

研究所

分析师: 肖承志
SAC 登记编号: S1340524090001
Email: xiaochengzhi@cnpsec.com
研究助理: 冯昱文
SAC 登记编号: S1340124100011
Email: fengyuwen@cnpsec.com

近期研究报告

《基金 Q1 加仓有色汽车传媒，减仓电新食品饮料——公募基金 2025Q1 季报点评》 - 2025. 04. 30

《泛消费打开连板与涨幅高度，ETF 资金平铺机器人、人工智能与芯片——行业轮动周报 20250427》 - 2025. 04. 28

《国家队交易特征显著，短期指数仍交易补缺预期，TMT 类题材仍需等待——行业轮动周报 20250420》 - 2025. 04. 21

《小市值持续，高低波风格交替——中邮因子周报 20250413》 - 2025. 04. 14

《4 月是否还会有“最后一跌”？——微盘股指数周报 20250406》 - 2025. 04. 07

《“924”以来融资资金防守后均见到行情低点，仍关注科技配置机会——行业轮动周报 20250330》 - 2025. 03. 31

《英伟达召开 GTC 2025 大会，Skywork-R1V、混元 T1 等推理模型接连上线——AI 动态汇总 20250324》 - 2025. 03. 25

《反转效应强势，GRU 模型新高——中邮因子周报 20250323》 - 2025. 03. 24

《微盘领涨创下历史新高，4 月临近仍有调整压力——微盘股指数周报 20250316》 - 2025. 03. 17

《小市值强势，动量风格依旧——中邮因子周报 20250309》 - 2025. 03. 10

金工周报

通义千问发布 Qwen-3 模型，DeepSeek 发布数理证明大模型——AI 动态汇总 20250505

● 通义千问发布 Qwen-3 模型

4 月 29 日，通义千问发布 Qwen-3 模型，可以看到本次发布的旗舰模型为 Qwen3-235B-A22B，在诸多基准测试中表现超越 DeepSeek-R1、o1、o3-mini、Grok-3、Gemini-2.5-Pro 等模型。除此之外，小型模型 Qwen3-30B-A3B 的激活参数数量是 QwQ-32B 的 10%，表现更胜一筹，甚至像 Qwen3-4B 这样的小模型也能匹敌 Qwen2.5-72B-Instruct 的性能。

● DeepSeek 发布数理证明大模型 DeepSeek-Prover-V2-671B

4 月 30 日，DeepSeek 发布了应用于数理问题形式化定理证明场景的大模型 DeepSeek-Prover-V2-671B。所谓形式化定理证明是指在数理逻辑中，不以自然语言书写而是以形式化语言书写的论证：这种语言包含了由一个给定的字母表中的字符所构成的字符串。而证明则是一种由该些字符串组成的有限长度的序列。这种定义使得人们可以谈论严格意义上的“证明”，而不涉及任何逻辑上的模糊之处。研究证明的形式化和公理化的理论称为证明论。

● Grok 3.5 将于本周推出

4 月 29 日，马斯克在社交平台 X 上表示，Grok 3.5 早期测试版将于本周向 SuperGrok 订阅者发布，同时，马斯克称“Grok 3.5 是第一个能够准确回答有关火箭发动机或电化学技术问题的人工智能”且“Grok 是从第一性原理推理并得出互联网上根本不存在的答案”。

● 小米开源推理大模型 MiMo-7B-RL

4 月 30 日，小米开源其首个推理大模型 Xiaomi MiMo，该模型发布版本为 MiMo-7B-RL，虽然模型参数量不大，但在数学推理（AIME 24-25）和代码竞赛（LiveCodeBench v5）基准上超过了 o1-mini 模型和通义千问的 QwQ-32B-Preview 模型。目前模型已开源。

● 风险提示：

以上内容基于历史数据完成，在政策、市场环境发生变化时存在失效的风险；历史信息不代表未来。

目录

1	AI 重点要闻	4
1.1	通义千问发布 Qwen-3 模型	4
1.2	DeepSeek 发布数理证明大模型 DeepSeek-Prover-V2-671B	6
1.3	Grok 3.5 将于本周推出	8
1.4	小米开源推理大模型 MiMo-7B-RL	9
2	企业动态	11
2.1	亚马逊发布 Amazon Nova Premier 模型	11
2.2	Meta 发布 ReasonIR-8B 模型	13
2.3	FutureHouse 发布科学家智能体	15
3	AI 行业洞察	16
3.1	智能数据标注产业发展观察报告：传统模式迈向智能化新业态	16
3.2	InfiGUI-R1：利用强化学习，让 GUI 智能体学会规划任务、反思错误	17
4	技术前沿	19
4.1	Paper2Code：科研报告代码复现智能体	19
5	风险提示	23

图表目录

图表 1: Qwen-3-235B-A22B 测评	4
图表 2: Qwen-3-30B-A3B 测评	4
图表 3: Qwen-3 混合思考模式	5
图表 4: Qwen-3 训练框架	6
图表 5: DeepSeek-Prover-V2 模型表现	7
图表 6: DeepSeek-Prover-V2 冷启动	8
图表 7: 宣布 Grok-3.5 即将推出	9
图表 8: MiMo-7B 模型测评	10
图表 9: MiMo-7B 训练架构	10
图表 10: Seamless Rollout 系统	11
图表 11: Nova Premier 对比	12
图表 12: RAG 检索性能不足	14
图表 13: ReasonIR-8B 训练数据架构	14
图表 14: LiteQA 基准测试	15
图表 15: Actor2Reasoner 架构	17
图表 16: Paper2Code	19
图表 17: Paper2Code	20
图表 18: 第一阶段: 规划阶段	21
图表 19: 第二阶段: 分析阶段	22
图表 20: 第三阶段: 编码阶段	22

1 AI 重点要闻

1.1 通义千问发布 Qwen-3 模型

4 月 29 日，通义千问发布 Qwen-3 模型，可以看到本次发布的旗舰模型为 Qwen3-235B-A22B，在诸多基准测试中表现超越 DeepSeek-R1、o1、o3-mini、Grok-3、Gemini-2.5-Pro 等模型。除此之外，小型模型 Qwen3-30B-A3B 的激活参数数量是 QwQ-32B 的 10%，表现更胜一筹，甚至像 Qwen3-4B 这样的小模型也能匹敌 Qwen2.5-72B-Instruct 的性能。

图表 1：Qwen-3-235B-A22B 测评

	Qwen3-235B-A22B Total	Qwen3-32B Dense	OpenAI o1 2025-03-17	Deepseek-R1	Grok 3 Beta Total	Gemini 2.5 Pro	OpenAI o3-mini Instruct
ArenaHard	95.6	93.8	92.1	93.2	-	96.4	89.0
AIME'24	85.7	81.4	74.3	79.8	83.9	92.0	79.6
AIME'25	81.5	72.9	79.2	70.0	77.3	86.7	74.8
LiveCodeBench v5, 2024-10-2025-02	70.7	65.7	63.9	64.3	70.6	70.4	66.3
CodeForces Elo Rating	2056	1977	1891	2029	-	2001	2036
Aider PowerB2	61.8	50.2	61.7	56.9	53.3	72.9	53.8
LiveBench 2024-11-03	77.1	74.9	75.7	71.6	-	82.4	70.0
BFCL v2	70.8	70.3	67.8	56.9	-	62.9	64.6
MultitF 8 categories	71.9	73.0	48.8	67.7	-	77.8	48.4

1. ArenaHard: The sample 16 times for each query and report the average of the accuracy. ArenaHard consists of Part I and Part II, with a total of 36 questions.
2. AIME: The AIME is a series of math problems that test the ability of AI models to solve complex mathematical problems.
3. BFCL: The BFCL is a series of math problems that test the ability of AI models to solve complex mathematical problems.

