

Artificial Intelligence Newsletter

人工智能 前沿快报

2025 年第 08 期
总第 26 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. Discovering Cognitive Strategies with Tiny Tecurrent Neural Networks	1
二. A Survey of Self-Evolving Agents: On Path to Artificial Super Intelligence	4
三. Qwen-Image Technical Report	5
四. Why are LLMs' abilities emergent?	8
五. GLM-4.5: Agentic, Reasoning, and Coding (ARC) Foundation Models	9
六. R-Zero: Self-Evolving Reasoning LLM from Zero Data	12
七. A Comprehensive Survey of Self-Evolving AI Agents	14
八. Geometric-Mean Policy Optimization	17
九. From AI for Science to Agentic Science: A Survey on Autonomous Scientific Discovery	19
十. Deep Think With Confidence	21
十一. InternVL3.5: Advancing Open-Source Multimodal Models in Versatility, Reasoning, and Efficiency	24
十二. Exploring the role of large language models in the scientific method: from hypothesis to discovery	26
十三. Generalist Reward Models: Found Inside Large Language Models	28

技术动态

一. OpenAI 推出 gpt-oss 系列开源模型, Gpt-oss-120b 和 gpt-oss-20b 以推动开源推理模型领域的技术边界	31
二. 谷歌推出 Genie 3, 能够生成多样化交互环境的通用世界模型	32
三. OpenAI 推出 GPT-5, 让专业级智能触手可及	34
四. DeepSeek-V3.1 发布, 迈向 Agent 时代的第一步	35
五. 轻量级易开发, 8B 参数释放大实力! 科学多模态模型 Intern-S1-mini 开源	37
六. 语音模型第一次超越人类! OpenAI 再现 Her 时刻, 95 后华人研究员坐镇	38

资源分享

一. Kitten TTS	40
二. Canary 1B v2: Multitask Speech Transcription and Translation Model	40
三. Qwen3-235B-A22B-Instruct-2507	41

鸣谢支持

一. 深圳软牛科技集团股份有限公司	42
-------------------	----

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，促进学术交流和技術繁荣。本期摘选了近期全球人工智能领域的 13 篇最新学术动态、6 篇技术动态、以及 3 个资源推荐给广大读者，部分亮点动态包括：

- 阿里巴巴发布 Qwen-Image 图像生成模型技术报告（08 月 04 日）
- 谷歌推出生成多样化交互环境的通用世界模型 Genie 3（08 月 05 日）
- OpenAI 开源 GPT-oss，在 AIME2024/2025、Humanity Last Exam、GPQA 等数据集上性能超越 o3-mini（08 月 05 日）
- 智谱发布混合专家大模型 GLM-4.5、视觉推理大模型 GLM-4.5V，多个数据集上性能达到 SOTA（08 月 12 日）
- 自进化 AI 代理论文综述（08 月 12 日）
- 腾讯与华盛顿大学提出大模型自我进化推理训练范式 R-Zero（08 月 12 日）
- 基于 AI 智能体的自动化科学发现综述论文（08 月 18 日）
- DeepSeek-V3.1 正式发布，首次在同一模型内实现“思考/非思考”混合推理（08 月 21 日）
- 8B 参数科学多模态模型 Intern-S1-mini 开源（08 月 21 日）
- 上海 AI Lab 发布开源多模态模型 InternVL3.5（08 月 25 日）
- NVIDIA Canary-1B-v2、Qwen3-235B-A22B-Instruct-2507 模型资源（资源分享）

ACADEMIC TRENDS

学术动态

● 论文一

Discovering cognitive strategies with tiny recurrent neural networks

日期

2025-07-02

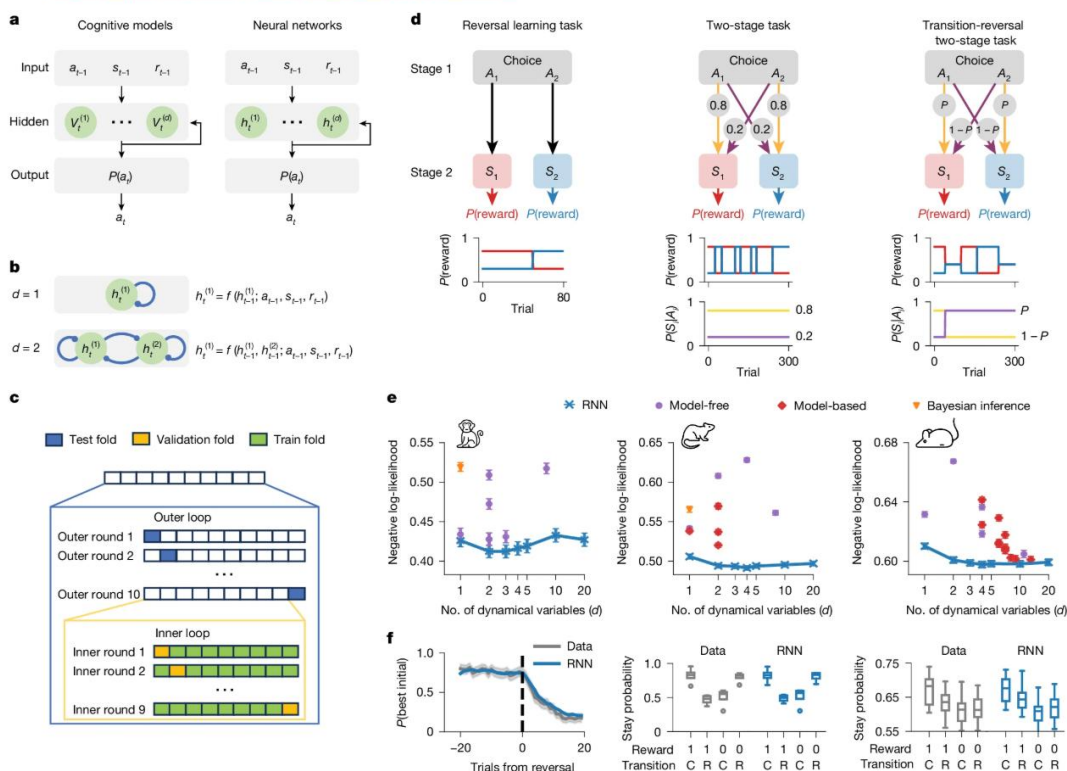
Article

Discovering cognitive strategies with tiny recurrent neural networks<https://doi.org/10.1038/s41586-025-09142-4> Li Ji-An¹, Marcus K. Benna¹ & Marcelo G. Mattar^{2,3}✉**内容摘要**

理解动物和人类如何通过经验学习做出适应性决策，是神经科学与心理学的核心目标。规范性建模框架，如贝叶斯推理与强化学习，揭示了适应性行为的部分原理，但其简洁性往往难以捕捉真实的生物行为，导致模型需要频繁手工调整，而这种操作容易受到研究者主观性的影响。本文提出了一种新的建模方法，利用循环神经网络（RNN）来自动发现支配生物决策的认知算法。结果表明，仅由 1–4 个单元组成的极小型 RNN，在六种经典奖励学习任务中对动物和人类个体选择进行预测时，不仅优于传统认知模型，甚至可与更大的神经网络相媲美。关键在于，这些小型网络可以通过动力系统的视角进行解释，使认知模型间具备统一的比较框架，并揭示选择行为的精细机制。本方法还能估计行为的维度性，并为理解元强化学习人工智能体所习得的算法提供见解。总体而言，我们提出了一种系统方法来发现可解释的认知策略，为揭示神经机制以及研究正常与异常的认知过程奠定了基础。

Fig. 1: Model overview and performance on animal tasks.

From: [Discovering cognitive strategies with tiny recurrent neural networks](#)



主要亮点

- **极小型 RNN 建模**: 仅需 1–4 个单元的小型循环网络, 就能优于经典认知模型, 并媲美大规模网络的预测能力。
- **动力系统解释**: 通过相图与动力学分析揭示新机制, 如状态依赖学习率、奖励诱发的中立效应和非对称选择偏置。
- **行为维度估计**: 将 RNN 单元数作为行为复杂度的量化指标, 发现多种任务中的动物与人类决策行为可用低维动力学刻画。
- **跨物种普适性**: 在猴子、鼠类与人类的八个数据集、六类奖励学习任务中均表现稳定优势, 显示方法的广泛适用性。

- **AI-生物对照：** 将 RNN 框架应用于元强化学习智能体，揭示人工智能与生物体在认知策略上的共性与差异。

相关链接

Paper: <https://www.nature.com/articles/s41586-025-09142-4>

● 论文二

A Survey of Self-Evolving Agents: On Path to Artificial Super Intelligence

日期

2025-08-01

A SURVEY OF SELF-EVOLVING AGENTS: ON PATH TO ARTIFICIAL SUPER INTELLIGENCE

Huan-ang Gao^{γ†}, Jiayi Geng^{α†}, Wenyue Hua^{ε†}, Mengkang Hu^{ω†}, Xinzhe Juan^{σμ†}, Hongzhang Liu^{ξ†},
Shilong Liu^{α†}, Jiahao Qiu^{αδ†}, Xuan Qi^{γ†}, Yiran Wu^{ρ†}, Hongru Wang^{τ^{α†}⊠}, Han Xiao^{τ†},
Yuhang Zhou^{λ†}, Shaokun Zhang^{ρ†}, Jiayi Zhang^π, Jinyu Xiang, Yixiong Fang^θ, Qiwen Zhao^ζ,
Dongrui Liu^σ, Qihan Ren^σ, Cheng Qian^β, Zhenhailong Wang^β, Minda Hu^τ,
Huazheng Wang^η, Qingyun Wu^ρ, Heng Ji^β, Mengdi Wang^{αδ⊠}

^αPrinceton University, ^δPrinceton AI Lab, ^γTsinghua University, ^θCarnegie Mellon University,
^ξUniversity of Sydney, ^σShanghai Jiao Tong University, ^ρPennsylvania State University, ^μUniversity of Michigan,
^ηOregon State University, ^τThe Chinese University of Hong Kong, ^λFudan University,
^πThe Hong Kong University of Science and Technology (Guangzhou), ^ωThe University of Hong Kong,
^εUniversity of California, Santa Barbara, ^ζUniversity of California San Diego,
^βUniversity of Illinois Urbana-Champaign,

Github Repo: <https://github.com/CharlesQ9/Self-Evolving-Agents>

[†]Equal contribution and the order is determined alphabetically, [⊠]Corresponding Author

内容摘要

大型语言模型（LLMs）已展现出强大的能力，但本质上仍是静态的，无法适应新任务、扩展知识领域或动态交互环境。随着 LLMs 越来越多地部署在开放性、交互性环境中，这种静态特性已成为关键瓶颈，迫切需要能够实时自适应推理、行动和进化的智能体。从扩展静态模型到开发自进化智能体的范式转变正引发学术界与产业界的广泛关注，并推动持续学习与适应性方法的研究。本综述首次对自进化智能体进行了系统而全面的回顾，围绕三个基础维度——进化什么、何时进化以及如何进化——进行组织。我们考察了智能体组件（例如模型、记忆、工具、架构）中的进化机制，根据阶段（例如测试内、测试间）对适应方法进行分类，并分析了指导进化适应的算法和架构设计（例如标量奖励、文本反馈、单智能体和多智能体系统）。此外，我们分析了为自进化智能体量身定制的评估指标和基准，突出了在编码、教育和医疗等领域的应用，并确定了在安全性、可扩展性和协同进化动态方面的关键挑战和研究方向。通过提供一个理解和设计自进化智能体的结构化框架，本综述为推进自适应智能体系统在研究和实际部署中的发展奠定了路线图，最终为实现人工超智能（Artificial Super Intelligence）提供指引，在这一阶段，智能体能够自主进化，在广泛任务中表现出或超越人类水平的智能。

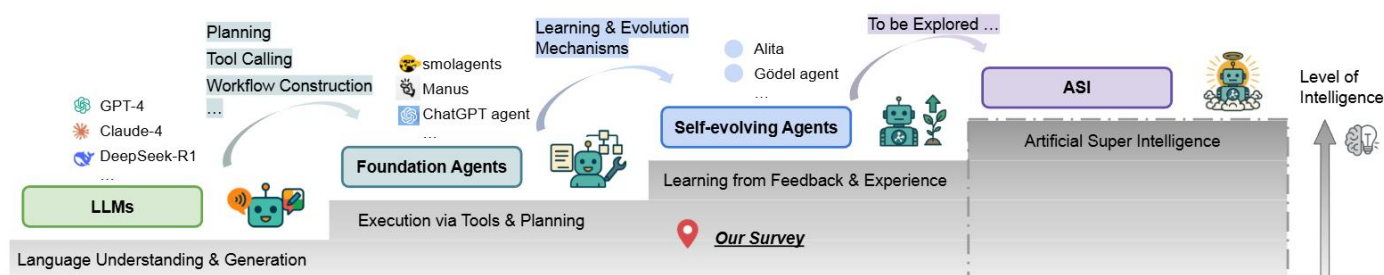


Figure 1: A conceptual trajectory illustrating the progression from large language models (LLMs) to foundation agents, advancing to self-evolving agents—our focus, and ultimately toward hypothetical Artificial Super Intelligence (ASI). Along this path, intelligence and adaptivity increase, marking a shift toward more autonomous and agentic AI systems.

