

科技评论

DeepSeek 技术创新与通用人工智能发展趋势

吴文峻^{1,2}, 廖星创¹, 赵金琨¹

摘要 概述了 DeepSeek 在通用人工智能领域的最新进展, 重点讨论了其在大语言模型、推理技术方面的创新。DeepSeek-V3 引入了新的模型架构和算法设计, 基于相对有限的智能硬件, 对模型训练方法进行了全面和深入的优化, 显著提升了模型训练效率。在推理技术方面, DeepSeek-R1 创新性地结合了强化学习 (RL) 与监督微调 (SFT), 提升了推理深度和逻辑推理能力。结合 DeepSeek 的创新工作, 讨论了通用人工智能发展趋势, 重点涉及 3 个问题: 开源开放生态对发展通用人工智能的作用; 依赖于模型规模扩展的“Neural Scaling Law”是否还能发挥作用; 如何基于 DeepSeek 这类基座模型, 以“通专结合”的方式实现行业大模型的落地等。

关键词 DeepSeek; 强化学习; 混合专家系统; 规模法则; 慢思考

近年来, 随着深度学习技术的快速发展, 大语言模型 (LLMs) 在自然语言处理、内容生成、推理能力等方面取得了显著进展。以 OpenAI 的 GPT 系列、Google 的 Gemini 和 Meta 的 LLaMA 为代表的模型, 不仅在通用任务上表现出色, 还在特定领域的应用中展现了强大的潜力。然而, 尽管这些大模型在规模和性能上不断突破, 其发展仍面临诸多挑战, 包括高昂的训练成本、推理效率的瓶颈, 以及模型在复杂任务中的逻辑推理能力不足等问题。

在此背景下, DeepSeek 应运

而生, 作为中国在通用人工智能 (AGI) 领域的重要探索, DeepSeek 不仅继承了现有大模型的优势, 还在模型架构、训练效率和推理能力等方面进行了创新性突破。DeepSeek 的推出标志着中国在人工智能领域的自主创新能力迈出了重要一步, 尤其是在大模型的低成本训练、推理能力提升以及开源生态建设方面, 为全球 AI 技术的发展提供了新的思路。

DeepSeek 于 2023 年 7 月由幻方公司提出, 2025 年 1 月 20 日, 伴随着 DeepSeek-R1 的正式发布, 其在社会各界引起了广泛关注和热烈讨论, 开启了基于国产的人工智能创新生态, 掀开了对通用人工智能进行深度探索的新篇章。目前, DeepSeek 模型家族主要分为 V 系列与 R 系列 2 大分支。其

中, V 系列大模型主要针对多领域对话与内容生成, 偏重通用性与自然语言覆盖广度; R 系列大模型主要强调推理与思维链, 以实现深度逻辑能力。

本文首先从 DeepSeek 的技术特点出发, 详细分析了其在模型架构和推理技术方面的核心创新。随后, 探讨了 DeepSeek 对通用人工智能未来发展的影响, 包括开源生态的作用、神经标度律 (Neural Scaling Law) 的局限性以及通专结合的行业模型。最后, 总结了 DeepSeek 的技术贡献, 展望了其在开源生态和垂直领域应用中的潜力。

1 DeepSeek-V3 的创新架构和低成本高性能训练方法

传统的大型模型训练, 无论是

1. 北京航空航天大学复杂关键软件环境
全国重点实验室, 北京 100191
2. 杭州市北京航空航天大学国际创新研究院, 杭州 311115

收稿日期: 2025-02-14; 修回日期: 2025-03-12

作者简介: 吴文峻, 教授, 研究方向为可信智能、群体智能、AI for Science, 电子邮箱: wwj09315@buaa.edu.cn

引用格式: 吴文峻, 廖星创, 赵金琨. DeepSeek 技术创新与通用人工智能发展趋势[J]. 科技导报, 2025, 43(6): 14-20; doi:10.3981/j.issn.1000-7857.2025.02.00175

闭源的 GPT 系列还是开源的 LLaMA 系列,都面临着巨大的图形处理器(GPU)资源消耗的挑战。然而,DeepSeek-V3^[1]模型以其创新的设计理念,巧妙地实现了低成本与高性能的完美融合,为这一难题提供了突破性的解决方案。该模型通过引入多头潜在注意力机制(MLA)、混合专家系统(MoE)的架构革新以及多 token 预测机制(MTP)等核心技术,不仅显著提升了推理效率,还有效优化了资源利用率,为大规模模型训练开辟了新的道路。

