 的增加, 记忆时间 τ_{mem} 近似线性增长, 而泛化时间 τ_{gen} 却基本保持不变, 由此形成一个随 n 增大而显著扩展的“泛化窗口”。在这一时间区间内, 模型即使参数远超数据量, 仍能通过早停 (early stopping) 获得优良的泛化能力, 避免过拟合和训练集复现。此外, 作者在真实图像数据和合成数据上均验证了上述规律, 并通过高维随机特征理论模型进行解析推导, 进一步解释了谱偏置 (spectral bias) 及损失景观的动态特性如何促成隐式动态正则化 (implicit dynamical regularization), 从而赋予扩散模型天然的抗记忆机制。文章不仅拓展了我们对扩散模型记忆-泛化转变本质的理解, 也为今后生成模型的安全应用和优化训练策略提供了理论指导和实践建议。

主要亮点

- **隐式动态正则化机制的揭示。** 作者首次系统地揭示了扩散模型训练过程中的隐式动态正则化现象: 即使在高度过参数化 (overparameterized) 情况下, 只要训练时间控制在泛化窗口内, 模型不会陷入记忆训练

集的陷阱。这种正则化并非源于传统的网络结构限制或显式正则项，而是由训练动态本身——尤其是优化路径和损失景观的谱特性——自然产生。理论分析显示，模型在早期阶段倾向于学习数据分布的低频部分（与真实分布接近），而高频、数据集特异的信息（导致记忆）仅在极长训练时间后才被捕捉，这与深度网络的谱偏置密切相关。该发现为扩散模型的泛化能力提供了全新解释，也为理解大模型泛化与记忆的本质提供了理论依据。

- **记忆与泛化时间尺度分离及其数据依赖性。**文章通过实证和理论模型证明了泛化(τ_{gen})与记忆(τ_{mem})时间尺度的分离现象，并揭示记忆时间与训练集规模之间的线性关系：随着数据量 n 增大，过拟合和记忆的发生被显著延后，而泛化时间几乎不变。这意味着在实际应用中，通过合理早停，不仅能实现高质量生成，还能规避隐私泄露和数据复现等风险。该规律在真实图像数据 (CelebA) 和合成高维数据上均得到验证，并在随机特征模型中获得精确数学刻画。该结果拓展了经典泛化理论在扩散模型中的适用性，强调了数据规模对训练动态的根本影响，并为安全可靠的生成模型设计提供了理论支撑。
- **理论与实验双重验证及广泛适用性。**本文采用多层次验证方法：一方面，在真实 U-Net 架构和多种数据集上进行了详尽实验，充分展示了泛化-记忆转变的普适性和动力学特征；另一方面，借助高维随机特征神经网络理论，精确推导出相关谱分布和训练时间尺度，为实验观察提供坚实数学基础。更重要的是，作者指出这一机制不仅适用于标准扩散模型，也可推广到更广泛的基于 score 的生成范式（如流匹配、随机插值等）。这显示了研究结论的广泛适用性和理论深度，为今后不同类型生成模型的训练与安全性分析提供了通用框架和方法论。

相关链接

Paper: <https://arxiv.org/pdf/2505.17638>

TECHNICAL UPDATES

技术动态

● 技术动态一

月之暗面发布 Kimi-K2，在 HLE 评测基准上超越现有所有开源闭源大模型

日期

2025-11-06



内容摘要

月之暗面推出的 Kimi K2 Thinking 是一款突破性的开源思维模型，其核心设计为一个能够进行复杂、多步推理的智能体，在自主使用工具的情况下可连续执行 200 至 300 次操作而无需人工干预。该模型通过在测试时扩展思维令牌和工具调用步骤，在多项高难度基准测试中创造了新的性能记录，尤其在需要深度推理、智能体搜索和代码生成的领域表现卓越，例如在涵盖众多学科的“人类终极考试”中借助工具达到 44.9% 的准确率，