主要亮点

- **首个系统综述：**本研究首次系统性地梳理了自演化智能体（self-evolving agents），填补了该领域缺乏综述的空白。

- **三维分析框架**: 提出“演化什么、何时演化、如何演化”三大核心维度，对模型、记忆、提示词、工具与架构的演化机制进行了全面归纳。
- **方法论总结**: 将现有方法分为奖励驱动、模仿与示范学习、进化式方法，并结合在线/离线、on/off-policy 等跨切维度进行系统比较。
- **评测与应用**: 提出适用于自演化智能体的新型评估指标与基准，涵盖静态评测、短期适应与长期终身学习，并展示了在编程、教育、医疗等领域的实践潜力。

相关链接

Paper: <https://arxiv.org/pdf/2507.21046v3>

github: <https://github.com/CharlesQ9/Self-Evolving-Agents>

● 论文三

Qwen-Image Technical Report

日期

2025-08-04



August 5, 2025

Qwen-Image Technical Report

Qwen Team



<https://qwen.ai>



<https://huggingface.co/Qwen/Qwen-Image>



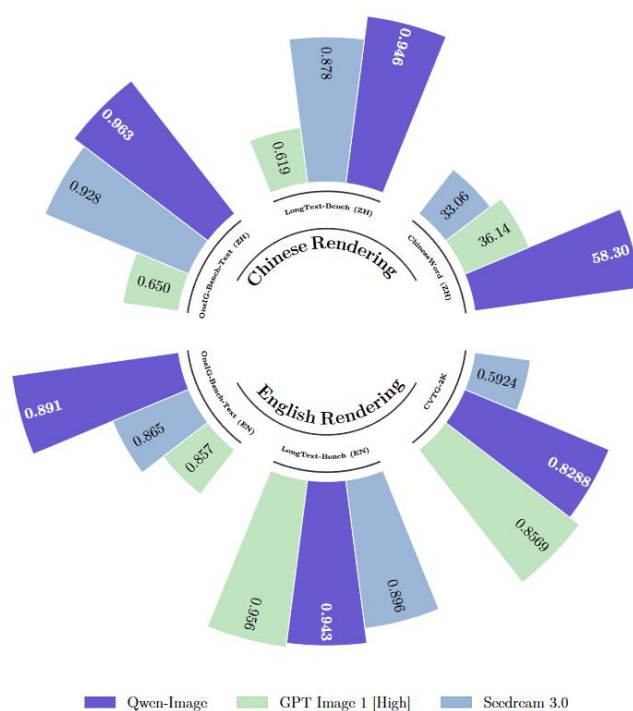
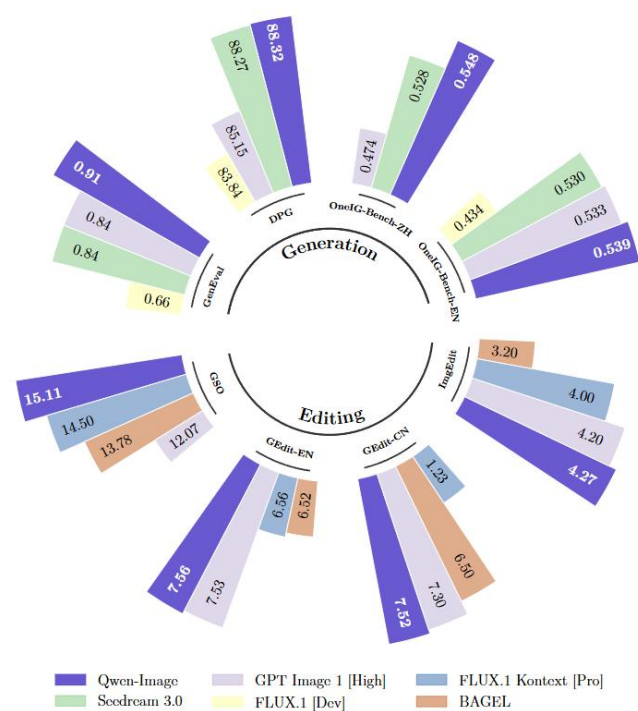
<https://modelscope.cn/models/Qwen/Qwen-Image>



<https://github.com/QwenLM/Qwen-Image>

内容摘要

Qwen-Image 是 Qwen 系列中的一个图像生成基础模型，在复杂文本渲染和精确图像编辑方面取得了显著进步。为了应对复杂文本渲染的挑战，我们设计了一个全面的数据管道，包括大规模数据收集、过滤、标注、合成和平衡。此外，我们采用了一种渐进式训练策略，从非文本渲染开始，逐步从简单的文本输入演变为复杂的文本输入，并逐渐扩展到段落级别的描述。这种课程式学习方法显著增强了模型的原生文本渲染能力。因此，Qwen-Image 不仅在英语等字母语言中表现优异，还在中文等更具挑战性的表意语言上取得了显著进步。为了增强图像编辑的一致性，我们引入了一种改进的多任务训练范式，该范式不仅包括传统的文本到图像（T2I）和文本图像到图像（TI2I）任务，还包括图像到图像（I2I）重建，有效地对齐了 Qwen2.5-VL 和 MMDiT 之间的潜在表示。此外，我们将原始图像分别输入 Qwen2.5-VL 和 VAE 编码器，以分别获取语义表示和重构表示。这种双重编码机制使编辑模块能够在保持语义一致性和维持视觉保真度之间取得平衡。Qwen-Image 在多个基准测试中实现了最先进的性能，展示了其在图像生成和编辑方面的强大能力。



(a) Image generation and editing benchmarks

(b) Text rendering benchmarks

Figure 1: Qwen-Image exhibits strong general capabilities in both image generation and editing, while demonstrating exceptional capability in text rendering, especially Chinese.

主要亮点

- **文本渲染进步**: 通过课程学习和合成数据, 大幅提升复杂文本生成能力, 特别是在中文渲染上取得领先。
- **一致性图像编辑**: 提出双通道语义+重构机制, 在编辑中兼顾语义一致性与视觉保真度。
- **高质量数据管线**: 构建大规模多类别数据集, 设计七阶段过滤与三类文本合成策略, 确保多样性与质量。
- **架构创新**: 基于 MMDiT 的扩散架构, 并引入 MSRoPE 位置编码, 优化跨模态对齐。
- **跨基准领先**: 在多项生成与编辑基准上取得 SOTA, 并在 AI Arena 人类评价中超过 GPT Image 1、FLUX 等闭源模型。

相关链接

Paper: <https://arxiv.org/pdf/2508.02324>

Code: <https://github.com/QwenLM/Qwen-Image>

Blog: <https://huggingface.co/Qwen/Qwen-Image>

● 论文四

Why are LLMs' abilities emergent?

日期

2025-8-6

Why are LLMs' abilities emergent?

Vladimír Havlík¹

内容摘要

大型语言模型（LLMs）在生成任务上表现卓越，其能力来源的根本性问题引起我们的思考：许多能力似乎在没有显式训练的情况下“涌现”出来。本文通过理论分析与实证观察，探讨了深度神经网络（DNNs）的涌现特性，并指出这种“创造而未完全理解”的现象是当代人工智能发展的核心困境。与符号计算不同，神经网络依赖非线性与随机过程，使得模型的宏观行为无法直接从微观神经元活动推导。通过分析模型的规模法则、理解现象和相变，本研究表明：LLM 的涌现能力源于复杂非线性系统的动力学，而不仅仅是参数规模的扩展。研究揭示，当前关于指标、预训练损失阈值和情境学习的争论，都未能触及 DNNs 涌现的形而上学本质。主张应把 DNN 看作一种复杂动力系统，其能力类似于物理、化学和生物中的普遍涌现现象。最终，作者认为理解 LLM 能力需要转向研究内部动态转化过程，而不仅通过外在现象来观察。

主要亮点

- **问题定位：** 系统性探讨 LLM 能力“涌现”现象，回应“为何能力会在未显式训练时突然出现”的核心疑问。
- **方法论视角：** 对比符号计算与神经计算，指出 DNN 的非线性与随机性决定了宏观行为不可简化为微观神经元的运算。
- **关键机制：** 分析尺度定律、grokking、能力相变等现象，揭示能力涌现来自复杂系统的动力学过程。
- **跨学科类比：** 将 LLM 的涌现与物理、化学、生物中的复杂系统现象进行对比，主张建立统一的“普遍涌现”框架。

相关链接

Paper:<https://arxiv.org/pdf/2508.04401>

● 论文五

GLM-4.5: Agentic, Reasoning, and Coding (ARC) Foundation Models

日期

2025-8-12

GLM-4.5: Agentic, Reasoning, and Coding (ARC) Foundation Models

GLM-4.5 Team

Zhipu AI & Tsinghua University

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

Abstract

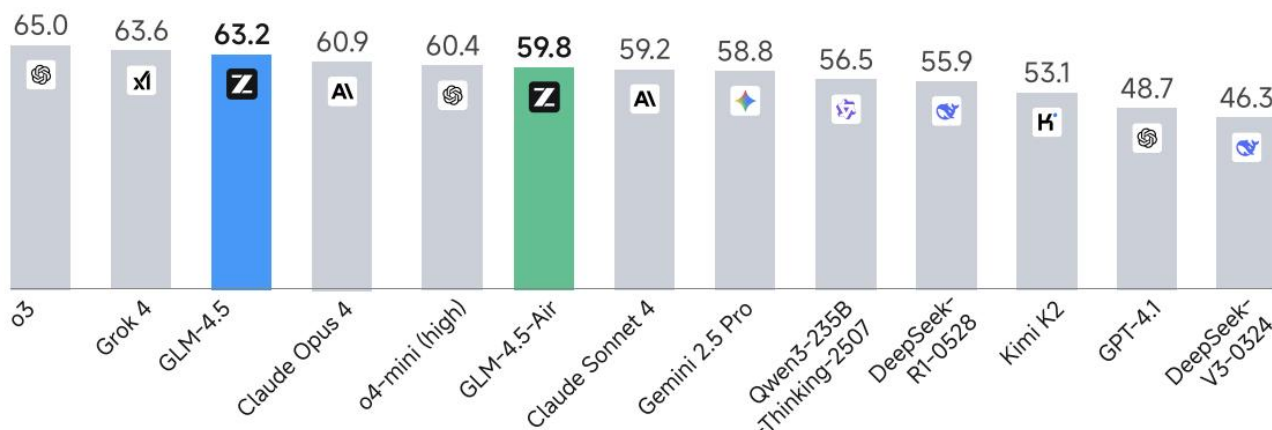
内容摘要

智谱华章近期推出了其开源混合专家 (MoE) 大型语言模型 GLM-4.5, 该模型总参数达 3550 亿, 激活参数为 320 亿, 并支持混合推理模式以兼顾思考与直接响应。经过 23 万亿 token 的多阶段训练和全面的后期优化, GLM-4.5 在智能体、推理和编码 (ARC) 任务上表现卓越, 尽管参数量远少于多个竞争对手, 其在所有参评模型中仍位列综合第三, 并在智能体基准测试中排名第二。为推动相关研究, 智谱华章同步发布了 GLM-4.5 及其小型版本 GLM-4.5-Air, 两款模型均采纳 MoE 架构以提升计算效率, 并在全球开源及闭源模型中均占据领先地位。