首先,多头潜在注意力通过低秩联合压缩技术减少了推理过程中的键值(KV)缓存需求,有效降低了内存占用,并在保持性能的同时提升了系统整体效率。其次,混合专家无损负载均衡策略通过引入精简的辅助损失和动态调节的偏置机制,解决了传统 MoE 模型中 token 分配不均的问题,确保了更平衡的负载分配,优化了计算资源利用,避免了少数专家负担过重的问题。最后,MTP 方法通过一次性预测多个后续 token,提升了预测效率和上下文理解能力,虽然该方法主要用于训练阶段以增强模型学习能力,但推理阶段并不使用此技术。通过这些创新,DeepSeek-V3 在保证高性能的同时,有效平衡了推理效率与资源消耗。

DeepSeek 通过创新性采用 FP8 混合精度训练框架,实现了大模型算法与智能硬件的高度协同优化,这一技术路径对推动中国 AI 基础设施自主化,实现基于国产软硬件的协同优化设计具有重要启示。FP8 格式凭借其 8 位浮

点计算位宽,较传统 FP16/FP32 方案能够显著降低显存占用与带宽压力,而国产芯片未来可通过定制化架构与自适应量化指令集,将 FP8 计算效率进一步提升。这种软硬件深度协同将推动中国在 AI 算力底座领域逐步形成从指令集扩展、编译器优化到框架调度的全栈自主创新能力。

DeepSeek-V3 以高达 671B (6710 亿)的参数量,将训练成本大幅降低至约 557 万美元(表 1),

与传统大模型动辄上亿美元的投入形成鲜明对比。例如, GPT-4o 和 Claude 3 的训练成本均约为 1 亿美元, Llama-3 约为 9000 万美元,而 Gemini Ultra 更是高达 1.9 亿美元;即便与国内模型相比, DeepSeek-V3 也展现出显著的成本优势,规模仅为 13B 的讯飞星火 X1,其训练成本仍高达 780 万美元。这一突破不仅彰显了 DeepSeek 的技术实力,更为行业提供了高效、经济的训练范式。

表 1 各种大模型训练成本对比

模型名称	Gemini Ultra	Llama-3	Claude 3	DeepSeek-V3	GPT-4o	讯飞星火 X1-13B
训练成本/美元	约1.9亿	约9000万	约1亿	约557万	约1亿	约780万

2 DeepSeek-R1 提升“慢思考”的认知推理能力

在人工智能领域,认知推理能力是衡量模型智能水平的重要指标之一。DeepSeek-R1^[2]通过创新的推理技术和架构设计(图 1),显著提升了“慢思考”能力——即系统化、逻辑化的深度推理能力。本节详细介绍 DeepSeek-R1 在推理技术上的突破,以及基于神经符号系统的推理架构如何实现更高效、更准确的认知推理。

2.1 DeepSeek-R1 推理技术简述

DeepSeek 研究团队发现大模型的推理能力可通过纯强化学习过程涌现出来,甚至可以完全绕过传统的监督微调过程。DeepSeek-R1-Zero 使用准确度奖励和格式奖励来指导模型推理的强化训练,尽管其推理表现尚不如其他模型,但通过生成中间思维步骤,成功展示了推理能力生成的可行性。这

一成果表明,纯强化学习方法在推理行为的自我涌现方面具有巨大潜力。

另外,DeepSeek-R1 进一步通过结合监督微调与强化学习来优化其推理性能。与 OpenAI o1^[3]通过将复杂任务拆解为若干步骤,并依次解决的过程奖励模型(PRM)不同,DeepSeek 在强化学习训练之前生成未经初始训练阶段的“冷启动”监督微调数据,随后通过指令微调进行训练,并引入准确度、格式一致性和语言一致性奖励机制避免语言混合等问题。通过采用监督微调和强化学习技术,DeepSeek 有效规避了 PRM 技术中依赖高质量的人工标注数据、难以明确定义细粒度步骤以及判断中间步骤正确性等关键问题,显著提升了模型在复杂推理任务中的性能表现。

