
“令学界振奋！”《自然》发文盛赞中国开源AI模型DeepSeek

作者：writer 来源：科学网

本文原地址：<https://www.iikx.com/news/progress/31539.html>

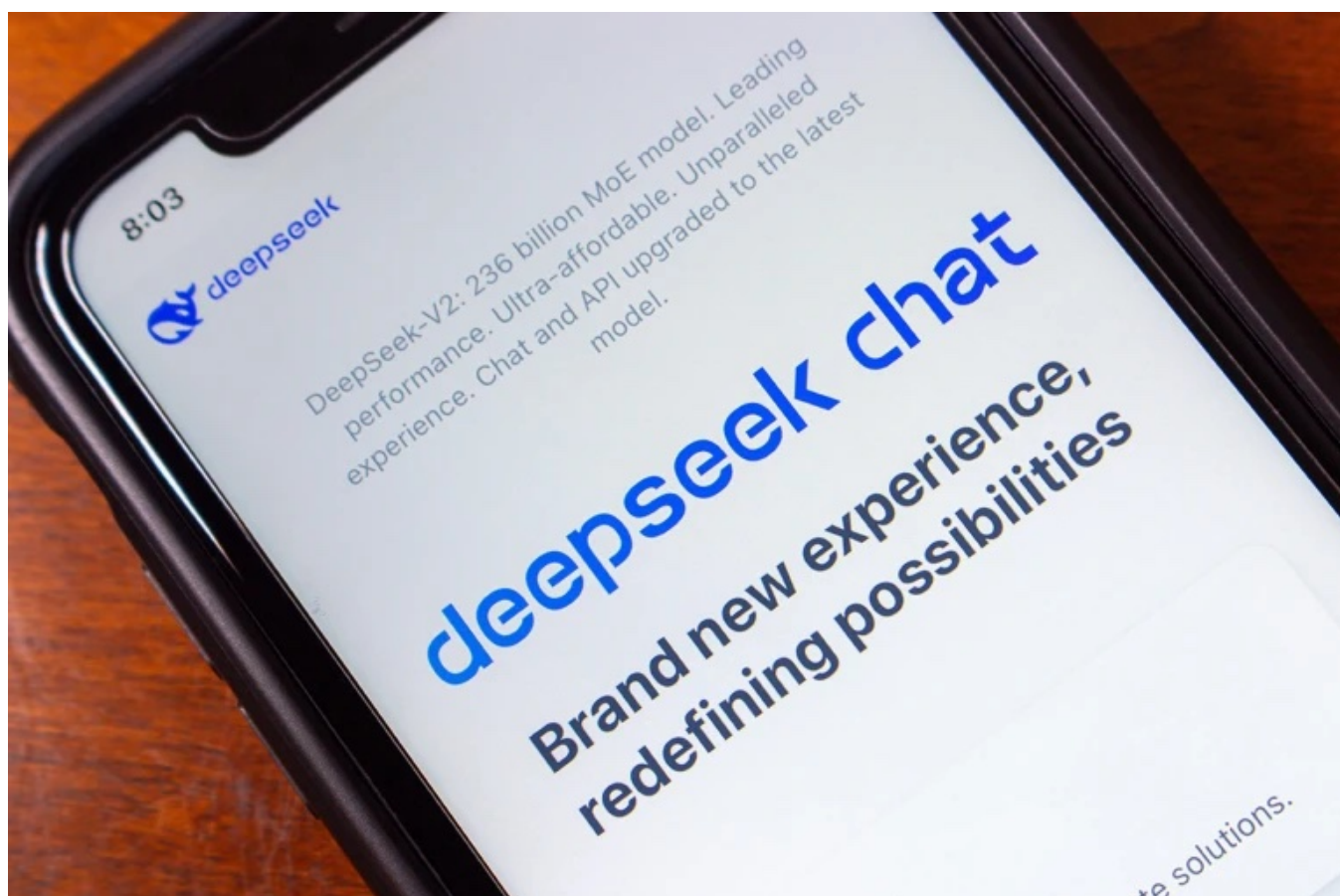
本文仅供学习交流之用，版权归原作者所有，请勿用于商业用途！

“令学界振奋！”《自然》发文盛赞中国开源AI模型DeepSeek

。最近，由来自杭州的“深度求索”初创团队开发的DeepSeek系列AI模型，引发了全球范围的关注。1月24日，知名学术期刊《自然》也发文关注该模型及相应产品，称“中国开发的大语言模型DeepSeek-R1以亲民价格和开放性挑战了OpenAI的推理模型GPT-o1的地位，令科学家们感到兴奋”。

《中国科学报》了解到，深度求索团队于1月20日在线发布DeepSeek-R1，并同步开源模型权重。同日，DeepSeek官网和对应APP同步更新上线。

对于DeepSeek-R1的表现，《自然》文章评论称，初步测试显示，R1在化学、数学和编程领域的特定任务表现与2024年9月令学界惊叹的GPT-o1旗鼓相当。



《自然》文章配图。图源：Nature

“这完全超出预期”，英国AI咨询公司DAIR.AI联合创始人埃尔维斯·萨拉维亚（Elvis Saravia）在社交平台上赞不绝口。

《自然》文章认为，这类模型通过类人推理的逐步响应生成机制，在解决科学问题方面展现出超越早期语言模型的能力，具有科研应用潜力。

“真正的‘Open-AI’”