图表 2：Qwen-3-30B-A3B 测评

	Qwen3-30B-A3B Total	QwQ-32B	Qwen3-4B Dense	Qwen2.5-72B-Instruct	Gemini 2.5 Pro	DeepSeek-V3	GPT-4o 2024-11-20
ArenaHard	91.0	89.5	76.6	81.2	86.8	85.5	85.3
AIME'24	80.4	79.5	73.8	18.9	32.6	39.2	11.1
AIME'25	70.9	69.5	65.6	15.0	24.0	28.8	7.6
LiveCodeBench v5, 2024-10-2025-02	62.6	62.7	54.2	30.7	26.9	33.1	32.7
CodeForces Elo Rating	1974	1982	1671	859	1063	1134	864
GPQA	65.8	65.6	55.9	49.0	42.4	59.1	46.0
LiveBench 2024-11-03	74.3	72.0	63.6	51.4	49.2	60.5	52.2
BFCL v2	69.1	66.4	65.9	63.4	59.1	57.6	72.5
MultitF 8 categories	72.2	68.3	66.3	65.3	69.8	55.6	65.6

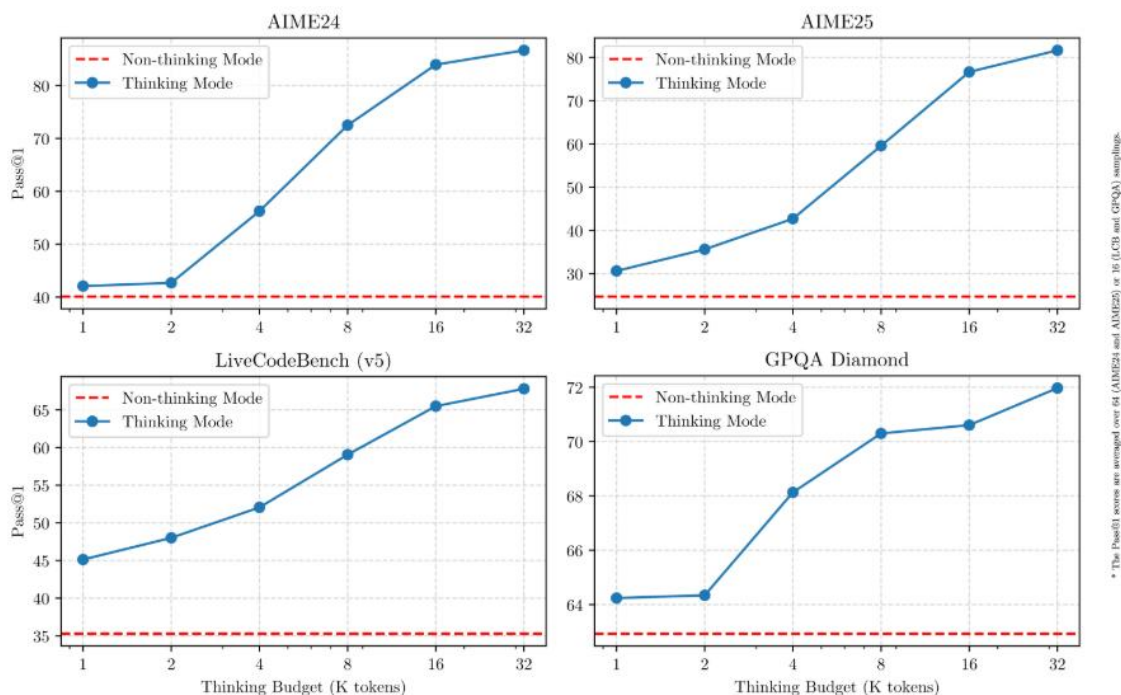
1. ArenaHard: The sample 16 times for each query and report the average of the accuracy. ArenaHard consists of Part I and Part II, with a total of 36 questions.
2. AIME: The AIME is a series of math problems that test the ability of AI models to solve complex mathematical problems.
3. BFCL: The BFCL is a series of math problems that test the ability of AI models to solve complex mathematical problems.

资料来源：Qwen，中邮证券研究所

资料来源：Qwen，中邮证券研究所

本次开源了两个 MoE 模型：Qwen3-235B-A22B 以及 Qwen3-30B-A3B，其中 A22B 和 A3B 指的是最小激活参数为 220 亿和 30 亿。此外，Qwen3 的六个 Dense 模型也已开源，包括 Qwen3-32B、Qwen3-14B、Qwen3-8B、Qwen3-4B、Qwen3-1.7B 和 Qwen3-0.6B，均在 Apache2.0 许可下开源。

从模型思考模式看，Qwen-3 支持多种思考模式，用户可根据任务需求选择，以平衡成本效益和推理质量。这是继 Claude 3.7 Sonnet、Gemini 2.5 后全球第三个混合推理模型，也是国内首款混合推理模型。

图表 3：Qwen-3 混合思考模式


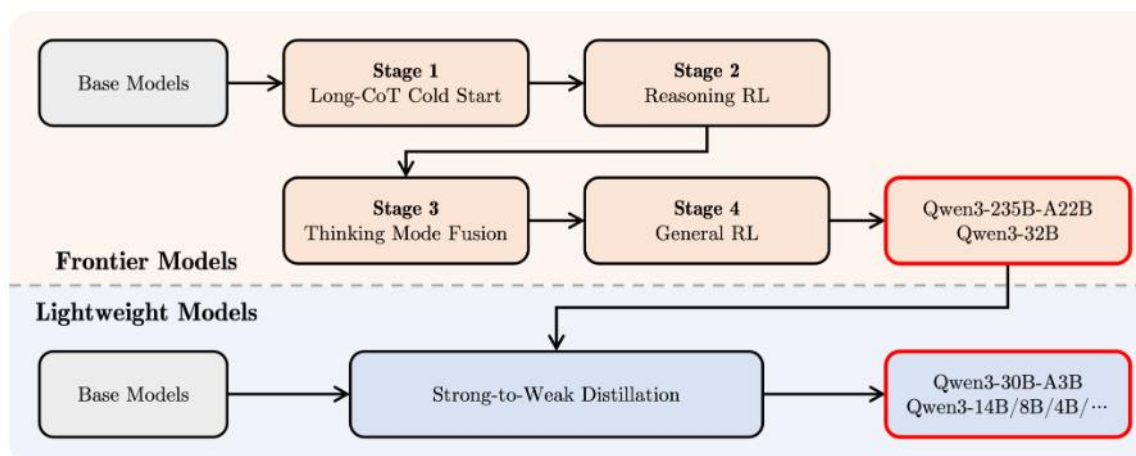
资料来源：Qwen，中邮证券研究所

在智能体开发方面，Qwen-3 模型针对 Agent 协同与代码能力进行了优化，同时也加强了对 MCP 的支持。Qwen-3 全面适配 AI Agent 生态。其架构不再局限于“问答能力”，而是面向 Agent 架构优化的执行效率、响应结构与工具泛化能力，做到了从底座模型到执行智能体的开发链路优化。

训练方面，Qwen3 的数据集相比 Qwen2.5 有了显著扩展，Qwen-3 训练使用的数据量达到了 36 万亿个 token，涵盖了 119 种语言和方言。整个预训练过程分为三个阶段。在第一阶段（S1），模型在超过 30 万亿个 token 上进行了预训练，上下文长度为 4K token。这一阶段为模型提供了基本的语言技能和通用知识。在第二阶段（S2），通过增加知识密集型数据（如 STEM、编程和推理任务）的比例来改进数据集，随后模型又在额外的 5 万亿个 token 上进行了预训练。在最后阶段，使用高质量的长上下文数据将上下文长度扩展到 32K token，确保模型能够有效地处理更长的输入。而由于 Qwen-3 是兼顾思考推理和快速响应的混合推理模型，因此在训练中也对模型进行了后训练。整个训练流程包括（1）长思维