在评估网络搜索与浏览能力的 BrowseComp 上获得 60.2% 的分数，并在软件工程任务 SWE-Bench Verified 上实现 71.3% 的通过率。Kimi K2 Thinking 不仅在学术和推理问题上表现出色，还能胜任从组件级网站开发到竞争性编程的各类编码任务，并在创造性写作中生成更具深度和情感共鸣的内容。目前，该模型已在 Kimi 官网的聊天模式中上线，其完整的智能体模式即将推出，并可通过 API 进行访问，标志着在实现长链条、自主决策的人工智能系统方面取得了重大进展。

相关链接

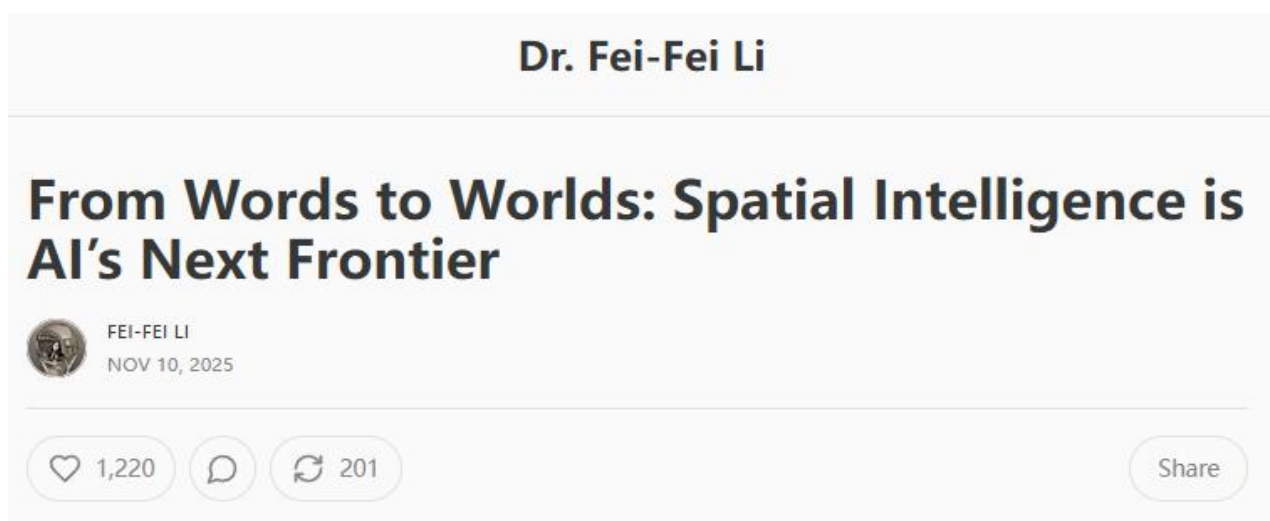
Blog: <https://moonshotai.github.io/Kimi-K2/thinking.html>

● 技术动态二

李飞飞关于空间智能的观点：From Words to Worlds: Spatial Intelligence is AI's Next Frontier

日期

2025-11-10



内容摘要

李飞飞教授提出，空间智能 是人工智能发展的下一个前沿，她认为当前以大型语言模型为代表的 AI 虽在文本和图像处理上表现出色，但缺乏对物理世界的真实理解，无法有效感知三维空间、推理物理规律或与环境进行动态交互；为实现突破，需构建具备生成性（生成符合几何与物理规则的虚拟世界）、多模态（处理图像、文本、动作等多维度输入）和交互性（预测动作结果并动态响应）能力的“世界模型”，其团队已通过 Marble 平台和 RTFM 模型等实践探索，推动 AI 在创意设计、机器人、科学模拟等领域的应用，最终目标是增强人类创造力与认知力，而非取代人类。

相关链接:

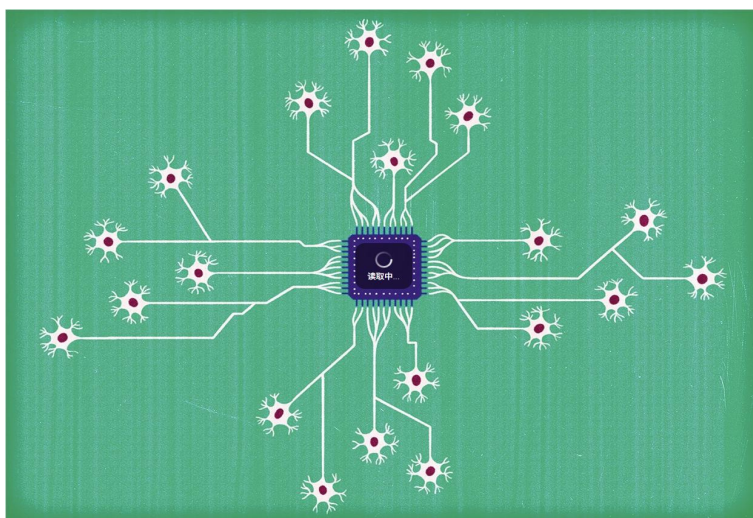
Blog: <https://drfeifei.substack.com/p/from-words-to-worlds-spatial-intelligence>

● 技术动态三

研究人员正在培育人类神经元，并将其转化为生物晶体管

日期

2025-11-11



内容摘要

科学家正在研发用人类神经元（脑细胞）来制造计算机。这些在实验室培养的、沙粒大小的脑细胞团块能够接收和处理电信号，其目标是实现强大的计算能力，同时能耗极低。目前，全球的研究团队可以远程利用这些“生物芯片”进行信息处理任务，这被视为替代传统硅基计算机的一种有潜力的新方向。

相关链接

Blog: <https://www.nature.com/articles/d41586-025-03633-0>

● 技术动态四

百度发布文心一言大模型系列新版本

日期

2025-11-11

LMarena

lmarena.ai

ERNIE 5.0 Preview
scores 1432 in Text Arena

Rank (UB)	Model	Score	95% CI (a)	Votes	Organization	License
1	gemin-2.5-pro	1452	±4	61,890	Google	Proprietary
1	claude-sonnet-4-5-20250929-thinking-32k	1448	±6	12,986	Anthropic	Proprietary
1	claude-opus-4-1-20250805-thinking-16k	1448	±5	28,595	Anthropic	Proprietary
2	gpt-4.5-preview-2025-02-27	1442	±6	14,644	OpenAI	Proprietary
2	claude-opus-4-1-20250805	1439	±4	41,049	Anthropic	Proprietary
2	claude-sonnet-4-5-20250929	1436	±8	5,484	Anthropic	Proprietary
2	ernie-5.0-preview-1022	1432	Preliminary ±13	2,132	Baidu	Proprietary
4	chatgpt-4o-latest-20250326	1438	±4	47,608	OpenAI	Proprietary
4	gpt-5-high	1437	±5	30,138	OpenAI	Proprietary
4	o3-2025-04-16	1433	±4	58,575	OpenAI	Proprietary
4	qwen3-max-preview	1433	±5	25,067	Alibaba	Proprietary
6	glm-4.6	1429	±6	10,511	Z.ai	MIT
8	qwen3-max-2025-09-23	1424	±6	9,302	Alibaba	Proprietary

内容摘要

2025 年 11 月，百度发布文心一言大模型系列两项重要更新：文本模型 ERNIE-5.0-Preview-1022 在全球权威评测 LMArena 中排名第二，即将正式开放；同时，多模态模型 ERNIE-4.5-VL-28B-A3B-Thinking 实现突破性升级，通过强化视觉推理、思维链分析及工具调用能力，在仅激活 30 亿参数下即可对标顶级旗舰模型性能，显著提升了复杂场景下的视觉-语言理解与交互水平。

相关链接

ERNIE-5.0-Preview-1022: <https://yiyao.baidu.com/blog/posts/ernie-5.0-preview-1022-release-on-lmarena>

ERNIE-4.5-VL-28B-A3B-Thinking: <https://yiyao.baidu.com/blog/posts/ernie-4.5-vl-28b-a3b-thinking/>

