

반도체

AI에 대한 다소 과장된 우려

SK증권 리서치센터



Analyst
한동희

donghee.han@sk.com
3773-8826



R.A
박제민

jeminwa@sk.com
3777-8884

DeepSeek, 저비용, 고성능 AI 모델 발전 방향 제시

중국의 DeepSeek는 R1 추론 모델 출시를 통해 저비용, 고성능 AI 모델의 발전 방향을 제시했다고 판단한다. R1 모델은 DeepSeek-V3를 기반으로 하며, 지도 학습 없이 자체 강화학습에 의한 CoT (Chain of Thought)를 통해 자체적인 추론 능력 강화를 강화한 R1-Zero의 문제점을 Cold start data를 통해 개선한 모델이다. R1의 전체 파라미터 수는 671B 개이나, MoE (Mixture of Expert)를 통한 실제 활성화 파라미터 수는 37B 개에 불과하다. Artificial Analysis에 따르면 DeepSeek-R1은 경쟁 모델 대비 Quality 측면에서는 최상위권, 토큰 당 가격 측면에서는 중위권 수준이라는 점에서 높은 가성비 모델이라고 평가할 수 있다. 알고리즘을 통한 AI 성능 제고는 AI 사이클의 지속 명분을 강화시키는 요소가 될 것으로 전망한다.

높은 컴퓨팅 파워와 강력한 거대 기반 모델 필요성 지속

알고리즘 발전을 통한 AI 성능 제고에도 높은 컴퓨팅 파워와 강력한 거대 기반 모델의 필요성은 지속될 것으로 전망한다. DeepSeek-R1의 Cold start에 활용된 data set은 프론티어급 모델에서 비롯되었으며, 이는 DeepSeek-R1 논문에서 언급된 대로 작은 모델은 높은 컴퓨팅 파워를 기반으로 한 거대 강화모델을 요구하며, 더 경제적, 효과적 Intelligence 역시 강력한 기반 모델과 Large-scale의 강화 학습에 기반하므로 하드웨어에 대한 시장의 우려는 다소 과장되어있다고 판단한다.

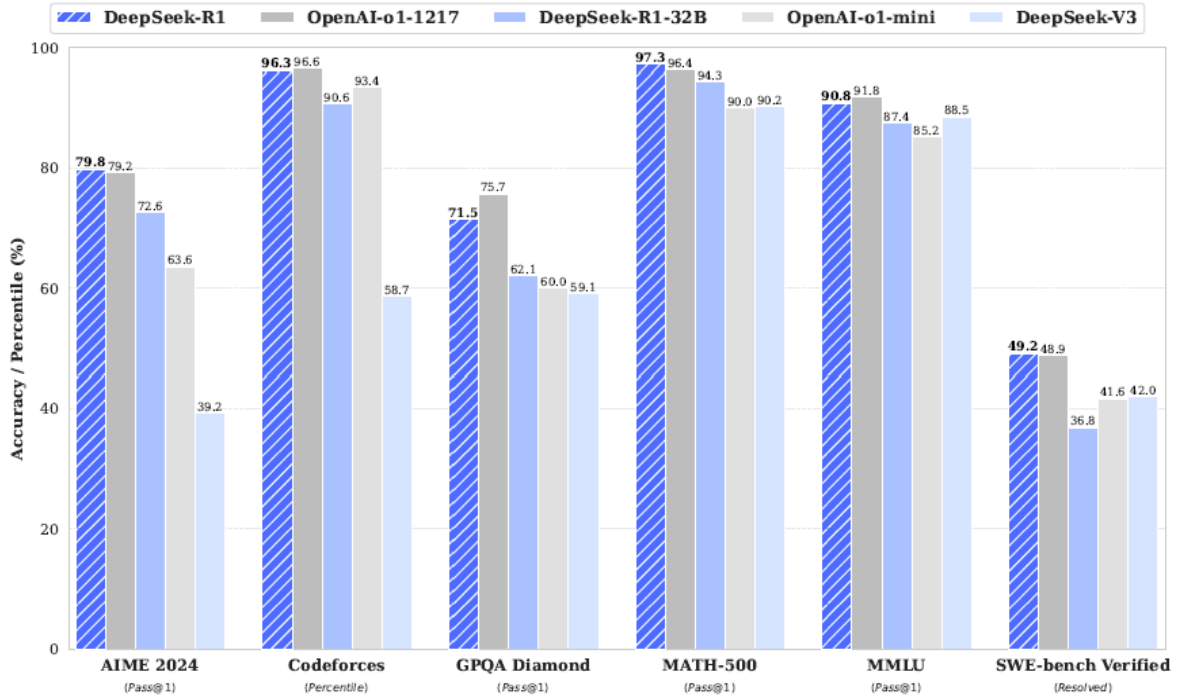
추론 고도화가 견인하는 Scaling 법칙의 지속: Test-time scaling