LLM Performance Evaluation: Agentic, Reasoning, and Coding Benchmarks



12 benchmarks: MMLU-Pro, AIME 24, MATH-500, SciCode, GPQA, HLE, LCB (2407-2501), SWE-Bench Verified, Terminal-Bench, TAU-Bench, BFCL V3, BrowseComp



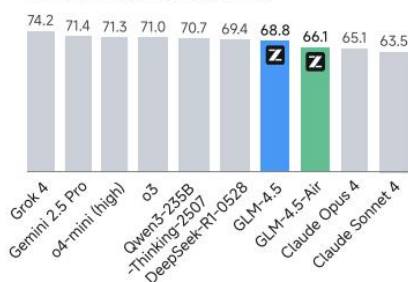
Agentic

Agentic Benchmarks: TAU-Bench, BFCL V3 (Full), BrowseComp



Reasoning

Reasoning Benchmarks: MMLU-Pro, AIME 24, MATH 500, SciCode, GPQA, HLE, LCB (2407-2501)



Coding

Coding Benchmarks: SWE-Bench Verified, Terminal-Bench

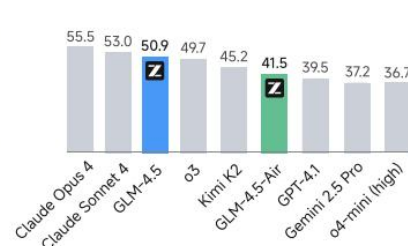


Figure 1: Average performance on agentic, reasoning, and coding (ARC) benchmarks. Overall, GLM-4.5 achieves a rank of 3rd, with GLM-4.5-Air following at rank 6th. The models listed are evaluated as of July 28, 2025.

主要亮点

- **卓越且高效的 ARC 能力：** GLM-4.5 在智能体、推理和编码三大核心任务上均展现出强大的性能。更重要的是，它以远少于竞争对手的参数量，在所有评估模型中综合排名第三，并在智能体基准测试中位居第二，体现了其出色的参数效率。这表明 GLM-4.5 不仅能力强，而且资源消耗更低。
- **创新的混合专家架构与混合推理方法：** GLM-4.5 是智谱华章首个采用 MoE 架构的模型，其总参数为 3550 亿，激活参数为 320 亿。这种架构显著提高了训练和推理的计算效率。此外，它引入了一种独特的混合推理方法，支持“思考模式”用于复杂推理和智能体任务，以及“直接响应模式”用于即时回复。这

种灵活的机制使得模型能够根据任务需求，在深度思考和快速响应之间切换，提高了解决问题的通用性和效率。

- **多阶段精细化训练与开放生态：** GLM-4.5 经过 23 万亿 token 的多阶段预训练，并结合专家模型迭代和强化学习进行全面的后期训练，使其具备了统一各类复杂任务的能力，尤其是在代码和科学推理方面进行了深入优化。同时，智谱华章慷慨地开源了 GLM-4.5 及其紧凑版本 GLM-4.5-Air 的模型权重和评估工具包，旨在推动大语言模型在推理和智能体 AI 系统应用及研究领域的发展，构建开放的 AI 生态。

相关链接

Paper: <https://arxiv.org/pdf/2508.06471>

Github: <https://github.com/zai-org/GLM-4.5>

● 论文六

AR-Zero: Self-Evolving Reasoning LLM from Zero Data

日期

2025-08-12

R-Zero: Self-Evolving Reasoning LLM from Zero Data

Chengsong Huang^{1,2+}, Wenhao Yu¹⁺, Xiaoyang Wang¹, Hongming Zhang¹, Zongxia Li^{1,3},
Ruosen Li^{1,4}, Jiaxin Huang², Haitao Mi¹, Dong Yu¹

¹Tencent AI Seattle Lab, ²Washington University in St. Louis,

³University of Maryland, College Park, ⁴The University of Texas at Dallas

† Core contributors

chengsong@wustl.edu; wenhaoyu@global.tencent.com

内容摘要

R-Zero 提出了一种创新的、完全自主的框架，旨在无需任何人类标注数据的情况下，通过自我演化显著提升大型语言模型（LLM）的推理能力。该框架的核心是一个协同演化循环：它从一个基础 LLM 启动，并实例化出两个相互独立的模型——“挑战者”（Challenger）和“求解器”（Solver）。挑战者被激励生成那些恰好处于求解器当前能力边缘的难题，而求解器则通过成功解决这些由挑战者提出的、难度逐渐增高的任务来提升自身。这种独特的互动机制生成了一个有针对性的、自我改进的训练课程。实验证明，R-Zero 能够显著增强不同基础 LLM 在数学和通用领域推理基准上的性能，例如使 Qwen3-4B-Base 在数学推理基准上平均提升了 +6.49 分，在通用领域推理基准上提升了 +7.54 分。

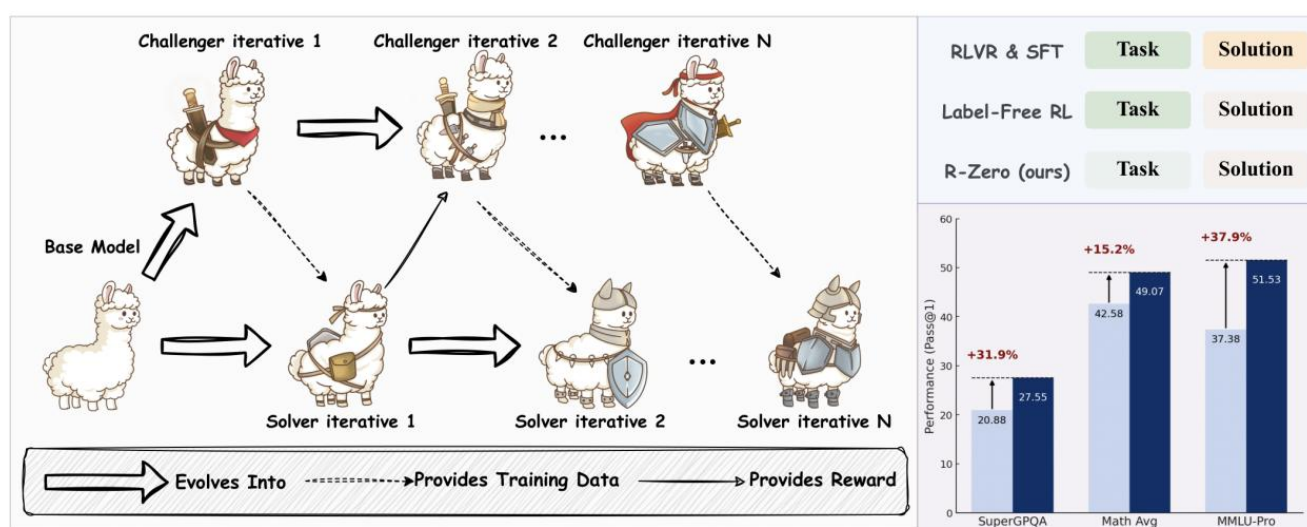


Figure 1: (Left): *R-Zero* employs a co-evolutionary loop between Challenger and Solver. (Right): *R-Zero* achieves strong benchmark gains without any pre-existing tasks or human labels.

主要亮点

- **完全自主的自演化框架：**R-Zero 解决了现有自演化 LLM 严重依赖大量人工标注数据的问题。它创新性地通过两个独立但协同演化的 LLM (挑战者和求解器) 来从零开始生成自身的训练数据和任务。挑战者生成难度适中的问题，求解器解答并反馈，形成闭环，实现了超越人类数据限制的自我学习路径。

- **对抗性协同演化机制：**R-Zero 引入了一种类似于生成对抗网络（GAN）的师生模式。挑战者的奖励机制是生成对当前求解器而言“刚好有挑战性”的问题，即求解器答对率约为 50% 的问题，以最大化学习效率。求解器则通过解决这些挑战者不断提出的难题来获得奖励并提升自身能力。这种动态平衡确保了训练课程的持续进步性和针对性，避免了模型在过易或过难任务上浪费学习资源。
- **显著的推理能力提升与泛化：**R-Zero 在数学推理基准上展示了强大的效果，能够持续、迭代地提升不同规模和架构的基础 LLM 的性能。例如，Qwen3-4B-Base 模型的数学推理平均分在三次迭代后提高了 +6.49 分。更重要的是，通过数学任务训练获得的推理能力能够泛化到通用领域，如 MMLU-Pro 和 SuperGPQA 等基准测试也展现出显著的性能提升，证明其增强的是模型底层的推理能力而非仅限于特定领域知识。

相关链接

Paper:<https://arxiv.org/pdf/2508.06471>

Github:<https://github.com/Chengsong-Huang/R-Zero>.

● 论文七

A Comprehensive Survey of Self-Evolving AI Agents

日期

2025-08-12

A Comprehensive Survey of Self-Evolving AI Agents

A New Paradigm Bridging Foundation Models and Lifelong Agentic Systems

Jinyuan Fang^{*1}, Yanwen Peng^{*2}, Xi Zhang^{*1}, Yingxu Wang^{*3}, Xinhao Yi¹, Guibin Zhang⁴, Yi Xu⁵, Bin Wu⁶, Siwei Liu⁷, Zihao Li¹, Zhaochun Ren⁸, Nikos Aletras², Xi Wang², Han Zhou⁵, Zaiqiao Meng¹✉

¹University of Glasgow, ²University of Sheffield, ³Mohamed bin Zayed University of Artificial Intelligence, ⁴National University of Singapore, ⁵University of Cambridge, ⁶University College London, ⁷University of Aberdeen, ⁸Leiden University

*Equal Contributor, ✉Corresponding Author

内容摘要

本综述全面探讨了自进化 AI 代理这一新兴范式,旨在通过自动化优化弥合基础模型的静态性和终身代理系统的适应性需求,为此提出了由“系统输入、代理系统、环境、优化器”组成的统一概念框架,并描绘了 LLM-centric 学习从静态预训练到多代理自进化的四大演进阶段。本文系统性地审视了针对代理系统各组件、多代理协作及特定领域(如生物医学、编程、金融法律)的广泛优化技术,并通过“忍耐、卓越、进化”的“自进化 AI 代理三定律”指导其安全与负责任的发展,最终指明了未来研究的挑战与前景,为构建更具适应性、自主性和韧性的 AI 代理系统提供了全面的理论与实践指南。