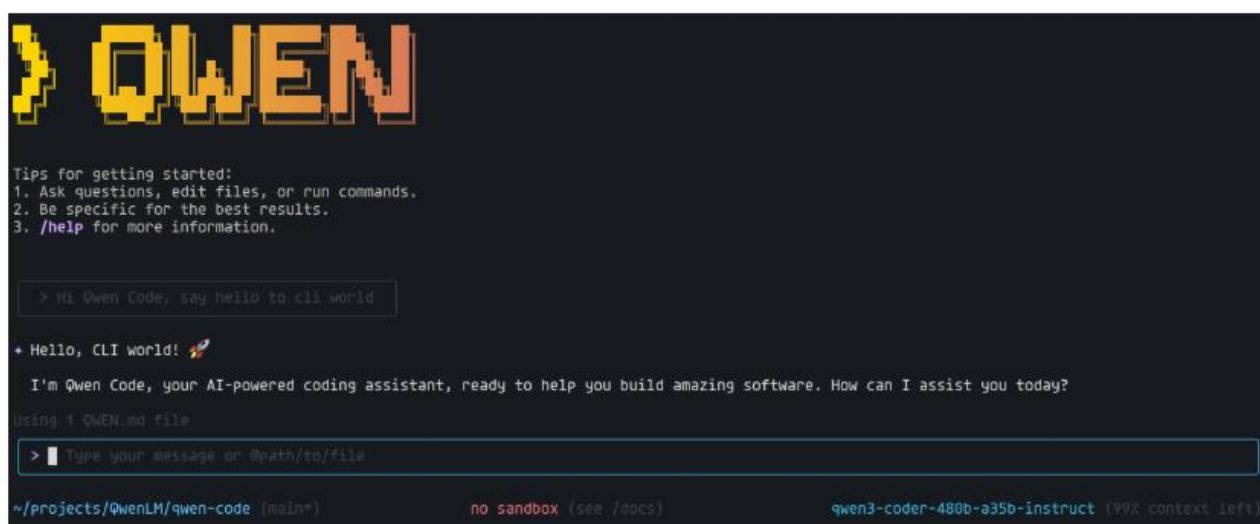
● 技术动态五

阿里巴巴发布 Qwen-Code 0.2.2

日期

2025-11-20

Qwen Code



```
> QWEN

Tips for getting started:
1. Ask questions, edit files, or run commands.
2. Be specific for the best results.
3. /help for more information.

> Hi Qwen Code, say hello to cli world

+ Hello, CLI world! 🚀

I'm Qwen Code, your AI-powered coding assistant, ready to help you build amazing software. How can I assist you today?

Using 1 QWEN.md file

> █ Type your message or @path/to/file

~/projects/QwenLM/qwen-code (main*)    no sandbox (see /docs)    qwen3-coder-480b-a35b-instruct (99% context left)
```

内容摘要

Qwen Code 是一款基于命令行的 AI 工作流工具，专为 Qwen3-Coder 模型优化，旨在提升开发效率。它支持代码理解与编辑、工作流自动化（如处理 Pull Request）、多模态分析（可自动识别图像并调用视觉模型），并提供了灵活的会话管理与视觉模式配置。用户可通过 npm 或 Homebrew 快速安装，推荐使用 Qwen OAuth 认证方式，每日可获得 2000 次免费请求。该工具适用于代码库探索、重构、自动化任务及调试等多种开发场景。

相关链接

GitHub: github.com/QwenLM/qwen-code

● 技术动态六

GPT-5.1 发布: A smarter, more conversational ChatGPT

日期

2025-11-20

2025年11月12日 产品 发布

GPT-5.1 全新上线：更智能、更具对话感的 ChatGPT

我们正在升级 GPT-5，同时让 ChatGPT 的自定义功能更易使用。从今天起开始陆续推出，首先面向付费用户开放。

试用 ChatGPT ↗

内容摘要

OpenAI 正式推出 GPT-5.1 系列，包括 GPT-5.1 Instant 和 GPT-5.1 Thinking 两个版本，显著提升模型的智能性、对话感和指令遵循能力，新增自适应推理功能，并在语气自定义方面优化预设选项，如新增“专业可靠”等风格。该更新将逐步向付费和免费用户推出，并计划集成至 API，旨在使 ChatGPT 更智能、自然且适应用户偏好，同时 GPT-5 系列将在三个月内逐步过渡下线。

相关链接

GPT-5.1: A smarter, more conversational ChatGPT

<https://openai.com/index/gpt-5-1/>

System Card: https://cdn.openai.com/pdf/4173ec8d-1229-47db-96de-06d87147e07e/5_1_system_card.pdf