链冷启动，(2) 长思维链强化学习，(3) 思维模式融合，以及 (4) 通用强化学习。

图表 4: Qwen-3 训练框架

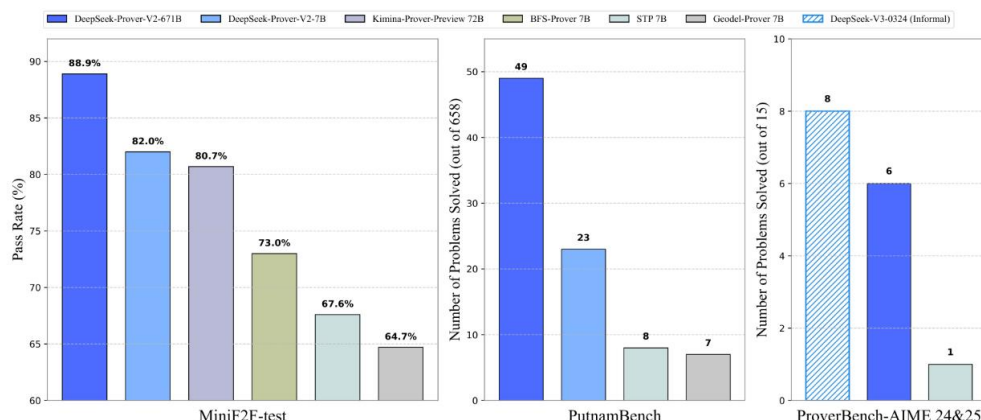


资料来源: Qwen, 中邮证券研究所

1.2 DeepSeek 发布数理证明大模型 DeepSeek-Prover-V2-671B

4月30日, DeepSeek 发布了应用于数理问题形式化定理证明场景的大模型 DeepSeek-Prover-V2-671B。所谓形式化定理证明是指在数理逻辑中, 不以自然语言书写而是以形式化语言书写的论证: 这种语言包含了由一个给定的字母表中的字符所构成的字符串。而证明则是一种由该些字符串组成的有限长度的序列。这种定义使得人们可以谈论严格意义上的“证明”, 而不涉及任何逻辑上的模糊之处。研究证明的形式化和公理化的理论称为证明论。

DeepSeek 早 在 之 前 就 基 于 DeepSeek-Math 模 型 优 化 出 DeepSeek-Prover-V1.5 模型, 而这次的 DeepSeek-Prover-V2 模型是以 DeepSeek V3 作为基础模型来进一步微调的。从结果上看, 该模型数理证明能力好于对标模型。

图表 5: DeepSeek-Prover-V2 模型表现


资料来源：DeepSeek，中邮证券研究所

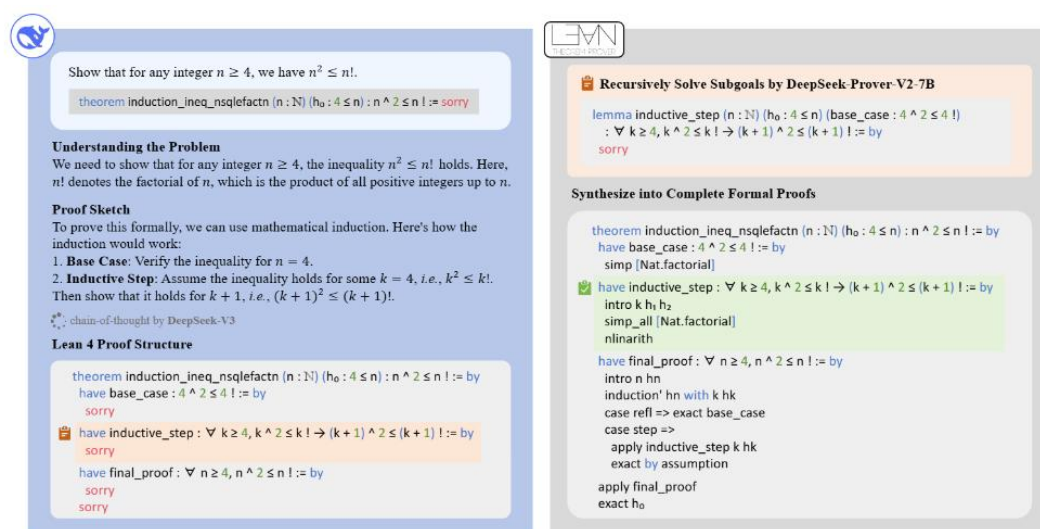
为了解决自然语言推理依赖启发式和数据驱动模式，缺乏形式化所需的严谨性；现有方法（如树搜索、整体证明生成）在复杂问题上效率低下的问题，整个 DeepSeek-Prover-V2 模型的训练思路是模仿人类数学家，构建递归子目标分解框架，复杂定理分解为可管理的子目标（Lemma），整个过程的实现步骤为：

- 子目标生成**：利用 DeepSeek-V3 生成自然语言证明草图，并同步将其形式化为 Lean 代码。
- 递归求解**：通过 7B 参数的小型证明器递归解决子目标，减少计算开销。
- 合成完整证明**：将子目标证明串联为完整形式化证明，与 DeepSeek-V3 的思考链结合生成冷启动数据。

在训练过程中，使用了课程学习与强化学习来增强大模型的推理解题能力。课程学习利用分解的子目标生成合成问题，逐步增加训练难度，引导模型学习复杂推理。强化学习则沿用了 V3 模型的 GRPO 算法，并且设置了一致性奖励，用于惩罚证明结构与子目标分解的偏差，确保逻辑连贯性。

在资源利用方面，DeepSeek-Prover-V2 模型引入双模式推理架构，对于简单的问题，采用非链式思考（Non-CoT）模式，快速生成简洁证明代码；对于复杂问题，采用链式思考（CoT）模式，详细分解中间步骤，提升复杂问题的证明准确性。

图表 6: DeepSeek-Prover-V2 冷启动



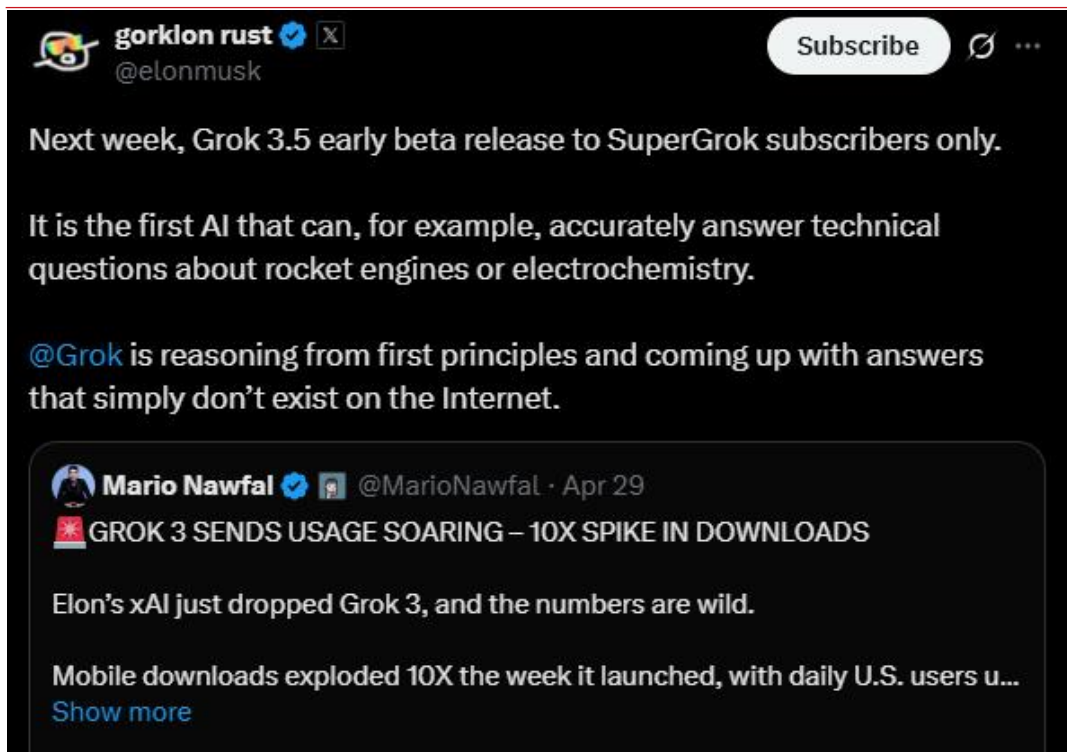
资料来源: DeepSeek, 中邮证券研究所

在整个训练中, 关键难点在于冷启动数据生成, 通过 DeepSeek-V3 生成证明草图→7B 模型递归求解子目标→合成完整证明→与思考链结合形成训练数据。如此处理可以结合非形式化推理与形式化验证, 缓解标注数据稀缺问题。从结果看, 该模型成功证明 AMC 竞赛题 (如 2024 年 AIME 第 15 题) 和高等数学问题 (如拓扑学定理)。