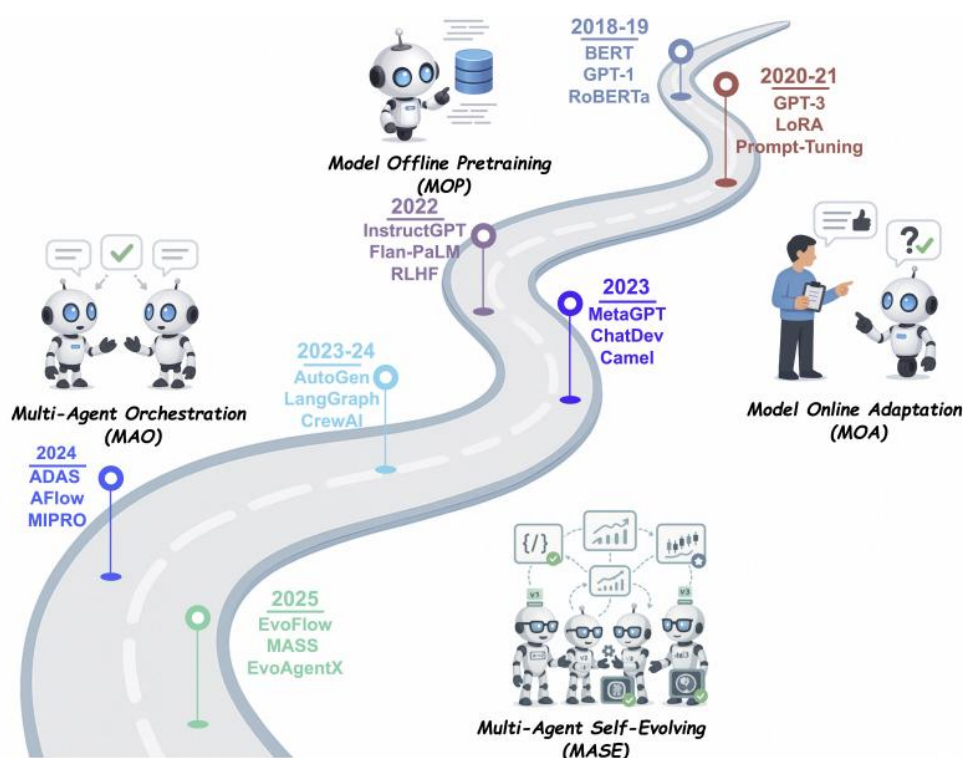


Figure1 LLM-centric learning is evolving from learning purely from static data, to interacting with dynamic environments, and ultimately towards lifelong learning through multi-agent collaboration and self-evolution.

主要亮点

- **提出并阐述了“自进化 AI 代理三定律”**：本综述创新性地借鉴了阿西莫夫的机器人三定律，为自进化 AI 代理提出了“忍耐 (Enudre) - 安全适应”、“卓越 (Excel) - 性能保持”和“进化 (Evolve) - 自主优化”这三条分层原则。这为未来设计和评估自适应、自主且值得信赖的终身代理系统提供了明确的道德和功能指导框架，强调了安全和性能是自主进化的前提，而非事后考量。
- **构建了统一的概念框架和 MOP→MASE 范式演化路线图**：本综述引入了一个全面的概念框架，将自进化代理的整个优化过程抽象为“系统输入、代理系统、环境、优化器”四大核心组件，并详细分析了它们之间的迭代反馈循环。在此基础上，论文进一步描绘了 LLM 中心化学习从“模型离线预训练 (MOP)”到“模型在线适应 (MOA)”，再到“多代理编排 (MAO)”，最终迈向“多代理自进化 (MASE)”的清晰演化路径。这一框架和路线图为理解、分类和系统化研究自进化代理提供了强大的理论基础和宏观视角。

- **全面系统地综述了多维度优化技术与挑战：**本综述不仅细致地分类和总结了针对单一代理（LLM 行为、提示词、记忆、工具）和多代理系统（提示词、拓扑结构、LLM 骨干）的多种进化及优化技术，还特别关注了生物医学、编程、金融和法律等领域特定的优化策略。这种详细且分层化的技术回顾，结合对评估方法、安全考量、现有挑战和未来方向的深入讨论，为研究人员和实践者提供了一份极为全面且具有前瞻性的资源指南，有助于推动更具适应性、自主性和韧性的 AI 代理系统发展。

相关链接

Paper <https://arxiv.org/pdf/2508.07407>

GitHub: <https://github.com/EvoAgentX/Awesome-Self-Evolving-Agents>

● 论文八

Geometric-Mean Policy Optimization

日期

2025-08-13

Geometric-Mean Policy Optimization

Yuzhong Zhao^{1*} Yue Liu^{1*} Junpeng Liu² Jingye Chen³ Xun Wu⁴ Yaru Hao⁴
Tengchao Lv⁴ Shaohan Huang⁴ Lei Cui⁴ Qixiang Ye¹ Fang Wan¹ Furu Wei⁴
¹UCAS ²CUHK ³HKUST ⁴Microsoft Research
<https://aka.ms/GeneralAI>

内容摘要

本文提出了针对大语言模型推理能力提升的新型强化学习方法 GMPO。该方法以几何平均替代传统算术平均，显著降低了因极端重要性采样比值导致的训练不稳定现象，有效缓解了策略更新过激和梯度波动的问题。理论与实验结果均表明，GMPO 不仅在数学推理、代码生成及多模态任务中表现出更高的训练稳定性，还能扩大探索空间，提升模型泛化能力。与主流 GRPO 算法相比，GMPO 在多个数学和几何推理基准上实现了明显性能提升，平均准确率提升 4.1%。该工作为大模型强化学习的稳定化与高效化提供了新范式和理论依据，对未来高可靠智能体和多模态推理系统的研发具有重要推动作用。

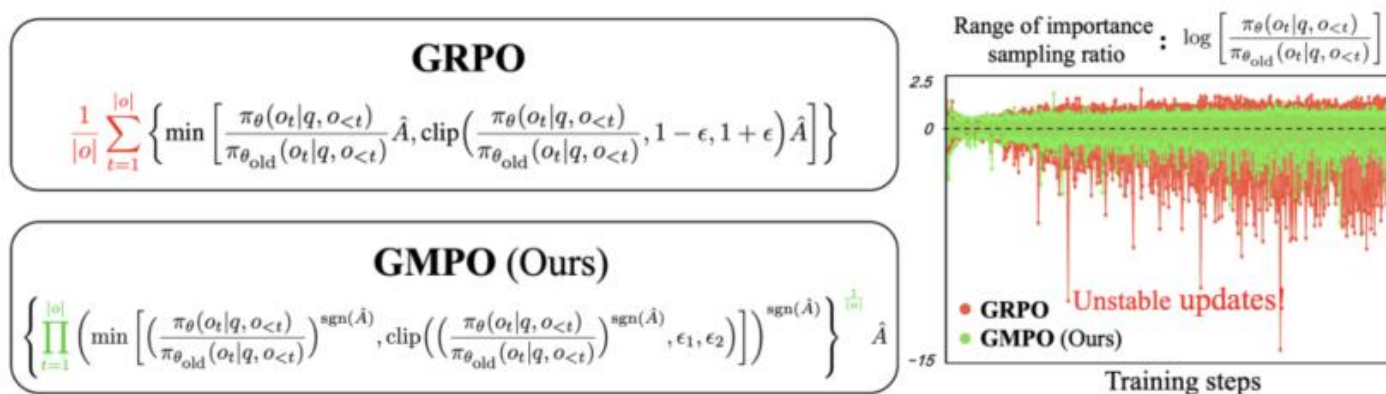


Figure 1: Comparison between vanilla Group Relative Policy Optimization (GRPO) and our proposed GMPO. GRPO optimizes the arithmetic mean of token-level rewards, while GMPO uses the geometric mean (left). For GRPO, extreme importance sampling ratios emerge frequently during training, which implies unstable policy updates. In contrast, GMPO results in more stable importance sampling ratios with less extreme values and narrower range (right).

主要亮点

- **创新性优化目标：**提出几何平均策略优化（GMPO），用几何均值替代传统算术均值，有效抑制极端重要性采样比带来的训练不稳定，大幅提升策略优化的稳健性。
- **理论与实验双重验证：**系统分析了 GMPO 在梯度平稳性、KL 散度和熵等指标上的优势，理论证明与大量实验均显示其在训练稳定性和探索能力上优于现有方法。

- **全面性能提升**: 在多项数学与多模态推理基准测试中, GMPO 显著超越主流 GRPO 算法, 平均提升 4.1%, 为大模型推理能力的可靠进化奠定了坚实基础。

相关链接

Paper: <https://arxiv.org/pdf/2507.20673>

Code: <https://github.com/callsys/GMPO>

● 论文九

From AI for Science to Agentic Science: A Survey on Autonomous Scientific Discovery

日期

2025-08-18

From *AI for Science* to *Agentic Science*: A Survey on Autonomous Scientific Discovery

Jiaqi Wei^{1,2*}, Yuejin Yang^{1,3*}, Xiang Zhang^{4*}, Yuhan Chen^{1,5*}, Xiang Zhuang^{1,2*},
Zhangyang Gao^{1*}, Dongzhan Zhou^{1*}, Guangshuai Wang^{1,3}, Zhiqiang Gao¹, Juntai Cao⁴,
Zijie Qiu^{1,3}, Xuming He^{1,2}, Qiang Zhang², Chenyu You⁶, Shuangjia Zheng^{7,8}, Ning Ding⁹,
Wanli Ouyang^{1,10}, Nanqing Dong¹, Yu Cheng^{1,10}, Siqi Sun^{1,3†}, Lei Bai^{1†}, Bowen Zhou^{1,9†}

¹ Shanghai Artificial Intelligence Laboratory, ² Zhejiang University, ³ Fudan University,

⁴ University of British Columbia, ⁵ Tongji University, ⁶ Stony Brook University,

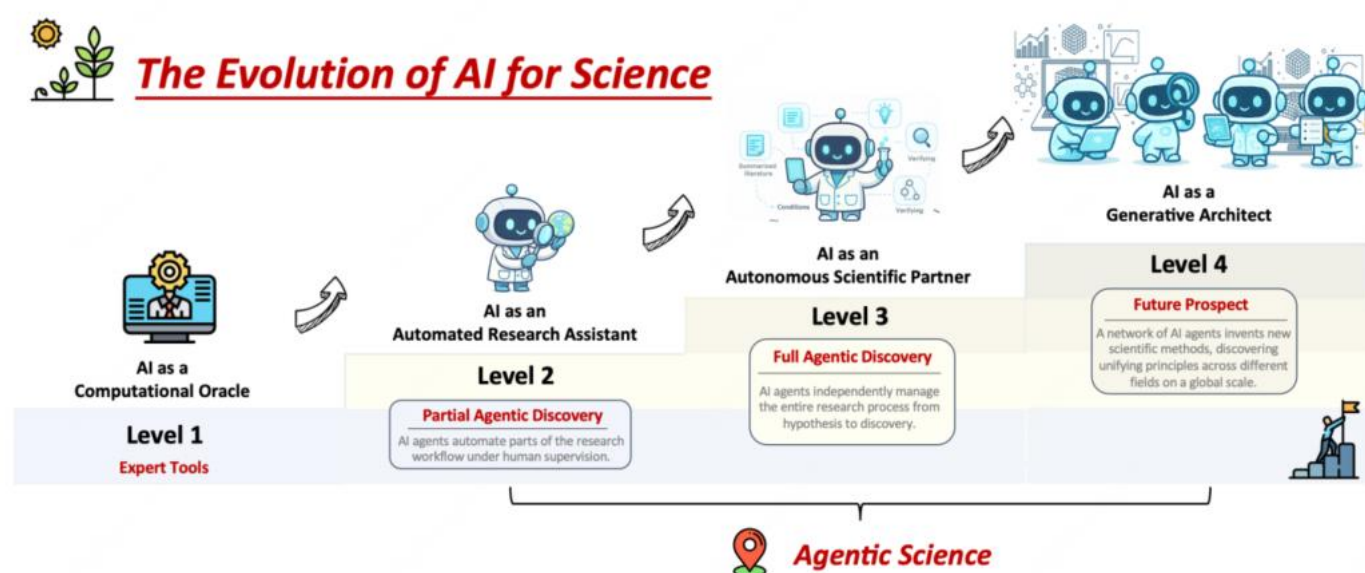
⁷ Shanghai Jiaotong University, ⁸ Lingang Laboratory,