● 技术动态七

英伟达文档解析大模型 NVIDIA Nemotron Parse v1.1，参数不到 1B，在表格解析、文档版面理解、公式识别等方面性能 SOTA

日期

2025-11-20

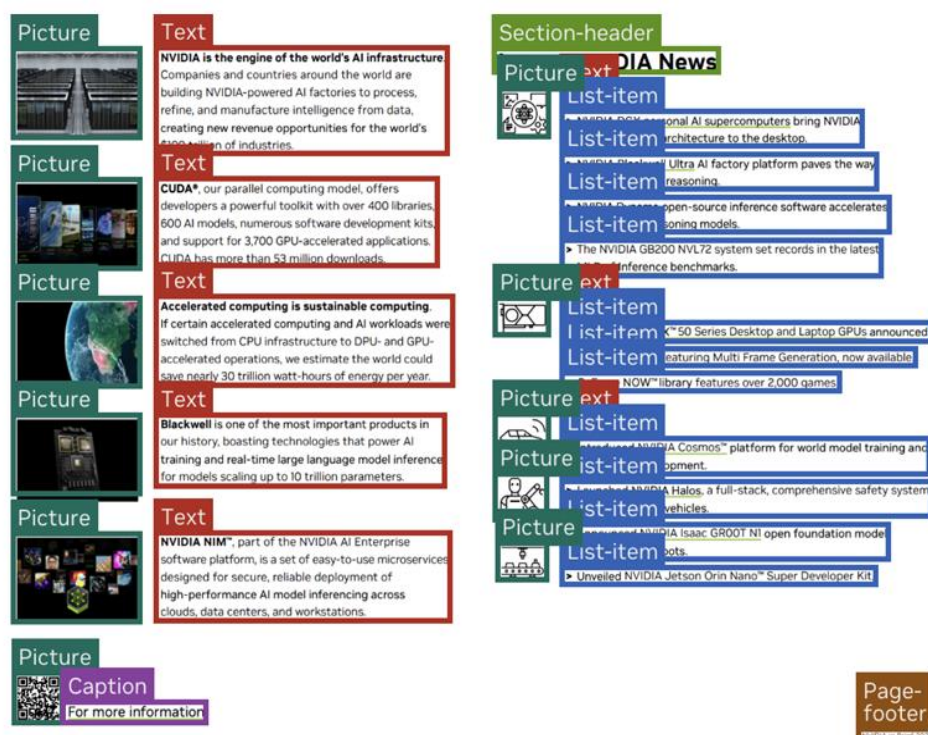


Figure 1 | Layout analysis: bounding box detection and prediction of semantic classes

内容摘要

Nvidia 提出 Nemotron-Parse-1.1，一款 885M 参数的编码器-解码器视觉语言模型，在通用 OCR、Markdown 排版、表格结构识别与语义区块分类等任务上全面超越 Kosmos-2.5、GOT 等主流方案。其创新包括：无位置编码的解码器设计，支持超长上下文；多令牌并行生成，推理速度提升 20%；自研 NVpdfTex 数据管线，实现字符级边界框与阅读顺序同步标注。模型与精简版 Nemotron-Parse-TC 已开源，单 H100 可达 5 页/秒，为大规模批处理与边缘部署提供高精度、低延迟的文档理解新基准。