2.2 基于神经符号系统的推理架构

DeepSeek-R1 和 OpenAI o1/o3

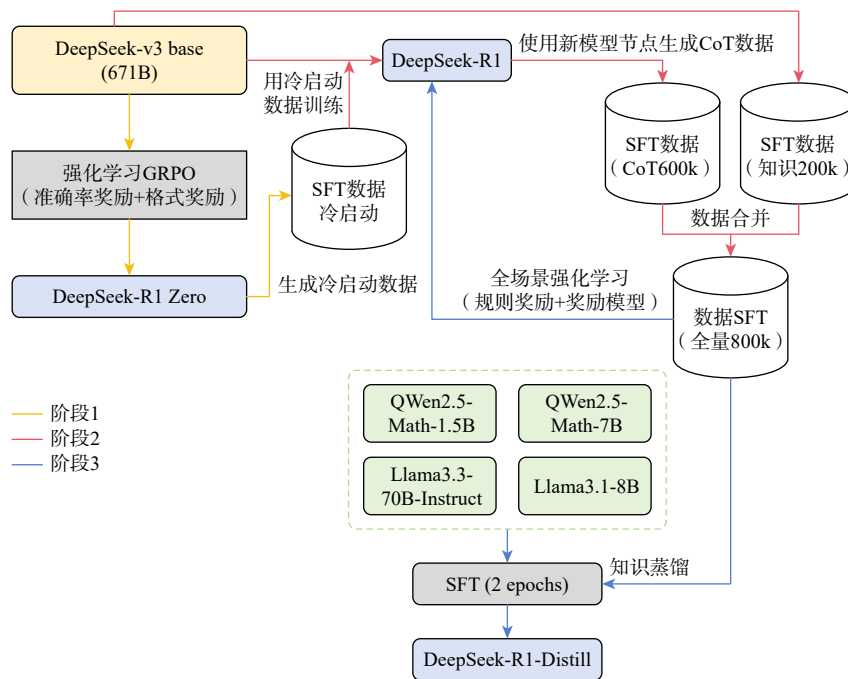


图1 DeepSeek 强化学习训练流程

这些大模型系统所取得的推理性能,标志着 LLMs 在推理方面的研究取得了新的突破,开启了这个领域的新范式,即系统 1(“快思考”)+系统 2(“慢思考”)。这个范式意味要对符号主义积累的成果和现有的大模型框架进行更深入的结合,可以在不同的情境中形成动态可变的、复杂思维链,以便在常识推理、数学推理、算法推理、科学推理、具身推理等方面持续提升,直至达到甚至超过人类的认知推理能力。

以常识推理为例,以前使用符号主义方法进行常识推理是非常困难的,因为经典的符号主义要依赖手工定义知识与规则,所以很难形成常识推理需要的海量知识。而大模型的出现为解决上述问题提供了契机。无论大语言模型还是多模态模型,都可以借助预训练-微调的方式,把互联网信息以训练语料的形态,灌输到大规模

transformer^[4]网络中,形成参数化的知识表达。这样大模型无论面对何种常识问题,都可以自如回答,在常识推理方面的能力有一定提升。

但是这些提升还远远不够。大模型虽然在各种自然语言问题的基准测试(benchmark)中表现不错,但是在真实的交互中,经常出现“幻觉”^[5]问题。核心原因是大模型的内容生成方式仍然遵循统计规律,只是以语言的 token 为单元的 next token 预测模式,与人类认知推理的抽象思维、因果推断等能力相差甚远。DeepSeek-R1 代表的强化推理能力,未来有望让大模型能够在合适的事实性和逻辑性约束的奖励函数限制下,通过自我反思和思维链回溯,大大减少自我幻觉的产生,更好地提升常识推理能力。