⁹ Tsinghua University, ¹⁰ The Chinese University of Hong Kong

内容摘要

本文系统梳理了人工智能 (AI) 在科学发现领域从工具化智能到具备自主科研能力的 “Agentic Science” 阶段的演进路径。文章提出, 得益于大模型、跨模态系统及集成科研平台的发展, AI 已从仅能执行部分辅助任

务的助手，跃升为能自主生成假设、规划实验、执行流程、分析数据并迭代优化的科学代理。这一新范式不仅覆盖生命科学、化学、材料科学和物理等自然科学主流领域，还构建了统一的理论框架，从过程导向、自治水平及机制视角融合分析，归纳出科学代理的五大能力及四阶段动态科研流程。作者详尽回顾了各领域前沿应用案例、技术突破及实际验证成果，深入剖析当前 AI 自主科研面临的可靠性、可复现性、科学解释性与伦理等核心挑战，并展望了未来 AI 科学家在自主创新、跨学科融合、全球协作及“诺贝尔-图灵测试”等方向的变革性潜力。该文不仅为 AI 驱动科学发现提供了全景式的理论与方法论基础，更激发了对“智能科学”时代科学本质与未来边界的深度思考。



主要亮点

- **创新性提出“自进化 AI 代理三定律”**：论文借鉴阿西莫夫机器人三定律，首次为自进化 AI 代理提出了“忍耐—安全适应”、“卓越—性能保持”、“进化—自主优化”三层原则，明确将安全和性能作为代理自主进化的先决条件。这一分层原则为未来自适应、终身学习型 AI 系统的设计与评估提供了道德与功能的双重指导框架。

- **构建统一的概念框架与演化路线图：**文章引入了涵盖“系统输入、代理系统、环境、优化器”四大核心组件的全面概念框架，并系统梳理了自进化 AI 代理从模型离线预训练（MOP）、模型在线适应（MOA）、多代理编排（MAO）到多代理自进化（MASE）的完整演化路径，为该领域的理论建模与方法分类提供了宏观视角与理论基础。
- **系统综述多维度优化技术与挑战：**论文详细分类和总结了单一代理及多代理系统中的优化与进化技术，覆盖行为建模、提示词、记忆机制、工具集成及系统拓扑等关键环节，并结合生物医学、金融、法律等领域的实际应用，深入讨论了安全性、评估标准及未来发展方向，为推动更具适应性和可信赖性的 AI 代理系统研究奠定了坚实基础。

相关链接

Paper: <https://arxiv.org/pdf/2508.14111>

Homepage: <https://agenticscience.github.io/>

Github: <https://github.com/AgenticScience/Awesome-Agent-Scientists>

● 论文十

Deep Think With Confidence

日期

2025-08-21

DEEP THINK WITH CONFIDENCE

Yichao Fu^{2†*}, Xuwei Wang¹, Yuandong Tian¹, Jiawei Zhao^{1†}

¹Meta AI, ²UCSD

[†]Equal contribution

Project Page: jiaweizzhao.github.io/deepconf

内容摘要

本文提出了一种高效增强大语言模型推理能力的新方法(DeepConf), 针对当前主流“自洽多数投票”推理存在的计算成本高、准确率提升受限等问题。DeepConf 利用模型内部的置信度信号, 在推理生成过程中动态过滤低质量推理轨迹, 无需额外训练或调参, 可直接集成于现有服务框架。该方法在离线和在线推理场景下均显著提升了推理准确率, 并大幅减少生成 token 数 (最高可减少 84.7%), 在如 AIME 2025 等高难度数学推理基准上达到 99.9% 的准确率。本文系统分析了多种置信度量 (如局部组置信度、尾部置信度、低分段置信度) 对推理质量的判别能力, 并通过自适应采样机制实现了推理效率和性能的最佳平衡。实验覆盖多个开源高性能基础模型和高难度数据集, 实验结果显示 DeepConf 在规模和性能上均具备极强扩展性和应用价值。

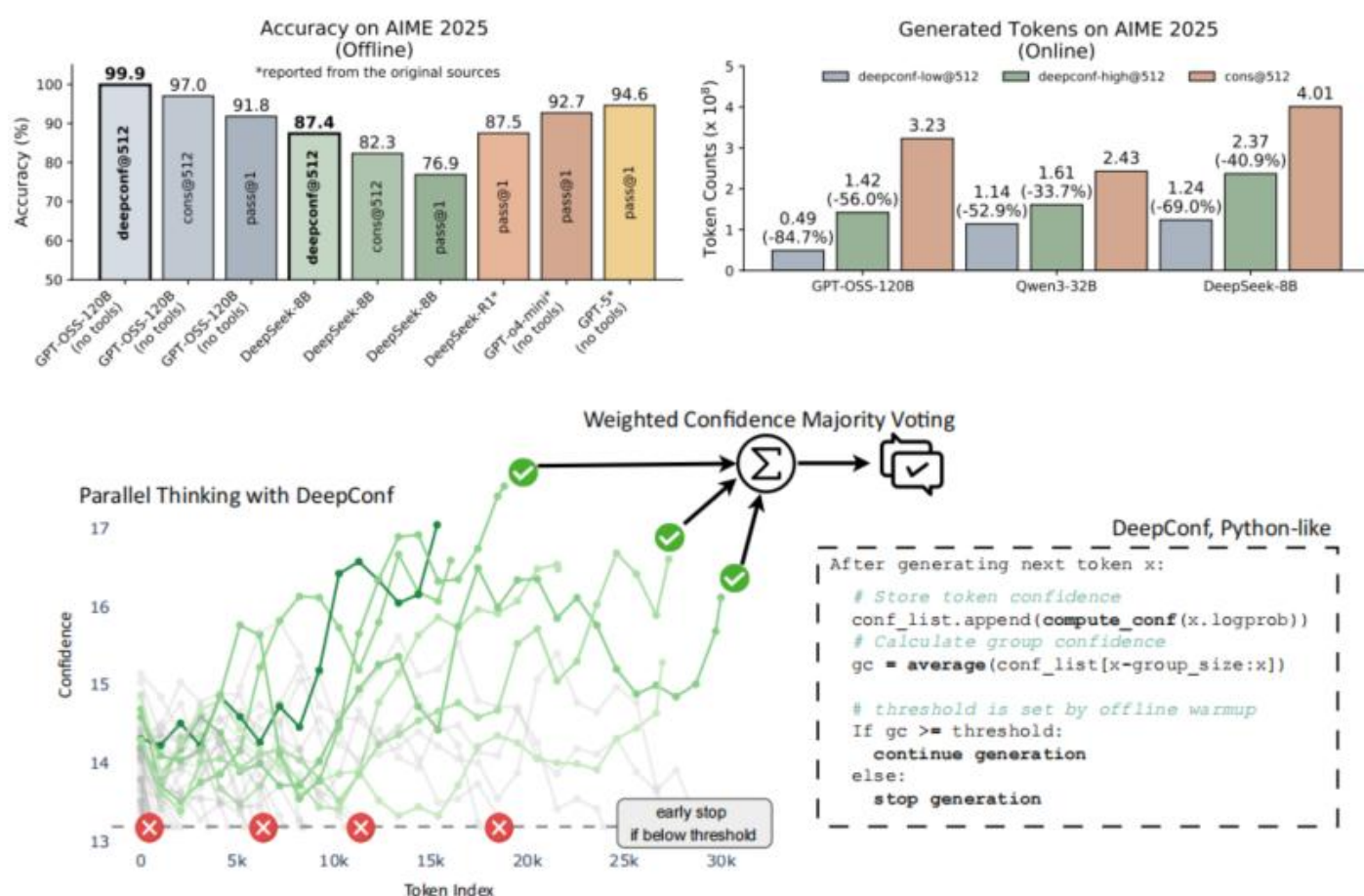


Figure 1: Up: DeepConf on AIME 2025. Down: Parallel thinking using DeepConf.

主要亮点

- **高效推理与大幅压缩计算**：DeepConf 利用模型内部置信度信号，有效过滤低质量推理路径，在不影响甚至提升准确率的前提下，极大减少了生成 token 的数量，最高节省计算量达 84.7%。
- **无需额外训练，兼容性强**：该方法无需重新训练模型或超参数调优，可直接集成到主流大模型推理服务流程，极大降低了部署和应用门槛。
- **多维置信度量，精准识别推理质量**：通过局部组置信度、尾部置信度等多种指标，实现对推理轨迹质量的细粒度判别，显著提升了模型在复杂数学等高难度任务上的推理表现。

相关链接

Paper: <https://arxiv.org/pdf/2508.15260>

● 论文十一

InternVL3.5: Advancing Open-Source Multimodal Models in Versatility, Reasoning, and Efficiency

日期

2025-08-25

InternVL3.5: Advancing Open-Source Multimodal Models in Versatility, Reasoning, and Efficiency

Weiyun Wang*, Zhangwei Gao*, Lixin Gu*, Hengjun Pu*, Long Cui*, Xingguang Wei*, Zhaoyang Liu*,
Linglin Jing*, Shenglong Ye*, Jie Shao*, Zhaokai Wang*, Zhe Chen*, Hongjie Zhang, Ganlin Yang, Haomin Wang,
Qi Wei, Jinhui Yin, Wenhao Li, Erfei Cui, Guanzhou Chen, Zichen Ding, Changyao Tian, Zhenyu Wu, Jingjing Xie,
Zehao Li, Bowen Yang, Yuchen Duan, Xuehui Wang, Zhi Hou, Haoran Hao, Tianyi Zhang, Songze Li, Xiangyu Zhao,
Haodong Duan, Nianchen Deng, Bin Fu, Yinan He, Yi Wang, Conghui He, Botian Shi, Junjun He, Yingdong Xiong,
Han Lv, Lijun Wu, Wenqi Shao, Kaipeng Zhang, Huipeng Deng, Biqing Qi, Jiaye Ge, Qipeng Guo, Wenwei Zhang,
Songyang Zhang, Maosong Cao, Junyao Lin, Kexian Tang, Jianfei Gao, Haian Huang, Yuzhe Gu, Chengqi Lyu,
Huanze Tang, Rui Wang, Haijun Lv, Wanli Ouyang, Limin Wang, Min Dou, Xizhou Zhu, Tong Lu, Dahua Lin,
Jifeng Dai, Weijie Su, Bowen Zhou[✉], Kai Chen[✉], Yu Qiao[✉], Wenhai Wang^{✉†}, Gen Luo^{✉†}

InternVL Team, Shanghai AI Laboratory

Code: <https://github.com/OpenGVLab/InternVL>

Model: https://huggingface.co/OpenGVLab/InternVL3_5-241B-A28B

内容摘要

本文介绍了 InternVL3.5 系列开源多模态大模型的最新进展。该模型通过创新性的“级联强化学习”（Cascade RL）策略，将离线和在线强化学习有机结合，极大提升了推理能力和训练稳定性；同时引入“视觉分辨率路由器”（ViR）与“视觉-语言解耦部署”（DvD）技术，实现视觉信息动态压缩与计算资源高效分配，

使推理速度提升至前代的 4 倍以上，且性能几乎无损。InternVL3.5 在多项主流多模态、推理、文本及智能体任务上表现卓越，最大规模版本在开源模型中综合得分领先，显著缩小与 GPT-5 等闭源商用模型的差距。该模型还支持 GUI 交互与智能体应用，展现出高度通用性与实际落地潜力。所有模型与代码均已开源，为多模态 AI 研究和工程应用提供了强大工具和新范式。