该文章谈到，DeepSeek-R1的另一突破在于其开放性。开发方深度求索团队采用开放权重模式发布，允许研究者研究并改进算法。虽然基于MIT许可证（MIT License）可自由复用，但因未公开训练数据，尚未达到完全开源标准。

德国马克斯·普朗克光科学研究所人工科学家实验室负责人马里奥·科瑞恩（Mario Krenn）则

评价称：“深度求索的开放程度令人瞩目”，相较之下，OpenAI的o1及其最新o3模型“本质仍是黑箱”。

“DeepSeek才是真正的‘Open-AI’！”在深度求索团队发布DeepSeek-R1的网络文章下面，这条评论获得了最高赞。

DeepSeek尚未公布训练R1的全部投入花费，但它向使用其界面的人收取的费用约为o1运行费用的三十分之一。据深度求索团队发布内容，DeepSeek-R1的API（应用程序接口）服务定价为每百万输入 tokens 1 元（缓存命中）/ 4 元（缓存未命中），每百万输出 tokens 16 元。同时，深度求索还创建了R1的迷你“蒸馏”版本，使计算能力有限的研究人员也能够使用该模型。

“使用o1的实验成本超过300英镑，而使用R1的实验成本不到10美元。”科瑞恩说：“这是一个巨大的差异，肯定会在未来的采用中发挥作用。”

“美国的领先优势在缩小”

《自然》文章认为，DeepSeek-R1是中国大型语言模型（LLMs）繁荣的一部分，并祭出DeepSeek以“小算力”驱动“大模型”、训练成本低廉为论据：DeepSeek训练模型所需的硬件大约需要6000万美元，而Meta的Llama 3.1 405B则需要60000万美元，后者使用了11倍于前者的计算资源。该文还提到，深度求索团队此前发布的DeepSeek-V3的表现也优于主要竞争者，“且其预算很小”。

对此，美国华盛顿州西雅图的人工智能研究员弗兰科伊斯·夏洛特（Francois Chollet）表示：“它来自中国的事实表明，高效利用资源比单纯的计算规模更重要。”

《自然》文章评述称，此前围绕DeepSeek的部分话题是，尽管美国的出口管制措施限制了中国公司获得为人工智能处理设计的高端计算机芯片，但DeepSeek还是成功地训练出了R1。

华盛顿贝尔维尤的技术专家阿尔文·王·格雷林（Alvin Wang Graylin）在社交平台上写道，DeepSeek的进展表明，“美国曾经的领先优势已经显著缩小”。

用“思维链”减少“幻觉”损害

人们对人工智能大模型常常诟病的一点是，它会产生“幻觉”。

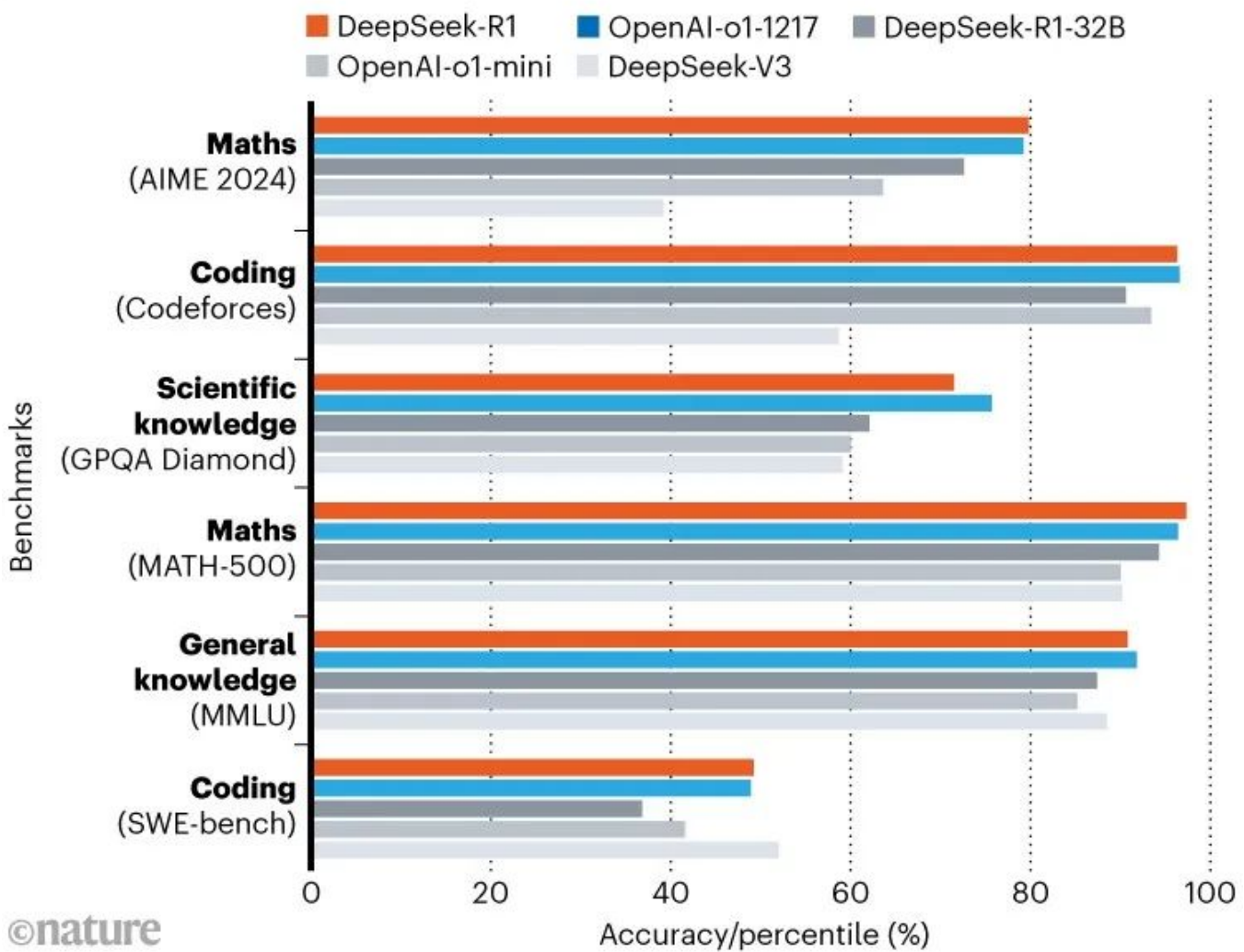
这是由于，大语言模型一般是在数十亿个文本样本上进行训练，将它们剪切成“标记”的单词部分并学习数据中的模式——这些关联使模型能够预测句子中的后续标记。但这些大模型倾向于编造事实，这是一种被称为“幻觉”的现象，并且经常难以通过推理解决或消除。