1.3 Grok 3.5 将于本周推出

4 月 29 日, 马斯克在社交平台 X 上表示, Grok 3.5 早期测试版将将于本周向 SuperGrok 订阅者发布, 同时, 马斯克称 “Grok 3.5 是第一个能够准确回答有关火箭发动机或电化学技术问题的人工智能” 且 “Grok 是从第一性原理推理并得出互联网上根本不存在的答案”。

图表 7：宣布 Grok-3.5 即将推出



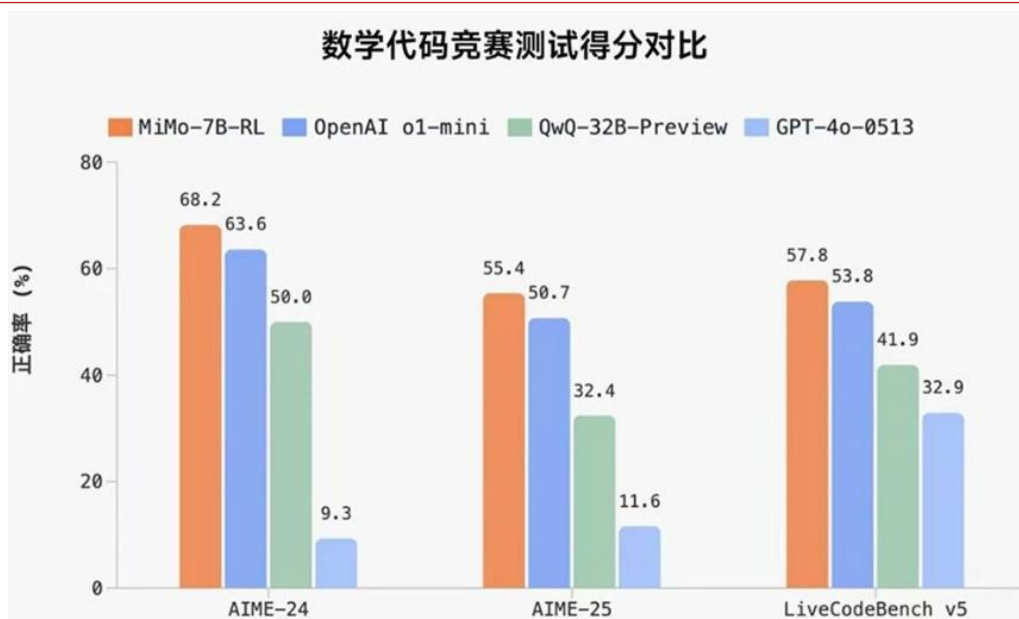
资料来源：X，中邮证券研究所

QuestBench 测试了包括 GPT-4o、Claude 3.5 Sonnet、Gemini 2.0 Flash Thinking Experimental 等领先模型，覆盖零样本、思维链和四样本设置。测试于 2024 年 6 月至 2025 年 3 月间进行，涉及 288 个 GSM-Q 和 151 个 GSME-Q 任务。

1.4 小米开源推理大模型 MiMo-7B-RL

4 月 30 日，小米开源其首个推理大模型 Xiaomi MiMo，该模型发布版本为 MiMo-7B-RL，虽然模型参数量不大，但在数学推理（AIME 24-25）和代码竞赛（LiveCodeBench v5）基准上超过了 o1-mini 模型和通义千问的 QwQ-32B-Preview 模型。目前模型已开源。

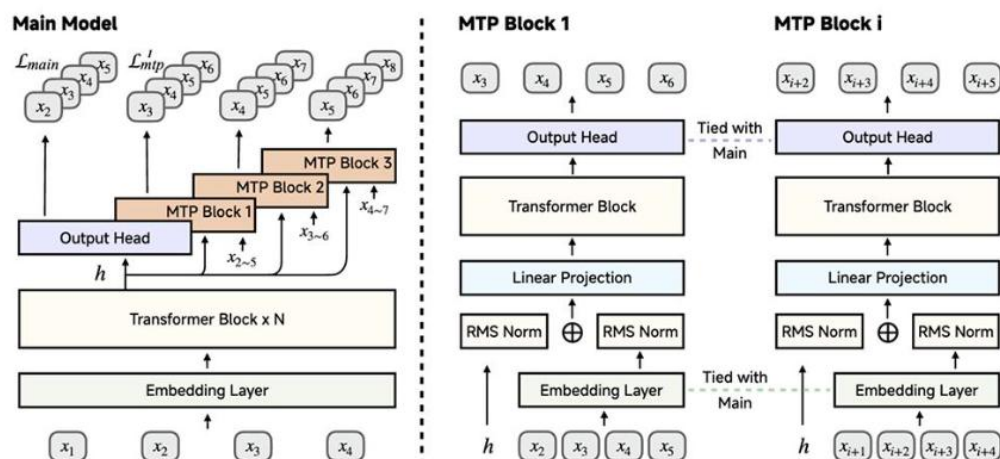
图表 8：MiMo-7B 模型测评



资料来源：Xiaomi，中邮证券研究所

MiMo-7B 模型的训练还是采取了预训练+后训练的架构，训练数据是该模型一大亮点：从网页、论文、代码、合成数据中提取了 25T tokens 的数据作为训练数据，着重挖掘富推理语料，并合成约 200B tokens 的推理数据。训练采用三阶段数据混合策略，逐步提升训练难度，MiMo-7B-Base 在约 25T tokens 上进行预训练；受 DeepSeek-V3 启发，将多 token 预测作为额外的训练目标，以增强模型性能并加速推理。

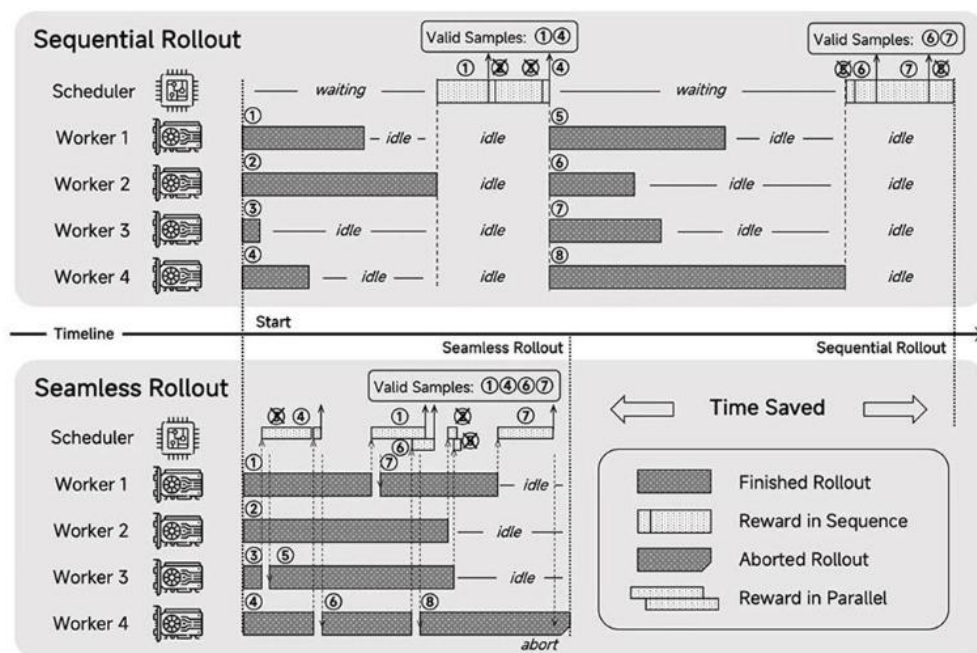
图表 9：MiMo-7B 训练架构



资料来源：Xiaomi，中邮证券研究所

在后训练部分，精选了 13 万道数学和代码题作为强化学习训练数据，可供基于规则的验证器进行验证。设计了 Seamless Rollout 系统，集成了连续部署、异步奖励计算和提前终止功能，以最大限度地减少 GPU 空闲时间，使得强化学习训练加速 2.29 倍，验证加速 1.96 倍。