SK 증권은 Test-time scaling을 통한 Scaling 법칙의 지속을 전망한다. 하드웨어 관점에서 가장 중요한 것은 컴퓨팅 파워 제고에 따른 Incentive 지속 여부이며, 이는 기존 방식의 한계점과 맞닿아있다는 점을 감안하면, 저비용 훈련 방법은 Scaling 법칙을 지속시킬 수 있는 명분이 될 것으로 판단하기 때문이다. CES2025 Keynote에서의 젠슨 황과, DeepSeek R1-Zero 논문, OpenAI 역시 o1을 통해 더 많은 추론 시간을 할애 할수록, 모델의 정확도가 높아진다고 언급하고 있다.

저비용 고효율 AI 모델의 대두는 시장 수요를 더욱 촉진시킬 것

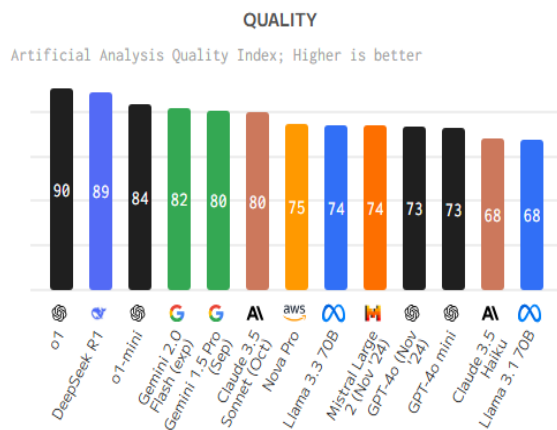
저비용 고효율 AI 모델의 대두는 AI에 대한 시장 수요를 더욱 촉진시키는 요소 중 하나가 될 것으로 전망한다. 또한 그 과정에서 해당 모델의 훈련 및 추론을 위한 Custom HBM 등 최적화 메모리 수요 역시 점증할 것이다. AI 경쟁 심화 지속과 저변 확대, 그 과정에서의 컴퓨팅 파워 제고의 방향성은 여전하다고 판단한다. 반도체 업종에 대해 비중확대 의견을 유지한다.

DeepSeek-R1의 Benchmark performance 비교



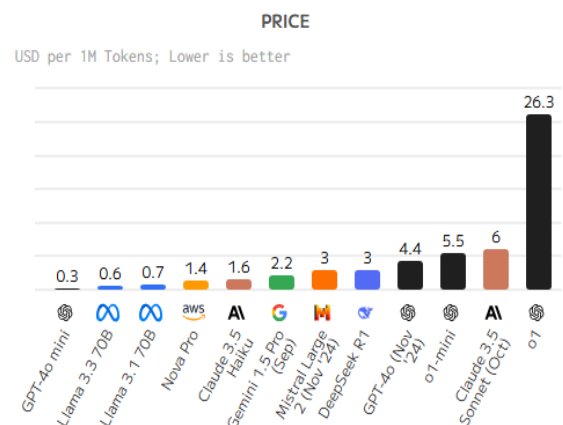
자료: DeepSeek-R1, SK 증권

Artificial Analysis Index (Quality)



자료: Artificial Analysis, SK 증권

Artificial Analysis Index (Price)



자료: Artificial Analysis, SK 증권

DeepSeek-V3 훈련 비용

Training Costs	Pre-Training	Context Extension	Post-Training	Total
in H800 GPU Hours	2664K	119K	5K	2788K
in USD	\$5.328M	\$0.238M	\$0.01M	\$5.576M