与OpenAI的o1一样，DeepSeek-R1也使用了“思维链”方法来提高大模型解决更复杂任务的能力，其机制包括过程回溯和策略评估。深度求索研发团队通过强化学习对V3模型进行“微调”训练：当模型获得正确答案或展示清晰“解题思路”时即给予正向反馈，从而塑造出R1的推理能力。

英国爱丁堡大学的人工智能研究员李文达（音）认为，正是计算能力有限的原因，促使深度求索团队“在算法上进行创新”。他指出，在强化学习阶段，团队会采用分阶段评估模型进展的监测方式，替代传统的独立验证网络法。英国剑桥大学的计算机科学家玛特亚·亚姆尼克（Mateja Jamnik）指出，这有助于降低训练和运行成本。研究人员还采用了混合专家（MoE）架构，使模型根据任务需求动态激活相应模块。

AI RIVALS

DeepSeek tested three versions of its large language models against OpenAI's o1 models on mathematics, coding and reasoning tasks. DeepSeek-R1 beat or rivalled o1 on maths and coding benchmarks.



DeepSeek-R1与OpenAI-o1-1217等模型在代码、科学知识、数学、常识方面的对比。图源：Nature

在加州大学伯克利分校设计的MATH-500数学题集上，DeepSeek-R1取得97.3%的准确率，并在Codeforces编程竞赛中超越96.3%的人类选手。这些成绩与o1相当（未纳入最新o3的对比测试）。

剑桥大学计算机科学家马可·多斯·桑多斯（Marco Dos Santos）指出，基准测试难以全面反映模型的真实推理与泛化能力，但得益于R1的开放性，研究者可解析其思维链条，“这极大提升了模型推理过程的可解释性”。

《自然》文章称，科研人员已展开实际测试。科瑞恩让这两个模型对3000个科研创意进行兴趣度排序，结果R1略逊于o1。但在量子光学特定计算中，R1展现出超越o1的实力，科瑞恩评价道，这确实令人印象深刻。

“DeepSeek 是化繁为简的大师”

1月26日，出门问问副总裁、Netbase前首席科学家李维发文表示，DeepSeek的创新和探索精神表现在，当社区把有监督的精调+强化学习（SFT+RL）当成是“后训练范式”的时候，他们做自主学习（Zero），完全排除人工数据，验证了纯粹的强化学习对于推理能力的学习潜力。

他指出，深度求索团队先是从Zero首先是学到了信心，体验了探索创新者的“啊哈时刻”（aha moment），然后又加入了一些用于冷启动的高质量人工数据做SFT，再做实用的R1就有底气了。

“两个模型都开源，供人研究和验证，做得煞是漂亮。”李维感叹：“DeepSeek是化繁为简的大师。”

相关文章链接：

<https://www.nature.com/articles/d41586-025-00229-6>

作者：赵广立 来源：中国科学报

更多 科学进展 请访问 <https://www.iikx.com/news/progress/>

本文版权归原作者所有，请勿用于商业用途，[爱科学iikx.com](https://www.iikx.com)转发