相关链接

Huggingface: <https://huggingface.co/nvidia/NVIDIA-Nemotron-Parse-v1.1>

● 技术动态八

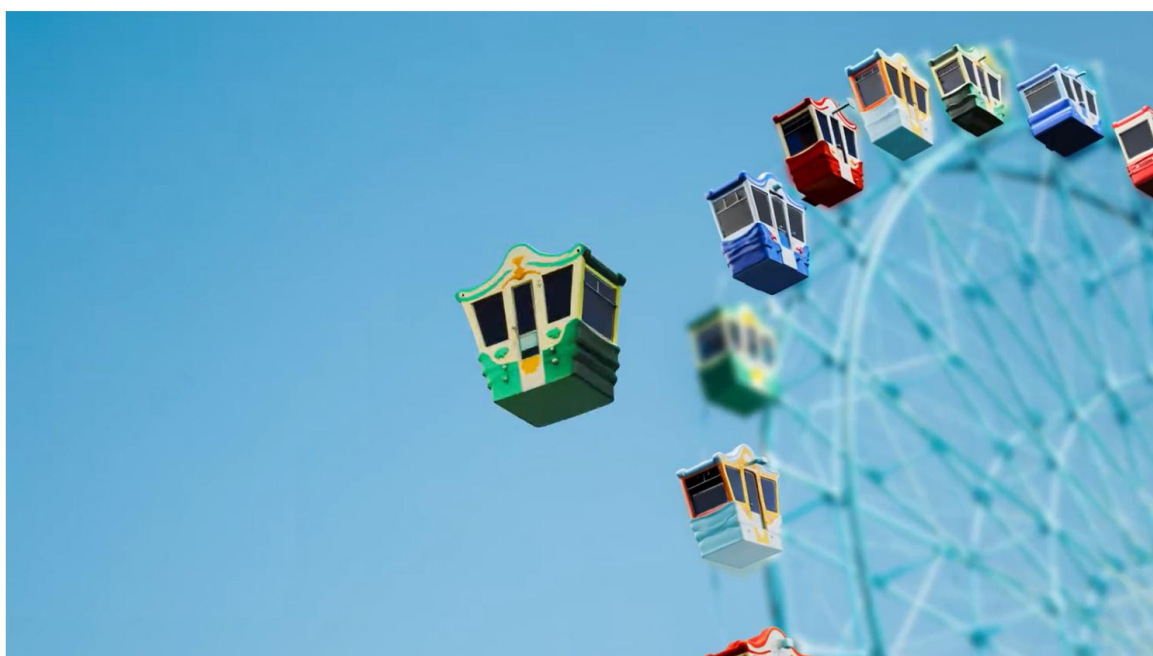
分割一切并不够，还要 3D 重建一切，SAM 3D 来了

日期

2025-11-20

Introducing SAM 3D: Powerful 3D Reconstruction for Physical World Images

November 19, 2025 • 10 minute read



内容摘要

Meta SAM 3D 博客官宣 Segment Anything Model 3D：把 2D 明星 SAM 升维到点云，一句话或一次点击即可在 VR/AR、机器人、建筑 BIM 等场景中零样本分割任意物体。新模型基于 Transformer，融合多视角 RGB-D 信息，训练数据覆盖 1200 万帧室内外交集，推理时无需体素化，直接输出实例级掩码，精度较此前最佳方法提升 18%，延迟降至 50 ms 内。Meta 同步开放跨平台 C++/Python SDK、PyTorch Hub 与 Oculus 示例，开发者可用 Quest 3 实时扫描客厅并指令“选沙发”即刻生成 3D 掩码，用于虚实遮挡、抓取规划或 AEC 标注。博客强调 SAM 3D 与 Llama 3、Ray-Ban Meta 眼镜协同，将推动“空间智能”生态，并承诺模型与数

据集均免费商用，加速下一代沉浸式 AI 应用落地。

相关链接

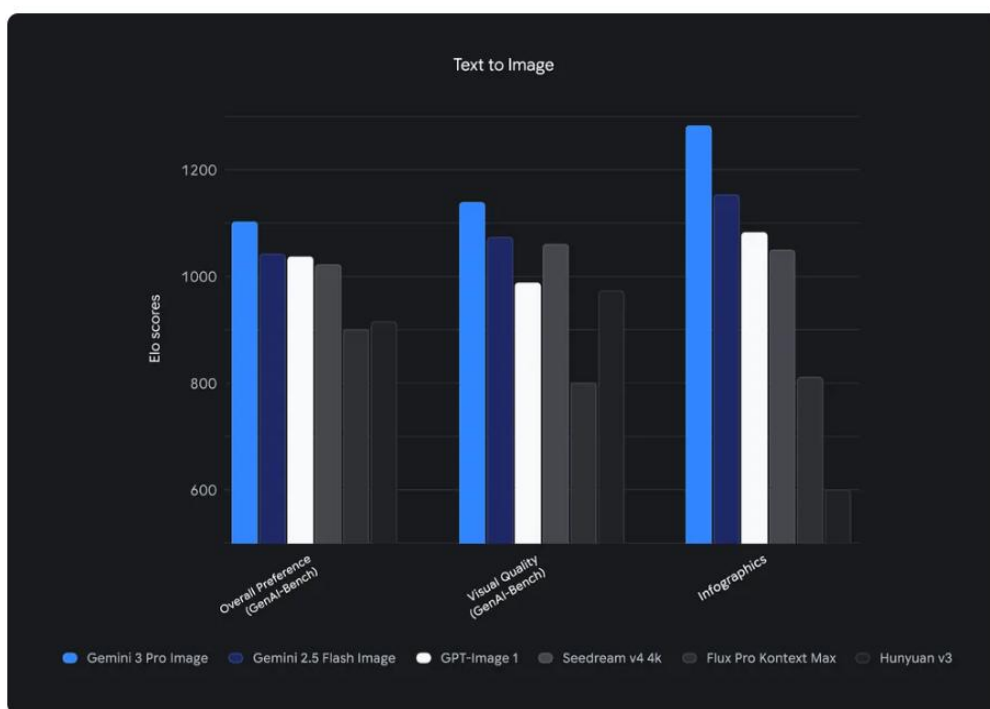
Blog: <https://ai.meta.com/blog/sam-3d/>