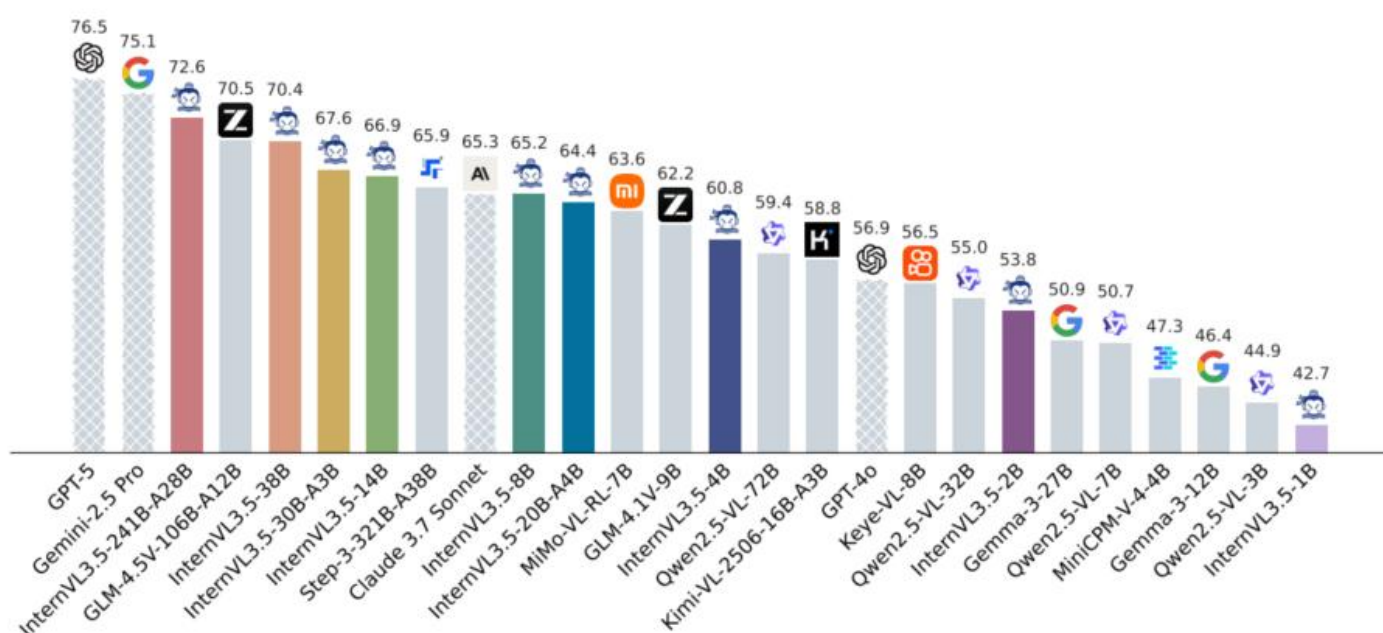


Figure 1: **Comparison between InternVL3.5 and leading MLLMs in general capabilities.** Hatched bars represent closed-source commercial models. We report average scores on a set of multimodal general, reasoning, text, and agentic benchmarks: MMBench v1.1 (en) [71], MMStar [11], BLINK [36], HallusionBench [41], AI2D [55], OCRBench [72], MMVet [168], MME-RealWorld (en) [178], MVBench [63], VideoMME [35], MMMU [170], MathVista [76], MathVision [134], MathVerse [175], DynaMath [189], WeMath [100], LogicVista [153], MATH500 [45], AIME24 [84], AIME25 [85], GPQA [106], MMLU-Pro [146], GAOKAO [177], IFEval [185], SGP-Bench [102], VSI-Bench [161], ERQA [121], SpaCE-10 [38], and OmniSpatial [50].

主要亮点

- 详细探讨了大型语言模型（LLMs）的发展历程、扩展规律和架构策略，揭示了数据和计算资源如何推动其性能提升。
- **创新级联强化学习框架**：提出了级联强化学习（Cascade RL），将离线与在线强化学习高效融合，显著提升模型的推理稳定性与泛化能力，在多项复杂多模态推理任务中实现领先性能。

- **高效视觉-语言协同机制**：引入视觉分辨率路由器（ViR）和视觉-语言解耦部署（DvD）技术，实现视觉信息的智能压缩与跨 GPU 并行推理，在保证性能的前提下，将模型推理效率提升至前代的 4 倍以上。
- **全面领先的开源多模态能力**：InternVL3.5 在多项通用、多模态、推理、智能体等基准任务上均取得开源模型领先成绩，最大规模版本综合表现接近商用顶级模型（如 GPT-5），并已实现全量开源，极大推动了多模态 AI 技术的开放与应用。

相关链接

Paper: <https://arxiv.org/pdf/2508.18265>

Code: <https://github.com/OpenGVLab/InternVL>

Model: https://huggingface.co/OpenGVLab/InternVL3_5-241B-A28B

● 论文十二

Exploring the role of large language models in the scientific method: from hypothesis to discovery

日期

2025-08-29

Exploring the role of large language models in the scientific method: from hypothesis to discovery

 Check for updates

Yanbo Zhang¹, Sumeer A. Khan^{2,3,4}, Adnan Mahmud^{5,6}, Huck Yang⁷, Alexander Lavin⁸, Michael Levin^{1,9},
Jeremy Frey¹⁰, Jared Dunnmon¹¹, James Evans^{12,13}, Alan Bundy¹⁴, Saso Dzeroski¹⁵,
Jesper Tegner^{2,16,17,18,19} & Hector Zenil^{20,21,22,23,24} ✉

内容摘要

本文系统梳理了大语言模型（LLMs）对科学方法的重塑作用，探讨其在假设提出到科学发现全过程中的潜力与局限。文章首先回顾了 LLMs 在科学研究中的多样化应用，包括文献综述、知识总结、实验设计、数据分析及自动化实验等，强调其在化学、生物等领域已显著提升科研效率。进一步，作者分析了基础模型（如 GPT-4、Evo、ChemBERT）在多模态科学数据处理、跨领域知识融合及零样本预测等方面的突破，展现出 LLMs 作为科学“创意引擎”的前景。文章还深入探讨了 LLMs 在科学发现中的挑战，如推理能力不足、可解释性有限及“幻觉”现象，并提出通过人机协作、算法置信度评估及多轮验证等方式加以应对。总体而言，作者认为 LLMs 有望推动科学范式转变，但其深度集成需以科学目标和严格评估体系为导向。

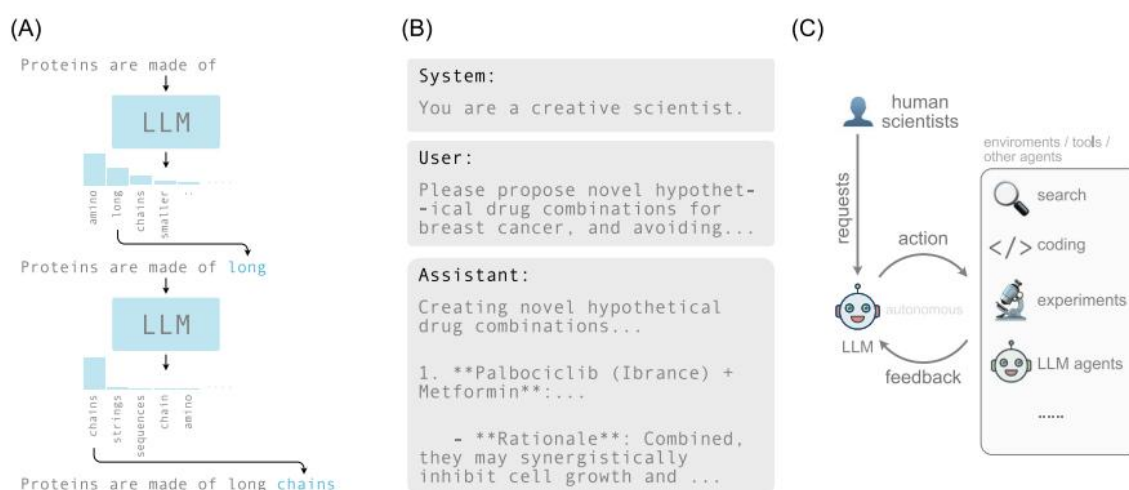


Fig. 1 | Auto-regressive generation, dialogue prompting, and agent frameworks in large language models. A LLMs generate sentences in an auto-regressive manner, sampling tokens from a predicted distribution at each step. B A typical prompt for LLMs consists of a system prompt and a user prompt. The LLM will then respond as an

assistant. A multi-round dialogue will repeat the user and assistant contents. C LLM agents are systems that use a large language model as its core reasoning and decision-making engine, enabling it to interpret instructions, plan actions, and autonomously interact with external tools, environments, or other LLM agents to fulfil a given goal.

主要亮点

- **全流程赋能科学研究：**文章系统梳理了 LLMs 在科学方法各环节中的应用，从文献综述、假设生成到实验设计和自动化执行，展示其对科学研究流程的深度重塑和效率提升。
- **基础模型推动跨领域创新：**通过分析 GPT-4、Evo、ChemBERT 等基础模型在多模态数据处理和零样本推理等方面的突破，强调 LLMs 跨学科知识融合与创新能力，为科学发现开辟新路径。

- **挑战与展望并重**：文章深入讨论了 LLMs 在推理、解释性和“幻觉”问题上的局限，提出人机协作、置信度评估与多轮验证等解决方案，呼吁建立科学目标引导和严谨评估体系以实现 AI 驱动的科学范式转型。

相关链接

Paper: <https://www.nature.com/articles/s44387-025-00019-5.pdf>

● 论文十三

Generalist Reward Models: Found Inside Large Language Models

日期

2025-08-29

GENERALIST REWARD MODELS: FOUND INSIDE LARGE LANGUAGE MODELS

A PREPRINT

Yi-Chen Li*, Tian Xu*, Yang Yu*,[†] Xuqin Zhang, Xiong-Hui Chen,
Zhongxiang Ling, Ningjing Chao, Lei Yuan, Zhi-Hua Zhou

National Key Laboratory for Novel Software Technology, Nanjing University, China
School of Artificial Intelligence, Nanjing University, China

* Equal contribution

[†] Correspondence: yuy@nju.edu.cn

内容摘要

本文提出了一种颠覆性的大语言模型（LLM）对齐方法，发现强大的通用奖励模型实际上已隐含于任何采用标准下一个词预测训练的 LLM 之中。作者通过理论证明，这一“内生奖励”与离线逆向强化学习（IRL）学习到的奖励函数在本质上是等价的，因此无需额外收集人类偏好数据或训练奖励模型，仅通过分析模型的 logits

即可直接提取高质量奖励信号。进一步，作者证明基于该内生奖励进行强化学习微调可使策略在误差界上优于原始模型，这是首次给出 LLM 强化学习有效性的理论证明。实验结果结果显示，该方法不仅超越现有

“LLM-as-a-judge” 无训练基线，还超过昂贵训练的人类标注奖励模型，展现了高度可扩展性和可个性化的对齐能力。本文开创性地推动了 LLM 自我对齐和多模态模型高效自评的新范式，对未来 AI 自主性与可控性具有重要意义。

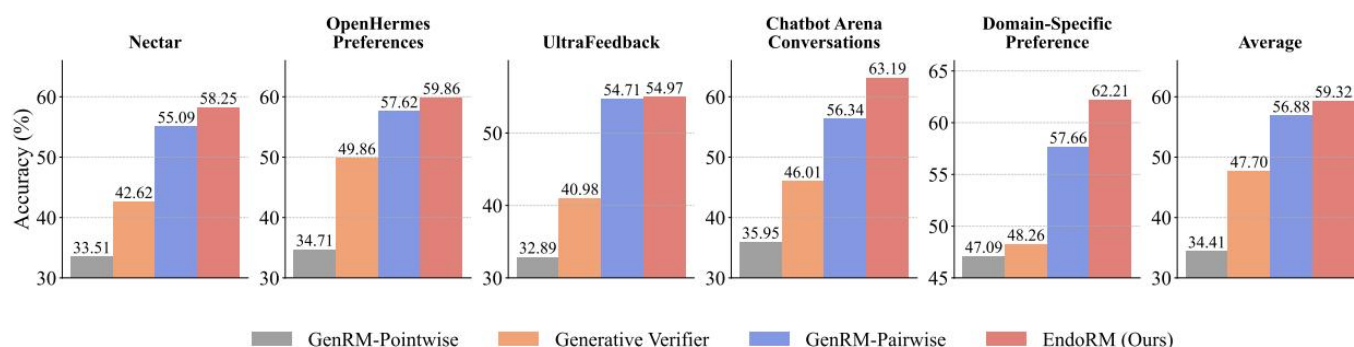


Figure 1: Response classification accuracy on the Multifacted-Bench.