图表 10: Seamless Rollout 系统



资料来源: Xiaomi, 中邮证券研究所

2 企业动态

2.1 亚马逊发布 Amazon Nova Premier 模型

5 月 3 日，亚马逊宣布正式推出 Amazon Nova Premier，Amazon Nova Premier 是 Nova 系列功能最强大的模型，适用于处理复杂任务，并可作为模型蒸馏的教师模型。目前 Nova Premier 已经可以在 Amazon Bedrock 中使用。

与 Nova Lite 和 Pro 类似，Premier 可以处理输入的文本、图像和视频（不包括音频）。凭借其先进能力，Nova Premier 擅长处理需要深度理解上下文、多步骤规划以及跨多工具和数据来源精确执行的复杂任务。凭借 100 万个 token 的上下文长度，Nova Premier 能处理超长文档或大型代码库。从测评结果看，Nova Premier 是 Nova 家族中功能最强的模型。

图表 11：Nova Premier 对比

	Nova Pro	Nova Premier
Text intelligence	Undergraduate level knowledge MMLU	85.9% 87.4%
	Science GPQA Diamond	50.0% 57.1%
	High school math competition AIME 2025	5.3% 16.0%
	Math problem-solving MATH-500	76.6% 82.0%
	Coding BigCodeBench Hard	22.3% 28.1%
	Coding MBXP (5 languages)	65.9% 78.4%
	Instruction Following IFEval	92.1% 91.5%
Visual intelligence	Visual understanding MMMU	62.0% 68.0%
	Document understanding OCRBench-v2	53.7% 56.9%
	Chart understanding CharXiv (Descriptive/Reasoning)	70.5%/40.6% 84.6%/48.8%
	Long-form video language understanding EgoSchema	72.1% 73.8%
	Visual counting TallyQA	54.0% 61.5%
Agentic workflows	Retrieval-augmented generation SimpleQA (SerpAPI)	84.6% 86.3%
	Function calling BFCL (2025-04-25)	60.8% 63.7%
	Web agents ScreenSpot Web Text	88.3% 91.7%
	Web agents ScreenSpot Web Icon	76.7% 84.0%
	Agentic coding SWE-Bench Verified (internal agentic scaffold)	- 42.4%
The numbers for other models are a mix of self-reported and measured evaluations. Several results overlap within the 95% confidence interval.		

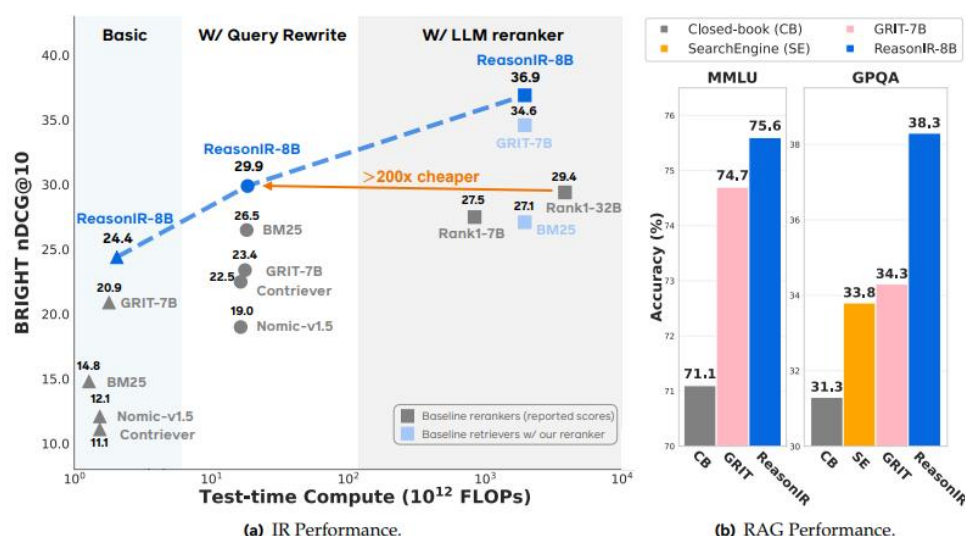
资料来源：AWS，中邮证券研究所

2.2 Meta 发布 ReasonIR-8B 模型

4 月 30 日，科技媒体 Marktechpost 报道，Meta AI 推出 ReasonIR-8B 模型，专为推理密集型检索设计，不仅在检索精度上取得突破，其低成本和高效率也使

其成为实际应用的理想选择。当前检索增强生成（RAG）系统在处理复杂推理任务时，常常因检索器性能不足而受限。传统检索器多针对简短事实性问题训练，擅长文档级别的词语或语义匹配，但面对长篇或跨领域查询时，难以整合分散知识，这种缺陷会导致错误信息传递，影响后续推理效果。ReasonIR-8B 模型直击这一痛点，基于 LLaMA3.1-8B 训练，结合创新数据生成工具 ReasonIR-SYNTHESIZER，构建模拟真实推理挑战的合成查询和文档对，更精准支持复杂任务。

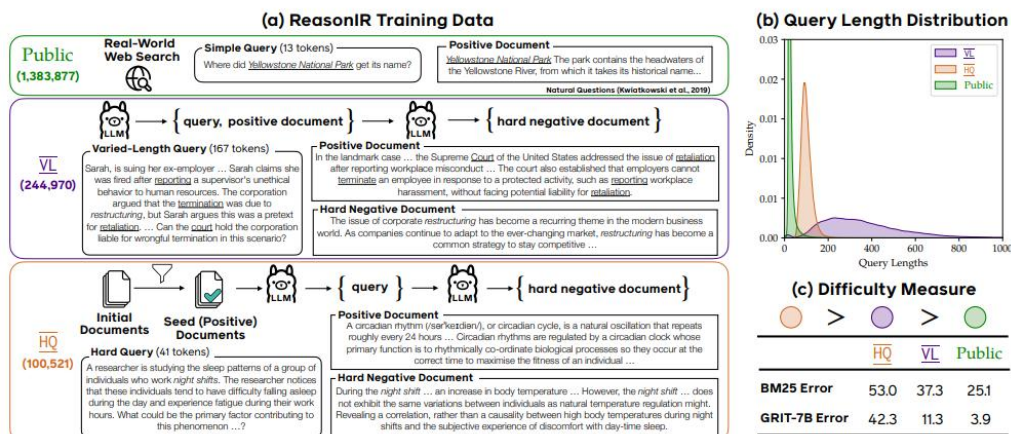
图表 12: RAG 检索性能不足



资料来源：Meta，IT 之家，中邮证券研究所

ReasonIR-8B 采用双编码器（bi-encoder）架构，将查询和文档独立编码为嵌入向量，通过余弦相似度评分。其训练数据包括长达 2000 个 token 的多样长度查询（VL Queries）和需逻辑推理的困难查询（HQ），有效提升模型处理长上下文和抽象问题的能力。Meta AI 目前已在 Hugging Face 上开源 ReasonIR-8B 模型、训练代码及合成数据工具，鼓励研究社区进一步探索多语言和多模态检索器的开发。