● 技术动态九

谷歌 Nano Banana Pro 炸了！硅谷 AI 半壁江山同框，网友：PS 已死

日期

2025-11-21



Gemini 3 Pro Image excels on Text to Image AI benchmarks.

内容摘要

Google 官方博客宣布推出 Gemini 3 Pro，面向开发者开放图像理解与生成 API。新模型支持高分辨率输入、多轮图文对话，原生输出 SVG、PNG、LaTeX，可一次性生成 10 张高质量图像并精准渲染文字、图表与代码。Gemini 3 Pro 在 MMLM、OCR、数学推理基准上刷新 SOTA，延迟降低 30%，成本下降 50%。谷歌

同步发布 Python SDK 与 Vertex AI 一键部署，开发者可用自然语言调用图像编辑、UI 原型、数据可视化等工具，无需 GPU 即可在 Colab 免费体验。博客展示实时生成可运行代码的仪表盘、游戏素材和科研示意图，强调安全过滤与版权检测已默认嵌入，标志着 Gemini 从“对话”走向“视觉创作”，开启云端多模态应用新范式。

相关链接

Blog: <https://blog.google/technology/developers/gemini-3-pro-image-developers/>

● 技术动态十

国内唯一！阿里千问斩获 NeurIPS 2025 最佳论文奖

日期

2025-11-27



Gated Attention for Large Language Models: Non-linearity, Sparsity, and Attention-Sink-Free

Zihan Qiu, Zekun Wang, Bo Zheng, Zeyu Huang, Kaiyue Wen, Songlin Yang, Rui Men, Le Yu, Fei Huang, Suozhi Huang, Dayiheng Liu, Jingren Zhou, Junyang Lin

内容摘要

NeurIPS 2025 最佳论文揭晓，阿里通义千问团队凭《注意力门控》研究从 5524 篇投稿中折桂，成为唯一获奖中国团队。论文首次在 1.7 B Dense 与 15 B MoE 上训练 3.5 T token，系统揭示：对注意力头输出施加轻量级门控，即可把首 token 注意力占比从 46.7 % 压至 4.8 %，最大激活值由 1053 降至 9，仅用

1 % 额外参数换来 MMLU +2、困惑度 -0.2，并显著提升低精度训练稳定性。该发现为“注意力池”与“巨量激活”两大顽疾提供通用解法，已集成至 Qwen3-Next 并全面开源；评委会预判其将成为新一代大模型标配，通义千问亦借此巩固全球开源生态领先地位。

相关链接

Blog: <https://blog.neurips.cc/2025/11/26/announcing-the-neurips-2025-best-paper-awards/>

RESOURCE SHARING

资源分享

● 资源分享一

《Machine-Learning-Systems》

类型：电子书

内容提要

这本《机器学习系统》教科书全面探讨了机器学习系统的工程实践，从基础理论到生产环境部署。内容分为系统基础、设计原则、性能工程、鲁棒部署和可信系统等核心部分，深入解析数据工程、模型训练、优化技术及安全伦理问题。书中通过嵌入式平台（如 Arduino、Raspberry Pi）的实验室练习强化实践，涵盖图像分类、目标检测等应用，并包含自检环节以巩固学习。作为开源资源，它旨在为学生、工程师及研究者提供构建可扩展、高效且负责任的 ML 系统所需的理论与实战知识。