主要亮点

- **理论突破与统一视角：**本文首次严格证明大语言模型在标准下一个词预测训练过程中已隐含通用奖励函数，并与离线逆向强化学习所得奖励理论等价，为奖励信号的提取和模型对齐建立了统一且坚实的理论基础。
- **方法创新与流程简化：**提出无需额外收集人类偏好数据或训练奖励模型，仅通过分析 LLM 的 logits 即可直接提取高质量奖励信号，实现训练和对齐流程的高度融合，大幅降低工程复杂度和资源消耗，推动了我对齐范式的发展。
- **性能提升与泛化能力：**实验证实该方法在多项主流基准任务上优于现有无监督及有监督奖励模型，并具备强大的指令可控性与偏好定制能力，展现出卓越的泛化性和可扩展性，为多模态、个性化 AI 对齐提供了高效解决方案。

相关链接

Paper: <https://arxiv.org/pdf/2506.23235>

TECHNICAL UPDATES
技术动态

● 技术动态一

OpenAI 推出 gpt-oss 系列开源模型，Gpt-oss-120b 和 gpt-oss-20b 以推动开源推理模型领域的技术边界

日期

2025-08-05

Model performance

[Read the research blog](#)

	gpt-oss-120b	gpt-oss-20b	OpenAI o3	OpenAI o4-mini
Reasoning & knowledge				
MMLU	90.0	85.3	93.4	93.0
GPQA Diamond	80.1	71.5	83.3	81.4
Humanity's Last Exam	19.0	17.3	24.9	17.7
Competition math				
AIME 2024	96.6	96.0	95.2	98.7
AIME 2025	97.9	98.7	98.4	99.5

内容摘要

OpenAI 发布了 gpt-oss-120b 和 gpt-oss-20b 两款开源推理模型，采用 Apache 2.0 许可与 gpt-oss 使用政策。模型兼容 Responses API，支持代理型工作流，具备强指令跟随、工具调用（如网页搜索与 Python 执行）、可调节推理能力，并提供完整链路推理（CoT）和结构化输出。

在安全方面，OpenAI 提示开源模型存在不同于专有模型的风险：一旦公开，攻击者可能通过微调绕过安全限制或用于有害目的。因此，部分应用场景下开发者和企业仍需额外实施防护措施。

OpenAI 对 gpt-oss-120b 进行了大规模能力评估，结果显示其在生化风险、网络安全、AI 自我改进三大类别下均未达到高能力阈值。即便在对抗性微调条件下，该模型仍未表现出高风险能力。此次发布体现了 OpenAI 致力于推动有益 AI 发展并提升行业安全标准的承诺。

相关链接

Blog: <https://openai.com/index/gpt-oss-model-card/>

Huggingface: <https://huggingface.co/openai/gpt-oss-120b>

● 技术动态二

谷歌推出 Genie 3，能够生成多样化交互环境的通用世界模型

日期

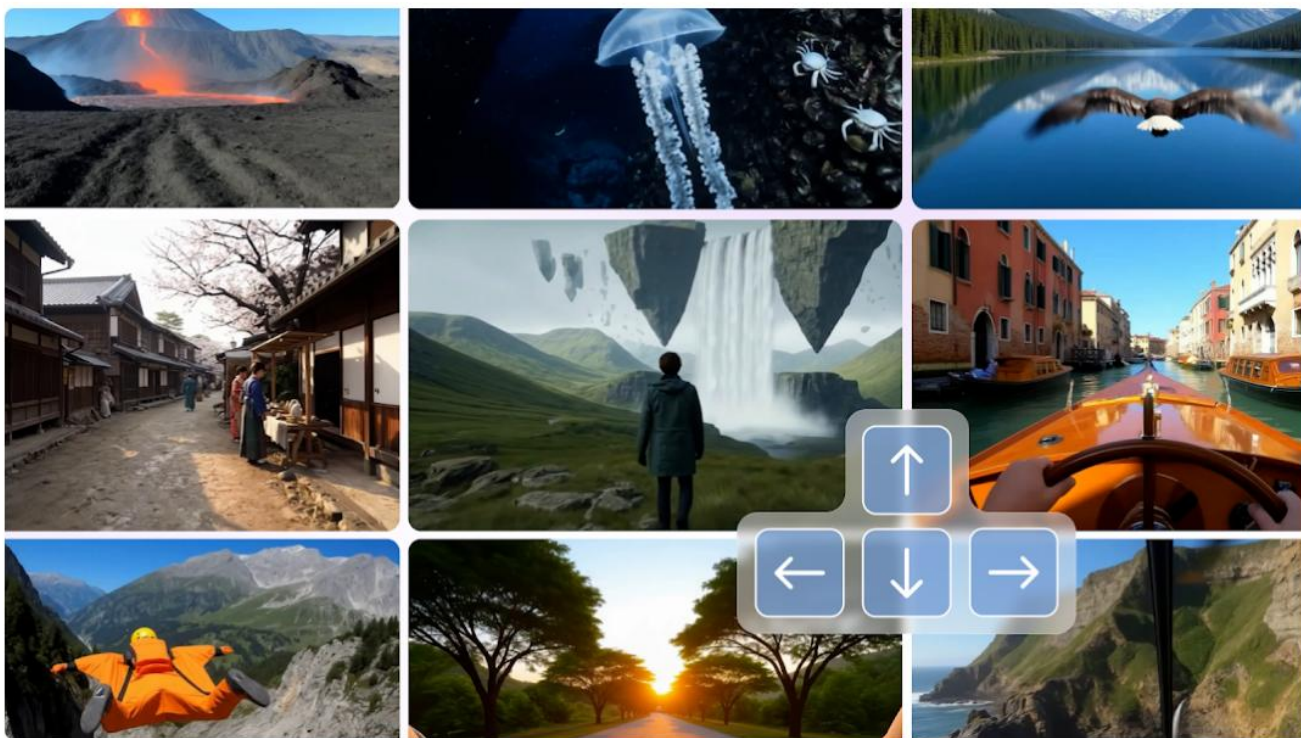
2025-08-05

MODELS

Genie 3: A new frontier for world models

5 AUGUST 2025

Jack Parker-Holder and Shlomi Fruchter

[Share](#)

内容摘要

DeepMind 于 2025 年 8 月推出了 Genie 3, 这是一个全新的通用世界模型, 能够根据用户的文本提示即时生成可探索、可交互的虚拟环境。与以往只能输出图像或短视频的生成模型不同, Genie 3 可以在 720p、24fps 的帧率下生成数分钟的动态世界, 并保持物体与环境的一致性, 例如即使物体暂时离开视野, 重新出现时依旧存在原位, 体现出“物体永久性”的特征。用户还可以在体验过程中输入新的提示, 实时改变天气、添加角色或修改环境, 模型能够即时响应而无需重新生成整个场景。

这一进展标志着“世界模型”研究的重要突破。Genie 3 不仅展示了更强的因果建模与物理一致性能力，也被认为是迈向更通用人工智能（AGI）的关键一步。它的潜在应用场景十分广泛，包括游戏设计、机器人训练、教育模拟和历史环境重构等，能大幅降低虚拟世界构建的门槛。作为 DeepMind “Genie” 系列的最新版本，Genie 3 延续了从 2D 到 3D、从静态生成到高度互动的演进路线，被视为交互式生成式 AI 的新前沿。

相关链接:

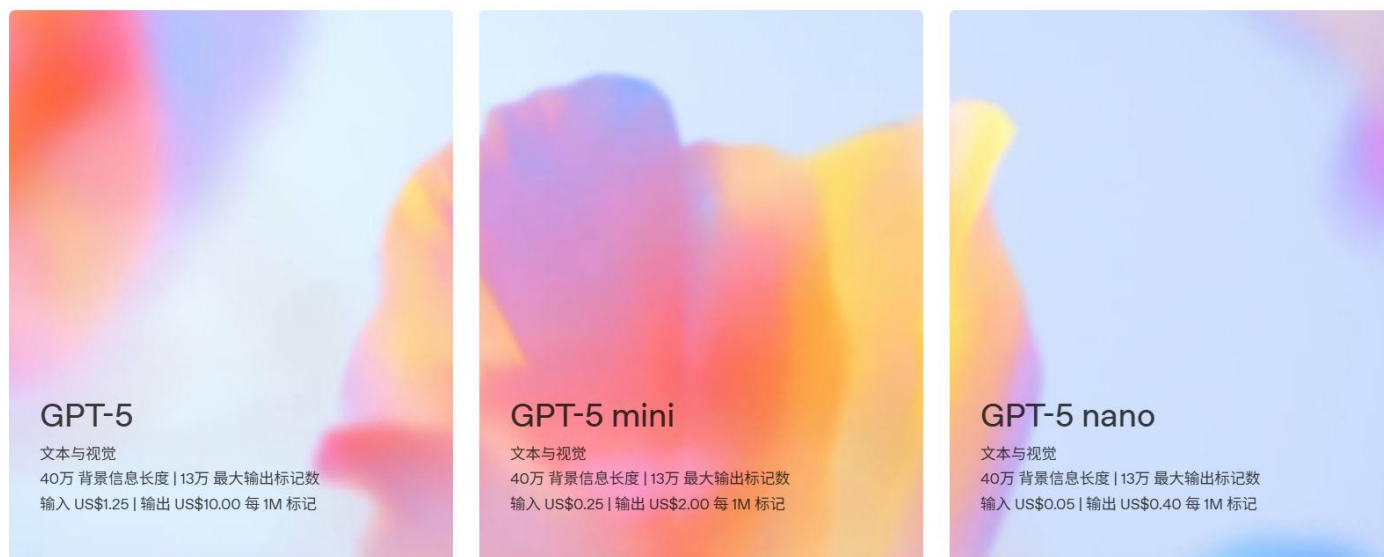
Blog: <https://deepmind.google/discover/blog/genie-3-a-new-frontier-for-world-models/>

● 技术动态三

OpenAI 推出 GPT-5，让专业级智能触手可及

日期

2025-08-07



内容摘要

在功能方面，GPT-5 显著提升了指令跟随、代码生成、工具调用和多步推理能力。例如在 SWE-Bench Verified 上达到了 74.9% 的成绩，在多语言编程和前端开发任务中表现优于前代模型。同时，它在多模态理解上也有进步，能够更好地处理数学、科学、写作和视觉任务，事实准确性和安全性均有所增强。

GPT-5 可通过 ChatGPT 与 API 使用，提供不同版本以满足性能、成本与延迟的平衡。普通用户可免费体验有限额度，Plus 和 Pro 用户可获得更高的使用权限，其中 Pro 版本具备更强推理能力。多家企业已在实际应用中采用 GPT-5，反馈其在处理模糊语境时更精准、更可靠。