图表 13: ReasonIR-8B 训练数据架构



资料来源：Meta，IT 之家，中邮证券研究所

2.3 FutureHouse 发布科学家智能体

5 月 3 日，非盈利组织 FutureHouse 发布四个超人类的 AI 科学家智能体，分别承担科研全链条的关键角色：

Crow（乌鸦）：基于优化搜索算法实现精确文献检索，支持全文级科学文献访问，整合多源信息提取关键数据。

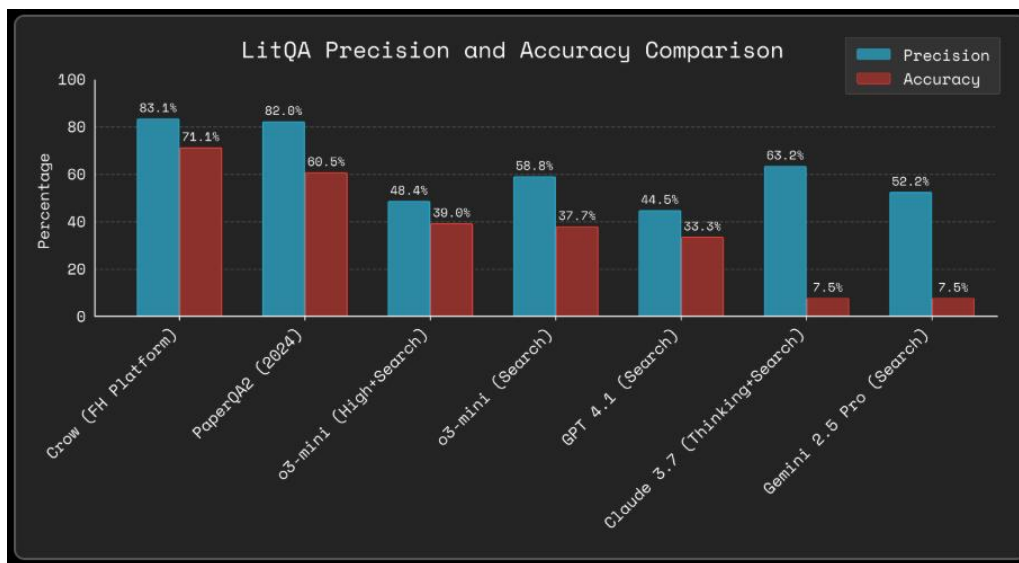
Falcon（猎鹰）：采用深度语义理解模型，完成深度文献综述与知识整合，基准测试中检索准确率达 95% 以上，超越 GPT-4.5 等主流模型。

Owl（猫头鹰）：聚焦科研历史脉络梳理，擅长先例搜索与研究差异分析，优化实验设计逻辑。

Phoenix（凤凰）：多模态药物研发引擎，可预测化学反应路径、筛选新分子结构，提升药物研发效率。

Crow、Falcon 和 Owl 通过了严格的基准测试，在搜索精度和准确性上已经超越了目前顶级搜索模型，比如 o3-mini，GPT-4.5，Claude-3.7。

图表 14：liteQA 基准测试



资料来源：FutureHouse，中邮证券研究所

3 AI 行业洞察

3.1 智能数据标注产业发展观察报告：传统模式迈向智能化新业态

根据最新发布的《智能数据标注产业发展观察报告》，传统的人力密集型标注模式正在加速向智能驱动和平台化的新业态转型。

产业转型的核心路径主要依托于技术驱动革新与平台化新业态。传统依赖人工标注的模式正加速向“智能驱动型”转变，人工智能辅助标注、合成数据生成、领域自适应算法等技术广泛应用，显著提升了标注效率和数据处理精度。企业通过构建智能化标注平台，整合标注工具链与数据治理流程，实现从单一标注服务向标准化、全流程解决方案的升级。

从四部门联合发布的《关于促进数据标注产业高质量发展的实施意见》可以看出，该方向已成为国家级战略布局。实施意见中明确提出，到 2027 年产业规模年均复合增长率超 20%，并推动数据标注基地建设与产学研联动创新。除了政策端发力，市场生态也在发生重构，国有企业、政府部门加速释放标注需求，公共数据标注目录编制及政务大模型协同开发成为重点方向。

未来主要挑战在标准化与质量监管与人才培养与就业拉动方面。需加快建立标注质量评估体系、行业治理规则，解决当前标注质量参差不齐的问题；与此同时，标注人才缺口或达百万量级，需推动高学历、跨领域人才向知识密集型标注岗位流动，形成“技术+专业”复合型人力结构。

3.2 InfiGUI-R1：利用强化学习，让 GUI 智能体学会规划任务、反思错误

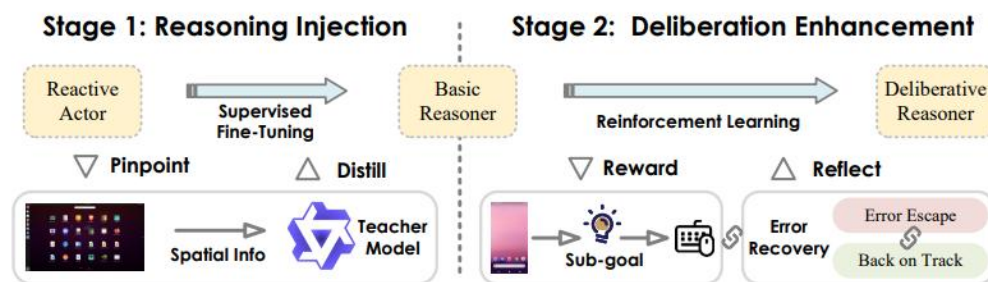
浙江大学、香港理工大学等机构的研究者们提出了 InfiGUI-R1，一个基于其创新的 Actor2Reasoner 框架训练的 GUI 智能体，旨在让 AI 像人一样在行动前思考，行动后反思。

现在智能体更类似于“反应式行动者”（Reactive Actors），主要依赖隐式推理，面对需要复杂规划和错误恢复的任务时常常力不从心。为了让 GUI 智能体更可靠、更智能地完成复杂任务，它们需要具备深思熟虑的推理能力。这意味着智能体的行为模式需要从简单的“感知 → 行动”转变为更高级的“感知 → 推理 → 行动”模式。这种模式需要智能体可以：

- 理解任务意图：将高层指令分解为具体的执行步骤；
- 进行空间推理：准确理解界面元素的布局 and 关系，定位目标；
- 反思与纠错：识别并从错误中恢复，调整策略。

为了实现这一目标，研究团队提出了 Actor2Reasoner 框架，一个以推理为核心的两阶段训练方法，旨在逐步将 GUI 智能体从“反应式行动者”培养成“深思熟虑的推理者”。

图表 15: Actor2Reasoner 架构



资料来源：InfiGUI-R1，中邮证券研究所

第一阶段：推理注入（Reasoning Injection）。此阶段的核心目标是完成从「行动者」到「基础推理者」的关键转变。研究者们采用了空间推理蒸馏（Spatial Reasoning Distillation）技术。他们首先识别出模型在哪些交互步骤中容易因缺乏推理而出错（称之为「推理瓶颈样本」），然后利用能力更强的「教师模型」生成带有明确空间推理步骤的高质量执行轨迹。

第二阶段：深思熟虑增强（Deliberation Enhancement）。在第一阶段的基础上，此阶段利用强化学习（RL）进一步提升模型的「深思熟虑」能力，重点打磨规划和反思两大核心能力。研究者们创新性地引入了两种方法：

目标引导：为了增强智能体「向前看」的规划和任务分解能力，研究者们设计了奖励机制，鼓励模型在其推理过程中生成明确且准确的中间子目标。通过评估生成的子目标与真实子目标的对齐程度，为模型的规划能力提供有效的学习信号。

错误回溯：为了培养智能体「向后看」的反思和自我纠错能力，研究者在 RL 训练中有针对性地构建了模拟错误状态或需要从错误中恢复的场景。例如，让模型学习在执行了错误动作后如何使用“返回”等操作进行“逃逸”，以及如何在“回到正轨”后重新评估并执行正确的动作。这种针对性的训练显著增强了模型的鲁棒性和适应性。

基于 Actor2Reasoner 框架，研究团队训练出了 InfiGUI-R1-3B 模型（基于 Qwen2.5-VL-3B-Instruct）。尽管只有 30 亿参数，InfiGUI-R1-3B 在多个关键基准测试中展现出了卓越的性能。

4 技术前沿

4.1 Paper2Code: 科研报告代码复现智能体

韩国科学技术院和 DeepAuto.ai 开发了一个名为 PaperCoder 的框架，能够自动从机器学习论文中生成高质量的代码仓库，从而促进科学研究的可重复性和透明度。该框架旨在解决许多研究论文缺乏相应的代码实现，导致研究人员难以复现和验证结果的问题。