资源链接

<https://github.com/KittenML/KittenTTS>

● 资源分享二

《电子书资源网站整理大全》

类型：网站

内容提要

这是一个 2025 年 11 月整理的电子书资源网站大全，收录了 25 个优质平台，覆盖中文综合资源（如鸠摩搜书、文渊阁）、外文及专业资源（如 LibGen、Project Gutenberg），提供 PDF、EPUB 等多种格式免费下载，支持学术、小说、古籍等多样需求。部分站点可能需要 VPN 访问，建议支持正版阅读。整理由@Berryxia_ai 完成，旨在为读者提供全面、便捷的电子书获取指南。

资源链接

<https://youmind.site/01e0UcErPxD0bG>

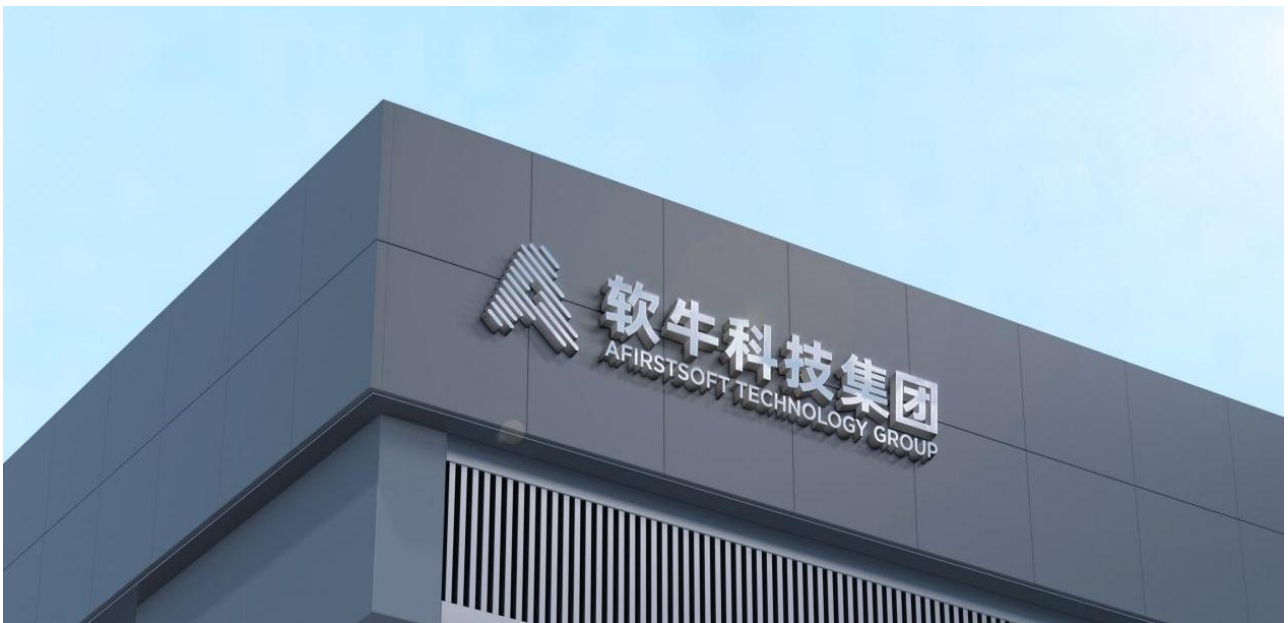


深圳软牛科技集团股份有限公司成立于 2014 年 6 月，作为国家级专精特新“小巨人”企业，始终以“让创意成为现实”为使命，专注于通过人工智能视觉大模型技术驱动全球数字创意生态。

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



2025 年第 11 期

总第 29 期

广东省图象图形学会

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
林泽柠	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

2025 年第 11 期

总第 29 期

— 2025.12.01

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，文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、科普宣传、科技成果鉴定