相关链接

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

● 技术动态四

DeepSeek-V3.1 发布，迈向 Agent 时代的第一步

日期

2025-08-21

Benchmarks	DeepSeek-V3.1	DeepSeek-V3-0324	DeepSeek-R1-0528
SWE-bench Verified	66.0	45.4	44.6
SWE-bench Multilingual	54.5	29.3	30.5
Terminal-Bench	31.3	13.3	5.7

表 1：编程智能体测评（SWE 使用内部框架测评，相比开源框架 OpenHands 所需轮数更少；Terminal Bench 使用官方 Terminus 1 framework）

内容摘要

DeepSeek-V3.1 正式发布，首次在同一 685B-MoE 模型内实现“思考/非思考”混合推理：官方 App、网页一键“深度思考”切换；API 端点 deepseek-reasoner 与 deepseek-chat 共享 128K 上下文，并新增 Beta 版支持 strict 模式 Function Calling 及 Anthropic 格式，可无缝嵌入 Claude Code。经过思维链压缩训练，思考模式在输出 token 减少 20-50%的情况下仍保持 R1-0528 级表现，非思考模式亦显著缩短回复长度。后续训练重点聚焦于工具调用和 Agent 能力优化，编程智能体在 SWE 代码修复与 Terminal-Bench 终端任务、搜索智能体在 BrowseComp 多步推理及 HLE 多学科难题评测中，性能全面超越历史版本。baseline 与后训练权重已同步开源，并首次采用面向下一代国产芯片的 UE8M0 FP8 精度优化，向通用智能体迈出关键一步。

相关链接

Blog: <https://mp.weixin.qq.com/s/WUbmBSapVyvxZe6HobD5Qw?scene=1>

● 技术动态五

轻量级易开发，8B 参数释放大实力！科学多模态模型 Intern-S1-mini 开源

日期

2025-08-21



[🤗 Huggingface](#) • [🔗 ModelScope](#) • [📄 Technical Report](#) • [💬 Online Chat](#)

[English](#) | [简体中文](#)

👉 join us on [Discord](#) and [WeChat](#)

内容摘要

上海人工智能实验室继 7 月开源科学多模态大模型 Intern-S1 后，近日推出轻量化版本 Intern-S1-mini。该模型拥有 8B 参数，兼具通用与专业科学能力，适合快速部署和二次开发。在 MMLU-Pro、AIME2025、MMMU 等权威基准上表现卓越，在化学、材料等科学任务中更是显著领先，同时在物理、地球、生物等学科也保持第一梯队水平。Intern-S1-mini 延续大模型设计理念，覆盖文本、图像、分子式、蛋白质等多模态数据领域，通过优化能力融合算法，实现轻量模型下优秀的综合实力。该模型参数规模适中，大幅降低了对高端计算设备的依赖，仅需 24GB 单卡即可完成 LoRA 微调，极大降低上手门槛。无论在科研、开发还是教育场景，Intern-S1-mini 均能提供强大支持。未来，上海 AI 实验室将持续推进该模型及工具链开源，支持免费商用，助力科学 AI 生态繁荣。

相关链接

Bolg: <https://mp.weixin.qq.com/s/4Mj8Df4kmTqTX06zVvYIng>

● 技术动态六

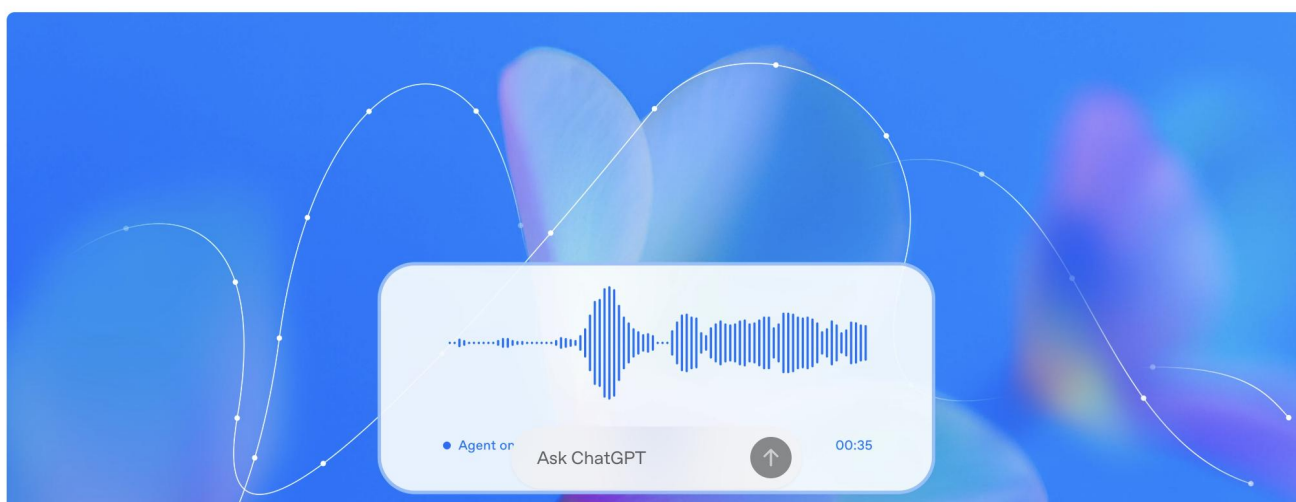
语音模型第一次超越人类！OpenAI 再现 Her 时刻，95 后华人研究员坐镇

日期

2025-08-29

Introducing gpt-realtime and Realtime API updates for production voice agents

We're releasing a more advanced speech-to-speech model and new API capabilities including MCP server support, image input, and SIP phone calling support.



内容摘要

OpenAI 在最新发布中推出了两项重磅 AI 语音更新：Realtime API 和 gpt-realtime，极大提升了语音智能体的能力。以下是其主要亮点总结：

- **Realtime API 的突破：**Realtime API 支持生产级实时语音智能体，具备高可靠性、低延迟和高音质的特点。它通过单一模型直接处理语音到语音任务，简化了传统多层处理链路。此外，API 可支持远程 MCP 服务器和图像输入，通过 SIP 协议实现电话功能，进一步扩展了智能体的适用场景和工具调用能力。
- **gpt-realtime 的领先性能：**作为最先进的语音到语音模型，gpt-realtime 在音质、智能理解、指令遵循及函数调用方面全面提升。它生成的语音自然流畅，具备情感表现力，可精准复述复杂指令、切换语言和捕捉非语言线索（如笑声）。模型在多项评测中表现优异，例如推理能力的 Big Bench Audio（82.8% 准确率）和函数调用的 ComplexFuncBench（66.5% 得分）。
- **多模态支持与开发者友好：**gpt-realtime 引入图像输入功能，支持将视觉信息与语音交互结合，增强用户体验。同时，其异步函数调用优化、指令遵循能力提升以及内置工具调用能力，使开发者能够更高效地创建灵活且易扩展的语音智能体。

相关链接

Blog: <https://mp.weixin.qq.com/s/pFQV5OVJ2OjVXJQAcO-HzA>

RESOURCE SHARING

资源分享

● 资源分享一

《Kitten TTS》

类型：模型资源

内容提要

KittenTTS 是一个超轻量级（<25MB）、开源的高质量文本转语音（TTS）模型，其核心优势在于能在纯 CPU 环境下高效运行，无需 GPU 支持，并提供了多种自然的语音选项，特别适用于资源受限的设备或需要低成本、快速语音生成的边缘计算和嵌入式场景。

资源链接

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

● 资源分享二

《Canary 1B v2: Multitask Speech Transcription and Translation Model》

类型：模型资源

内容提要

NVIDIA Canary-1B-v2 是由 NVIDIA 开发并开源的一个高性能、多语言语音翻译与合成模型 (TTS)，参数规模为 10 亿 (1B)。它专为高质量的端到端语音翻译和语音合成而设计，支持英语、西班牙语、德语和法语，能够将源语言语音直接翻译并合成为目标语言的语音，或进行高质量的纯文本到语音合成，生成具有自然韵律和保真度的语音，适用于会议记录、媒体制作等专业场景

资源链接

<https://huggingface.co/nvidia/canary-1b-v2>

● 资源分享三

《Qwen3-235B-A22B-Instruct-2507》

类型：模型资源

内容提要

Qwen3-235B-A22B-Instruct-2507 是由 Qwen 团队发布的一个超大规模 (2350 亿参数)、高性能的开源指令微调大语言模型 (LLM)，其基于创新的“A22B”架构 (推测为混合专家 MoE 或类似变体) 进行设计。它专为复杂推理、多轮对话和中文任务优化，具备强大的中文理解与生成能力、超长上下文处理 (如 128K tokens) 以及优秀的代码、数学和工具使用能力

资源链接

<https://huggingface.co/Qwen/Qwen3-235B-A22B-Instruct-2507>

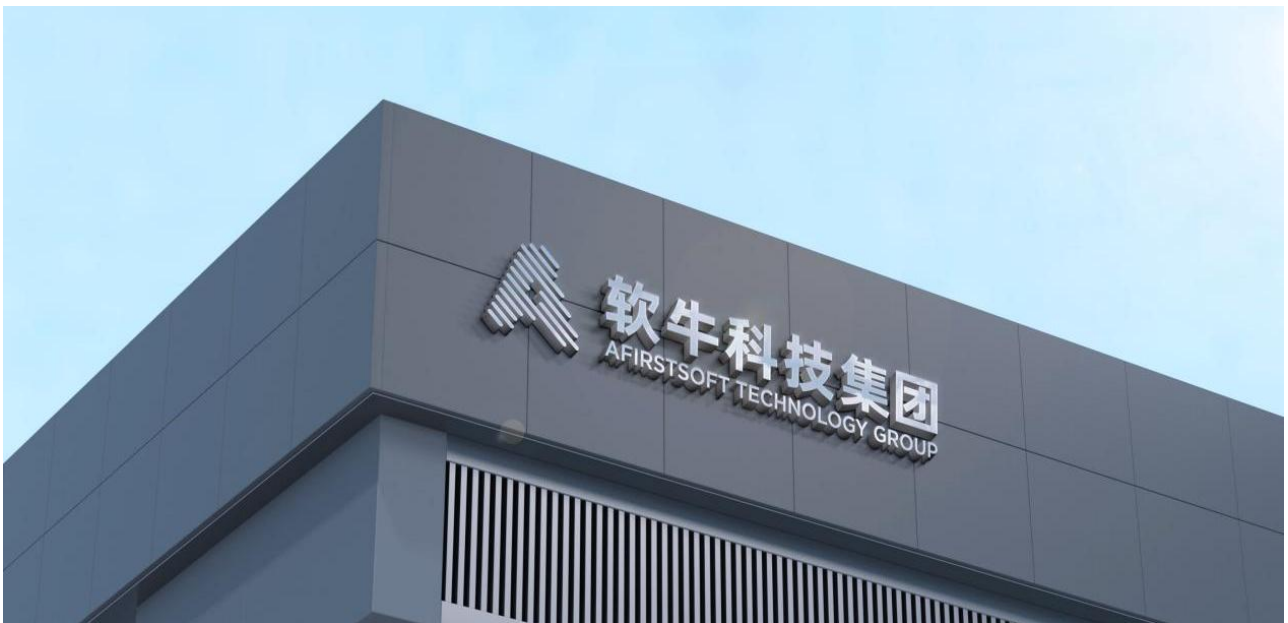


深圳软牛科技集团股份有限公司成立于 2014 年 6 月，作为国家级专精特新“小巨人”企业，始终以“让创意成为现实”为使命，专注于通过人工智能视觉大模型技术驱动全球数字创意生态。

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 08 期

总第 26 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
黄旭进	金山办公
胡晋武	华南理工大学
王宇丰	华南理工大学
王骞玥	华南理工大学
李成昊	华南理工大学
刘祎然	华南理工大学
许毅平	华南农业大学

人工智能 前沿快报

2025 年第 08 期

总第 26 期

— 2025.09.02

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，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