图表 16: Paper2Code

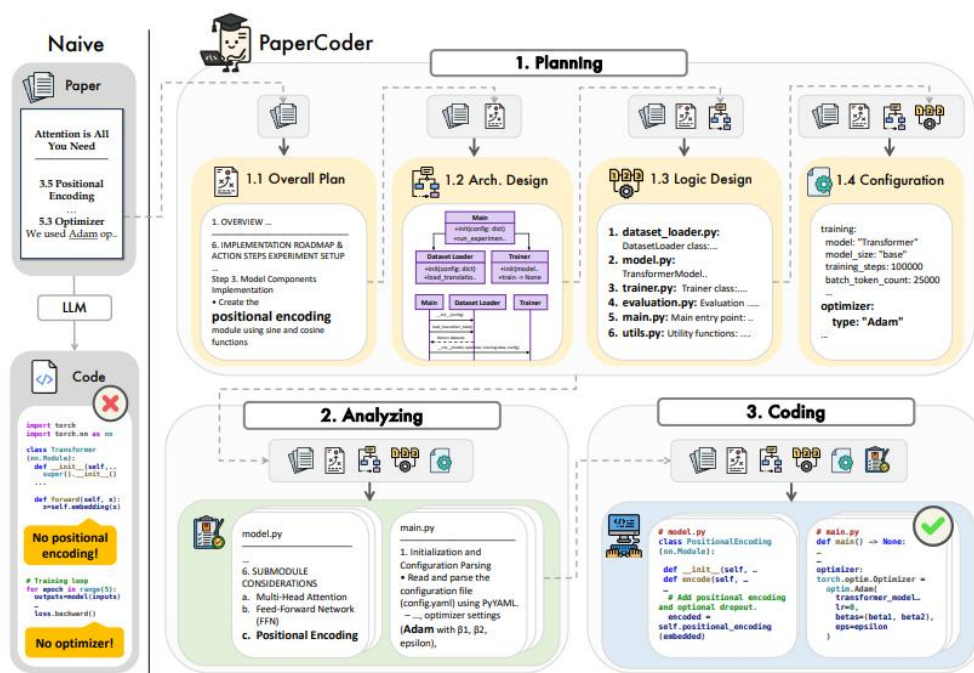
Paper2Code: Automating Code Generation from Scientific Papers in Machine Learning

Minju Seo¹, Jinheon Baek¹, Seongyun Lee¹, Sung Ju Hwang^{1,2}
KAIST¹, DeepAuto.ai²
{minjuseo, jinheon.baek, seongyun, sungju.hwang}@kaist.ac.kr

资料来源：论文 Paper2Code，中邮证券研究所

该研究要解决的核心问题是科学研究领域可重复性（Reproducibility）问题，该问题是科学进步的核心。因此 PaperCoder 采用多代理 LLM 框架，分为三个主要阶段：规划（Planning）、分析（Analysis）和生成（Generation）。每个阶段都通过一组专门设计的智能体来实例化这个过程。

图表 17: Paper2Code

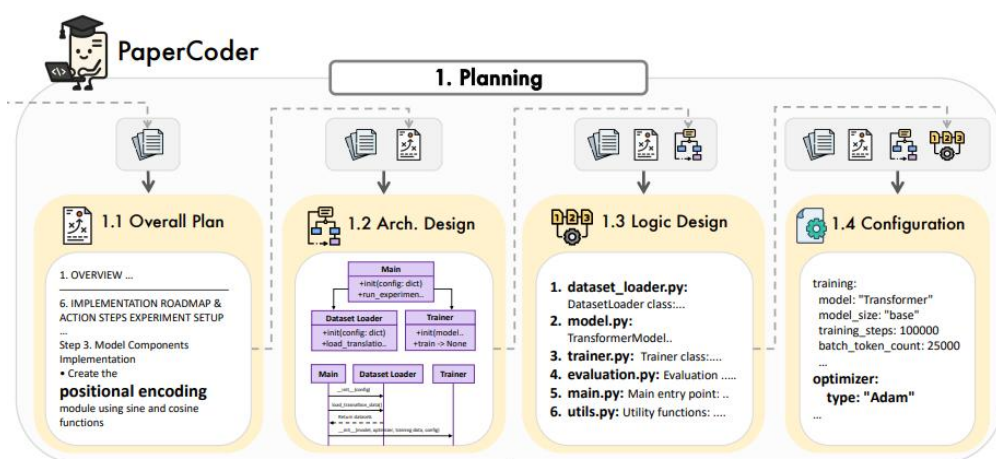


资料来源：论文 Paper2Code，中邮证券研究所

PaperCoder 旨在通过将任务分解为三个结构化阶段来模拟人类开发者和研究人员编写仓库级代码的典型生命周期：

首先，在规划阶段，框架构建了一个高层次路线图以确定要实现的核心组件，绘制了类图和序列图来建模模块之间的结构关系，识别了文件依赖关系及其执行顺序以指导正确的构建和执行流程，并生成了配置文件以使人类研究人员能够灵活定制实验 workflow。接下来是分析阶段，对每个文件和函数进行细致的解析，以理解其预期功能，例如所需的输入和输出、与其他模块的交互，以及从源论文中得出的任何算法或架构约束。最后，在生成阶段，框架根据先前确定的执行顺序以及前几个阶段产生的工件来合成整个代码库。

图表 18：第一阶段：规划阶段



资料来源：论文 Paper2Code，中邮证券研究所

在规划阶段，研究人员意识到将论文直接输入模型以生成完整的代码仓库是极具挑战性的。研究论文的主要目的是记录发现和说服读者，而非作为软件开发的结构性输入。因此，论文中常常包含大量补充信息，这些信息虽然有助于传达核心概念，但可能与实际实现无关，甚至引入干扰。为了解决这一问题，PaperCode 建议将论文分解为一个结构化的多方面计划，而不是单纯依赖论文本身。这种方法将关键的实现相关元素组织成四个不同的组件，确保生成的仓库结构良好，并与论文的方法论保持一致。因此整个规划阶段又包含以下部分：

• 架构设计

架构设计是根据总体计划和研究论文来设计软件架构。设计一个结构良好的架构对于确保多个功能模块的无缝交互至关重要。此阶段的关键是识别必要的组件并定义它们之间的关系，以确保仓库的组织性和功能性。为此，PaperCode 要求创建定义软件架构的关键工件，包括文件列表和类图。文件列表概述了软件的模块化结构，而类图则使用统一建模语言（UML）符号，直观地展示了不同组件之间的交互。

• 逻辑设计

在软件开发中，文件之间通常存在依赖关系。为了处理这些依赖关系，此阶段将研究论文以及前两个阶段生成的工件作为输入，分析每个文件及其组件的逻辑，确定必要的依赖关系和最优执行顺序。作为输出，它生成一个有序的文件列

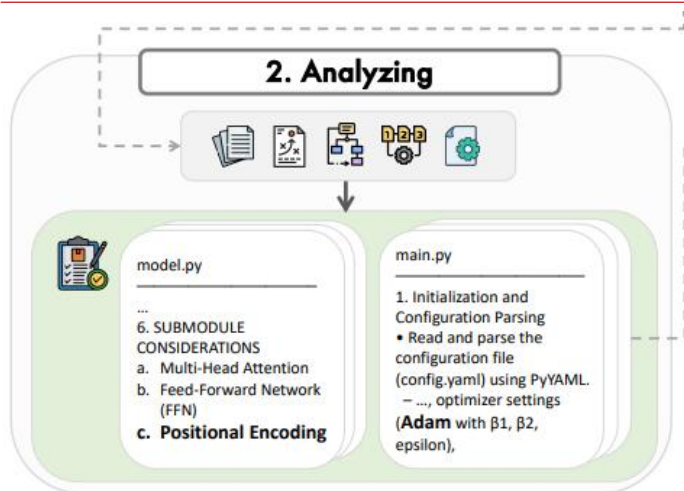
表，详细说明每个文件的角色及其在仓库中的依赖关系。这种方法确保了仓库生成不仅考虑单个文件结构，还考虑文件间的通信，从而促进了一个组织良好且逻辑连贯的实现。

• 配置文件生成

最后，配置文件生成步骤综合所有先前确定的输出，生成一个包含模型训练所需超参数和配置的配置文件（`config.yaml`）。此过程以前一阶段产生的总体计划、架构设计和有序文件列表作为输入。用户可以审查和修改 `config.yaml` 文件，以识别和纠正任何缺失或错误指定的细节，减少生成过程中可能出现的错误。