除了常识推理,研究系统 1+系统 2 的神经符号架构,正在逐步

演进到逻辑性更严、思维深度更高的专用领域逻辑推理,典型的领域包括数学推理和科学推理等。

数学推理是目前 LLMs 研究者都高度关注的领域,充分体现了神经符号融合的技术内涵。DeepSeek 就有专门针对数学的模型版本:DeepSeekMath 7B^[6]和 DeepSeek-Prover-V1.5^[7]。DeepSeekMath 7B 是在 2024 年春节发布的,对标 GPT-4^[8] 数学能力的领域 LLMs,只基于网上收集的数学的语料数据进行监督微调 and 强化学习训练,其性能达到 51.7% 的 math 测试基准水平,接近当时的 Gemini-Ultra^[9]与 GPT-4。其比较大的创新点是 DeepSeek 提出的组相对策略优化(group relative policy optimization, GRPO)强化学习算法,取代经典的近端策略优化(Proximal Policy Optimization, PPO)方法^[10],在优化了内存开销的同时,增强了数学推理能力。DeepSeek-Prover-V1.5 则是在 2024 年 8 月发布的,针对数据推理要求更高、推理链条更长的数学证明类问题。这类问题通常要结合 LLMs 的答案生成能力与专用数学证明求解器的符号推理能力,实现自然语言与数学形式化表述之间的融会贯通,把 LLMs 的思维链引导、证明思路的高效搜索、证明步骤的快速生成等集成到推理系统。DeepSeek-Prover 把自然语言的逻辑思维与数学证明器 Lean^[11]结合,基于 DeepSeekMath 的基本数学能力,在证明数据集上进一步微调,并基于 Lean 的逻辑验证功能定义奖励函数,以 GRPO 对模型进行强化学习。在证明推理过程

中, DeepSeek-Prover 又引入了基于蒙特卡洛树的搜索方法, 并深入优化了证明步骤生成—中间结果验证—回溯修正的自动化证明过程。

前文论述的 DeepSeek 和其他推理系统所构造出的神经符号系统, 为未来实现跨领域、超长链条的科学推理提供了很有价值的探索思路。科学推理是科学家从已有的理论知识和模型假说中, 结合新的实验观测数据, 综合运用归纳、演绎、溯因推理方法, 推导科学结论, 形成新观点和新假说的过程。经典的科学理论和建模强调 Occam's Razor(奥卡姆剃刀)法则, 通过构建简约模型来刻画复杂系统, 以便抓住系统的本质。但是正如著名物理学家理查德·费曼提出的“费曼极限”所指明的那样, 使用简单的数学模型来描述复杂系统时会面临局限性。因为这些精简模型尽管能有效提供近似解, 并让我们理解和计算复杂现象, 但是它们无法完全捕捉到所有细节。未来, 基于类似 DeepSeek-R1 的推理技术, 可以从多个方面赋能跨领域的科学推理, 从而突破费曼极限。首先, 可以在海量科技文献的基础上, 实现最新学科知识的深度整合和挖掘, 形成对复杂科学领域探索的知识基础。其次, 可以在海量科学数据的基础上, 对这些数据蕴含的内在规律和模式进行分析, 归纳总结出本质的科学规律, 建立融合神经网络和数理机理符号的新型科学模型, 来刻画和表征复杂系统的关键因素, 理解复杂系统辩论之间的深层次关系, 突破传统模型无法领悟的因果链条。

3 通用人工智能的未来趋势

DeepSeek 的成功开启了通用人工智能深度探索的新篇章, 特别是对大语言模型、多模态大模型和具身智能大模型的发展带来新的动力。大语言模型的能力仍呈现不断提升的趋势, DeepSeek 的强大推理能力将使得它更好地理解语言背后的语境、隐含的逻辑和因果关系, 从而胜任更为专业的工作场景。多模态大模型则进一步将视觉、听觉、触觉等多种感知数据融合到大模型中, 使其不仅在更多维度上进行学习和推理, 而且能够更好地生成丰富的多媒体内容, 更自如地与人类和周围环境进行互动。从多模态大模型出发, 把大模型和机器人本体结合起来, 构造具身智能体, 使得数字世界的大模型拥有感知、认知、决策、运动等综合能力, 能够进入三维的物理环境中, 实现 AI 系统与物理世界全方位的交互。

为推动人工智能技术的持续演进, 亟需深入探讨以下 3 个关键问题: 首先, DeepSeek 所采用的开源路线将在通用人工智能的发展中扮演何种角色? 其次, 在迈向通用人工智能的过程中, Neural Scaling Law^[12]是否仍具备其理论指导意义? 最后, 如何通过“通专结合”的方式构建行业大模型, 从而充分发挥大模型技术在产业应用中的价值?

3.1 开源开放生态对发展通用人工智能的作用

DeepSeek 的成功本质上体现了开源路线在推动生成式人工智能技术进步中的重要作用。自

OpenAI 发布 ChatGPT 以来, 大模型领域逐渐分化为 2 大阵营: 其一是以 OpenAI 为代表的闭源大模型生态, 该生态拒绝公开模型权重、相关代码以及关键技术细节; 其二是以 LLaMA^[13]为首的开源大模型生态, 模型的权重和部分(或全部)的代码开源, 并公布非常详细的技术报告。开源大模型生态虽然帮助 AI 领域的研究者和开发者能够更好、更快地进行技术创新和应用领域适配, 降低了生成式人工智能技术应用的门槛, 但是其技术水平一直无法达到 OpenAI 大模型的性能, 使得 OpenAI 形成了所谓的技术“护城河”。所以当 LLMs 推理技术取得突破之时, OpenAI 的 o1/o3 模型不公布任何技术细节, 同时给访问调用其功能设置了比较高的收费。而 DeepSeek-R1 系统超过了 o1 模型的性能, 无疑给开源大模型社区提供了全新的技术选择, 打破了 OpenAI 的技术垄断, 让每个研究者和开发者都能站在这个全新的起点上, 以开源进化的技术演进模式, 持续加速探索 LLMs 推理的新思路和新创新。

DeepSeek 的开源实践或许印证了发展通用人工智能的必然规律, 必须通过开放的技术创新生态, 打造开放的通用人工智能系统, 实现开源创新驱动、模型滥用风险防范、可持续商业模式之间的微妙平衡。DeepSeek 选择有限度开放模型权重和推理代码, 而不是完全开源模型架构和训练代码, 可以有效地破解大模型创新与伦理安全之间的悖论。这样既保持了开源 LLMs 技术前进的动力, 又

为广大研究者做更深入的安全对齐和评测评价提供了可能。

3.2 Neural Scaling Law 对通用人工智能发展的影响

在大模型研发中,通常认为模型的性能提升和模型参数规模之间满足幂律关系,也就是 Neural Scaling Law,即模型规模越大,输入的训练数据越多,模型预测能力就越强。因此,有人认为只要继续扩大模型规模,就能够在不久的将来实现通用人工智能。但这种指数级增长的算力需求,给智算集群系统带来了极大的开销,使得规模的可持续扩展遇到瓶颈。例如:最大模型 GPT-4 的参数规模已到万亿级别,构建和运行如此大规模的智算集群,需要克服集群计算通信开销、供电和散热等能耗难题。

此外,高质量和高密度数据语料库的稀缺性也成为制约模型规模扩展的关键因素。大模型的性能提升依赖于大量高质量训练语料,而现有语料库主要依赖于互联网公共领域数据的汇集。有研究预计到 2028 年大模型训练将耗尽所有互联网公共数据资源,大模型增长不可避免地遭遇数据危机^[14]。为此,需要面向垂直领域,深度挖掘私域数据,扩大高质量数据共享,以支撑大模型更好地适应于垂直领域的需求。

综上所述,单纯依赖模型规模的扩展来实现通用人工智能,无论在技术可行性还是经济成本方面,均难以构成可持续的技术路径。当前预训练方法已接近其极限,进一步扩大模型规模和数据量的边际效益正在递减。因此,研究重点正在从单纯的预训练扩展转向更

高效、更精细化的技术路径。

目前,业界的研究重点正在转向后训练(Post-Training)^[15]和测试时间缩放(Test-Time Scaling, TTS) 2 个方面。1) 监督微调和强化学习等 Post-Training 方式成为优化模型性能的关键手段。通过监督微调,模型可以在特定任务上进行精细化调整。例如,DeepSeek-R1 通过结合监督微调和强化学习,避免了传统 PRM 技术在细粒度步骤定义和中间步骤判断上的局限性,显著提升了模型的推理能力和任务适应性。2) TTS 是一种在推理阶段通过增加计算资源或时间来提升大模型性能的技术。通过动态搜索和优化技术(如 Beam Search、Sampling、Best-of-N weighted 等),模型可以在生成过程中实时调整输出,以提高结果的准确性和一致性。