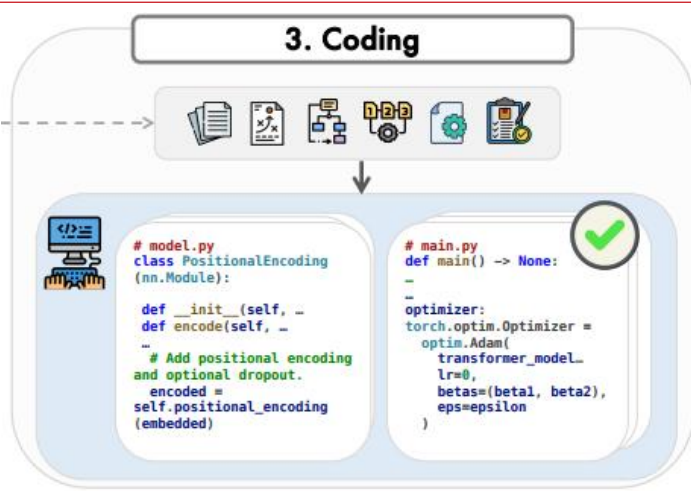
第二阶段为分析阶段，虽然规划阶段主要关注设计整体仓库结构和概述高层路线图，但分析阶段则深入到每个单独文件的具体实现细节。在这个阶段，会彻底分析仓库中每个文件的详细目的和必要考虑因素。此阶段生成的输出明确指定了每个文件应实现的目标，并强调了成功实施所需的关键因素。

图表 19：第二阶段：分析阶段



资料来源：论文 Paper2Code，中邮证券研究所

图表 20：第三阶段：编码阶段



资料来源：论文 Paper2Code，中邮证券研究所

第三阶段也是编码阶段，该阶段生成构成研究仓库的代码。每个文件的生成都由前几个阶段的综合输出指导，包括研究论文本身、总体计划、架构设计、逻辑设计、配置文件、特定文件分析以及先前生成的代码。由于仓库文件之间经常存在导入依赖关系，PaperCode 严格遵循规划阶段建立的有序文件列表，以确保顺序一致性。

经过一系列严格的实验和评估，在所有会议和两种评估模式下，PaperCoder 的表现遥遥领先于其他基线模型。在基于参考的评估中，PaperCoder 在 ICML、NeurIPS 和 ICLR 论文上的平均正确性得分分别达到了 3.72、3.83 和 3.68；在无参考评估里，得分更是高达 4.73、4.77 和 4.73，直接碾压其他模型。

5 风险提示

以上内容基于历史数据完成，在政策、市场环境发生变化时存在失效的风险；历史信息不代表未来。

中邮证券投资评级说明

投资评级标准	类型	评级	说明
报告中投资建议的评级标准： 报告发布日后的 6 个月内的 相对市场表现，即报告发布日后的 6 个月内的公司股价（或行 业指数、可转债价格）的涨跌 幅相对同期相关证券市场基准 指数的涨跌幅。 市场基准指数的选取：A 股市 场以沪深 300 指数为基准；新 三板市场以三板成指为基准； 可转债市场以中信标普可转债 指数为基准；香港市场以恒生 指数为基准；美国市场以标普 500 或纳斯达克综合指数为基 准。	股票评级	买入	预期个股相对同期基准指数涨幅在 20%以上
		增持	预期个股相对同期基准指数涨幅在 10%与 20%之间
		中性	预期个股相对同期基准指数涨幅在-10%与 10%之间
		回避	预期个股相对同期基准指数涨幅在-10%以下
	行业评级	强于大市	预期行业相对同期基准指数涨幅在 10%以上
		中性	预期行业相对同期基准指数涨幅在-10%与 10%之间
		弱于大市	预期行业相对同期基准指数涨幅在-10%以下
	可转债 评级	推荐	预期可转债相对同期基准指数涨幅在 10%以上
		谨慎推荐	预期可转债相对同期基准指数涨幅在 5%与 10%之间
		中性	预期可转债相对同期基准指数涨幅在-5%与 5%之间
		回避	预期可转债相对同期基准指数涨幅在-5%以下

分析师声明

撰写此报告的分析师（一人或多人）承诺本机构、本人以及财产利害关系人与所评价或推荐的证券无利害关系。

本报告所采用的数据均来自我们认为可靠的目前已公开的信息，并通过独立判断并得出结论，力求独立、客观、公平，报告结论不受本公司其他部门和人员以及证券发行人、上市公司、基金公司、证券资产管理公司、特定客户等利益相关方的干涉和影响，特此声明。

免责声明

中邮证券有限责任公司（以下简称“中邮证券”）具备经中国证监会批准的开展证券投资咨询业务的资格。

本报告信息均来源于公开资料或者我们认为可靠的资料，我们力求但不保证这些信息的准确性和完整性。报告内容仅供参考，报告中的信息或所表达观点不构成所涉证券买卖的出价或询价，中邮证券不对因使用本报告的内容而导致的损失承担任何责任。客户不应以本报告取代其独立判断或仅根据本报告做出决策。

中邮证券可发出其它与本报告所载信息不一致或有不同结论的报告。报告所载资料、意见及推测仅反映研究人员于发出本报告当日的判断，可随时更改且不予通告。

中邮证券及其所属关联机构可能会持有报告中提到的公司所发行的证券头寸并进行交易，也可能为这些公司提供或者计划提供投资银行、财务顾问或者其他金融产品等相关服务。

《证券期货投资者适当性管理办法》于 2017 年 7 月 1 日起正式实施，本报告仅供中邮证券客户中的专业投资者使用，若您非中邮证券客户中的专业投资者，为控制投资风险，请取消接收、订阅或使用本报告中的任何信息。本公司不会因接收人收到、阅读或关注本报告中的内容而视其为专业投资者。

本报告版权归中邮证券所有，未经书面许可，任何机构或个人不得存在对本报告以任何形式进行翻版、修改、节选、复制、发布，或对本报告进行改编、汇编等侵犯知识产权的行为，亦不得存在其他有损中邮证券商业性权益的任何情形。如经中邮证券授权后引用发布，需注明出处为中邮证券研究所，且不得对本报告进行有悖原意的引用、删节或修改。

中邮证券对于本申明具有最终解释权。

公司简介

中邮证券有限责任公司，2002 年 9 月经中国证券监督管理委员会批准设立，注册资本 50.6 亿元人民币。中邮证券是中国邮政集团有限公司绝对控股的证券类金融子公司。

公司经营范围包括：证券经纪；证券自营；证券投资咨询；证券资产管理；融资融券；证券投资基金销售；证券承销与保荐；代理销售金融产品；与证券交易、证券投资活动有关的财务顾问。此外，公司还具有：证券经纪人业务资格；企业债券主承销资格；沪港通；深港通；利率互换；投资管理人受托管理保险资金；全国银行间同业拆借；作为主办券商在全国中小企业股份转让系统从事经纪、做市、推荐业务资格等业务资格。

公司目前已经在北京、陕西、深圳、山东、江苏、四川、江西、湖北、湖南、福建、辽宁、吉林、黑龙江、广东、浙江、贵州、新疆、河南、山西、上海、云南、内蒙古、重庆、天津、河北等地设有分支机构，全国多家分支机构正在建设中。

中邮证券紧紧依托中国邮政集团有限公司雄厚的实力，坚持诚信经营，践行普惠服务，为社会大众提供全方位专业化的证券投、融资服务，帮助客户实现价值增长，努力成为客户认同、社会尊重、股东满意、员工自豪的优秀企业。

中邮证券研究所

北京

邮箱：yanjiusuo@cnpsec.com

地址：北京市东城区前门街道珠市口东大街 17 号

邮编：100050

上海

邮箱：yanjiusuo@cnpsec.com

地址：上海市虹口区东大名路 1080 号邮储银行大厦 3 楼

邮编：200000

深圳

邮箱：yanjiusuo@cnpsec.com

地址：深圳市福田区滨河大道 9023 号国通大厦二楼

邮编：518048