未来的研究将更加注重模型的高效性和可持续性,而非单纯追求规模的扩展。通过结合 Post-Training 优化和 Test-Time Scaling 技术,生成式人工智能可以在有限的资源和数据条件下,实现更高效的性能提升。同时,垂直领域的数据挖掘和私域数据的利用也将成为支撑大模型发展的重要方向。

3.3 “通专结合”的行业模型

随着大模型训练范式的改变,尤其是推理和 Post-Training、Test-Time Scaling 逐渐成为发展的热点。在垂直领域走“通专结合”的技术路线成为必然,需要引入模块化的架构假设、强化式的能力提升,支持大模型与业务逻辑的紧密结合,在产业领域实现广泛落地与价值赋能。

首先,巨无霸式的通用模型必然给企业带来维护升级、训练成本等一系列的复杂性难题,必须引入模块化设计的理念,实现模型结构和业务功能的松耦合架构。在明确问题边界、配置所需算力的前提下,将大模型的专家模块按照功能业务场景进行划分,其中比较大规模的模块负责一些基础功能,具体业务领域设置不同的中小规模模型,并通过智能的语义路由把各模块串接在一起。在这种设计中,由于各模块之间采取稀疏性的链接方式,其训练代价、管理成本都会显著下降,这本质上是通过群体化机制解决复杂领域问题。

其次,业务场景需要对基座模型进行定向的蒸馏与微调,以提升其专业能力。采用 DeepSeek 这类开放的基座模型,并使用强化学习方式,将能够以自动化的方式持续提升模型在解决专业问题上的表现,涌现垂直领域所需的专业技能。例如:医疗、教育、金融等场景已经积累了各类行业数据和知识图谱等资源,将成为强化式的专业能力提升的基础,使得开发者基于 DeepSeek 的基座创造出更多衍生模型,充分解决行业模型落地问题。

自 DeepSeek 大模型发布以来,已在多个行业,尤其是保险领域,取得了显著的应用成效。截至 2025 年 2 月 18 日,已有多家领先保险公司接入 DeepSeek 并落地应用。例如,新华人寿保险股份有限公司在其“新华 e 家”App 中成功接入 DeepSeek-R1 和 DeepSeek-V3,推动了个人 AI 助理、销售支持、风控合规管理等多个场景的

智能化,提升了业务效率。北大方正人寿保险有限公司通过上线“方灵”智能展业助手,利用 DeepSeek 的知识智能检索功能,为代理人提供保险基础概念、法律法规等内容的快速解答,大幅提升了代理人的工作效率,并通过 AI 智能陪练和核保助手降低了培训和核保成本。东吴人寿保险股份有限公司借助腾讯云大模型引擎全面部署 DeepSeek-R1,优化了条款解析,提升准确率 40%,并建立了健康管理知识图谱,助力个性化保障方案的设计。

这种“通专结合”的模式不仅促进了大模型的落地应用,更推动了行业智能化的全面升级。

4 结论

DeepSeek 的推出标志着中国在通用人工智能领域迈出了重要的一步,开启了基于国产技术的人工智能创新生态新篇章。通过在大语言模型、推理技术等方面的创新,DeepSeek 不仅展示了其在多领域对话、内容生成以及深度逻辑推理方面的强大能力,还为未来通用人工智能的发展提供了新的思路 and 方向。DeepSeek-V3 和 R1 系列模型在低成本高性能训练、推理能力提升等方面的突破,展示了中国在人工智能技术上的自主创新能力。

DeepSeek 的成功不仅为开源生态注入了新的活力,也为行业模型的“通专结合”提供了可行的路径。随着模型规模的扩展和推理技术的不断优化,DeepSeek 有望在更多垂直领域实现广泛应用,推

动人工智能技术在医疗、教育、金融等行业的深度落地。同时,DeepSeek 的开源策略和技术创新生态,将为中国乃至全球的人工智能研究者提供更多的合作与创新机会,进一步加速通用人工智能的发展。

总的来说,DeepSeek 不仅是中国人工智能技术发展的里程碑,更是全球人工智能领域的重要贡献者。随着 AI 技术的不断进步和应用的深入,像 DeepSeek 一样的中国自主 AI 研究力量有望在更多领域引领原创技术突破,推动人工智能迈向新的高度,为人类社会可持续发展带来更多的创新与变革。

参考文献 (References)

- [1] Liu A, Feng B, Xue B, et al. DeepSeek-V3 technical report[J]. Computer Science, 2024, doi: 10.48550/arXiv.2412.19437.
- [2] DeepSeek-AI, Guo D Y, Yang D J, et al. DeepSeek-R1: Incentivizing reasoning capability in LLMs via reinforcement learning[J]. Computer Science, 2025, doi: 10.48550/arXiv.2501.12948.
- [3] Jaech A, Kalai A, Lerer A, et al. Openai o1 system card[J]. Computer Science, 2024, doi: 10.48550/arXiv.2412.16720.
- [4] Vaswani A. Attention is all you need[J]. Advances in Neural Information Processing Systems, 2017, doi: 10.48550/arXiv.1706.03762.
- [5] Huang L, Yu W J, Ma W T, et al. A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions[J]. ACM Transactions on Information Systems, 2025, 43(2): 1–55.
- [6] Shao Z H, Wang P Y, Zhu Q H, et al. DeepSeekMath: Pushing the limits of mathematical reasoning in open language models[J]. Computer Science, 2024, doi: 10.48550/arXiv.2402.03300.
- [7] Xin H J, Ren Z Z, Song J X, et al. DeepSeek-prover-V_{1.5}: Harnessing proof assistant feedback for reinforcement learning and Monte-Carlo tree search[J]. Computer Science, 2024, doi: 10.48550/arXiv.2408.08152.
- [8] Achiam J, Adler S, Agarwal S, et al. Gpt-4 technical report[J]. Computer Science, 2023, doi: 10.48550/arXiv.2303.08774.
- [9] Team G, Anil R, Borgeaud S, et al. Gemini: A family of highly capable multimodal models[J]. Computer Science, 2023, doi: 10.48550/arXiv.2312.11805.
- [10] Schulman J, Wolski F, Dhariwal P, et al. Proximal policy optimization algorithms[J]. Computer Science, 2017, doi: 10.48550/arXiv.1707.06347.
- [11] de Moura L, Ullrich S. The lean 4 theorem prover and programming language[C]//Automated Deduction – CADE 28: 28th International Conference on Automated Deduction, Virtual Event. Cham: Springer International Publishing, 2021: 625–635.
- [12] Kaplan J, McCandlish S, Henighan T, et al. Scaling laws for neural language models[J]. Computer Science, 2020, doi: 10.48550/arXiv.2001.08361.
- [13] Touvron H, Lavril T, Izacard G, et al. LLaMA: Open and efficient foundation language models[J]. Computer Science, 2023, doi: 10.48550/arXiv.2302.13971.
- [14] Jones N. The AI revolution is running out of data. What can researchers do?[J]. Nature, 2024, 636(8042): 290–292.
- [15] Liu Z H, Wang Y H, Han K, et al. Post-training quantization for vision transformer[J]. Advances in Neural Information Processing Systems, 2021, 34: 28092–28103.

DeepSeek: Technological innovations and development trends toward artificial general intelligence

WU Wenjun^{1,2}, LIAO Xingchuang¹, ZHAO Jinkun¹

1. State Key Laboratory of Complex and Critical Software Environment, Beihang University, Beijing 100191, China

2. International Innovation Institute, Beihang University, Hangzhou 311115, China

Abstract In this paper, the latest advancements of DeepSeek in Artificial General Intelligence (AGI) were reviewed, with a focus on innovations in large language models and reasoning technologies. Novel model structures and algorithmic designs were introduced in the newly developed DeepSeek-V3 architecture, , to implement a comprehensive optimization methodology which significantly enhanced training efficiency under constrained computing resources. In the reasoning system of DeepSeek-R1 framework, reinforcement learning (RL) was integrated with supervised fine-tuning (SFT) in an innovative way and breakthroughs in reasoning depth and logical coherence were achieved. Based on the innovations of DeepSeek, in the study, three critical challenges in AGI development were discussed: (1) the strategic role of open-source ecosystems in accelerating AGI progress, (2) the validity of neural scaling laws in the era of foundation models, and (3) practical approaches to implement industry-specific models through synergistic integration of general and domain-specific capabilities, demonstrated via DeepSeek-based implementations.

Keywords DeepSeek; reinforcement learning; MoE; Scaling Law; slow thinking ●



(责任编辑 王丽娜)