Lastly, we emphasize again the economical training costs of DeepSeek-V3, summarized in Table 1, achieved through our optimized co-design of algorithms, frameworks, and hardware. During the pre-training stage, training DeepSeek-V3 on each trillion tokens requires only 180K H800 GPU hours, i.e., 3.7 days on our cluster with 2048 H800 GPUs. Consequently, our pre-training stage is completed in less than two months and costs 2664K GPU hours. Combined with 119K GPU hours for the context length extension and 5K GPU hours for post-training, DeepSeek-V3 costs only 2.788M GPU hours for its full training. Assuming the rental price of the H800 GPU is \$2 per GPU hour, our total training costs amount to only \$5.576M. Note that the aforementioned costs include only the official training of DeepSeek-V3, excluding the costs associated with prior research and ablation experiments on architectures, algorithms, or data.

자료: DeepSeek-V3, SK 증권

DeepSeek, 높은 컴퓨팅 파워와 강력한 거대 기반 모델에 대한 필요성 언급

Model	AIME 2024		MATH-500	GPQA Diamond	LiveCodeBench
	pass@1	cons@64	pass@1	pass@1	pass@1
QwQ-32B-Preview	50.0	60.0	90.6	54.5	41.9
DeepSeek-R1-Zero-Qwen-32B	47.0	60.0	91.6	55.0	40.2
DeepSeek-R1-Distill-Qwen-32B	72.6	83.3	94.3	62.1	57.2

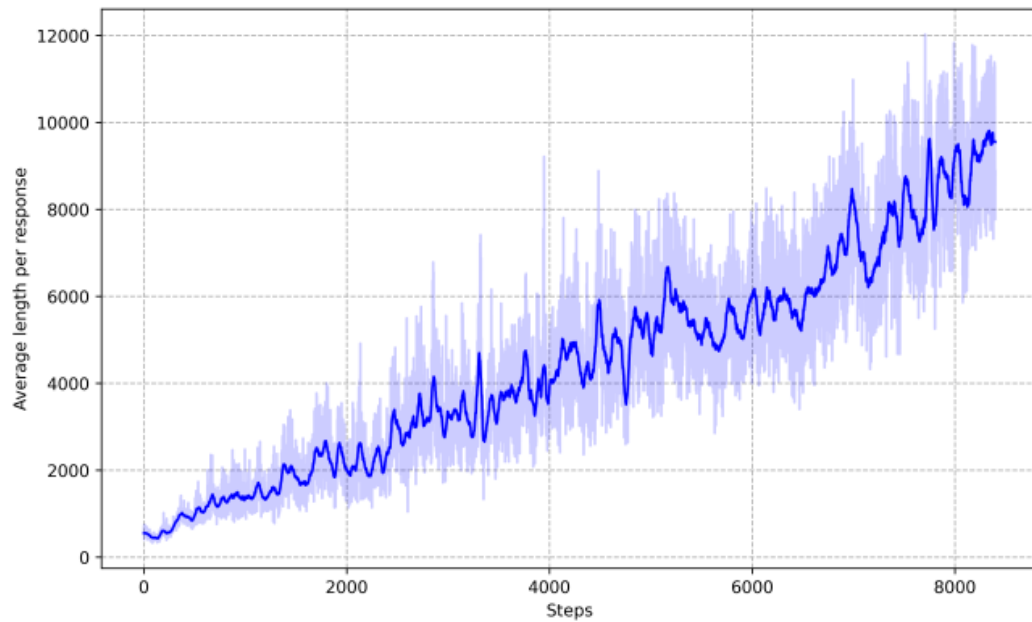
Table 6 | Comparison of distilled and RL Models on Reasoning-Related Benchmarks.

RL training, achieves performance on par with QwQ-32B-Preview. However, DeepSeek-R1-Distill-Qwen-32B, which is distilled from DeepSeek-R1, performs significantly better than DeepSeek-R1-Zero-Qwen-32B across all benchmarks.

Therefore, we can draw two conclusions: First, distilling more powerful models into smaller ones yields excellent results, whereas smaller models relying on the large-scale RL mentioned in this paper require enormous computational power and may not even achieve the performance of distillation. Second, while distillation strategies are both economical and effective, advancing beyond the boundaries of intelligence may still require more powerful base models and larger-scale reinforcement learning.

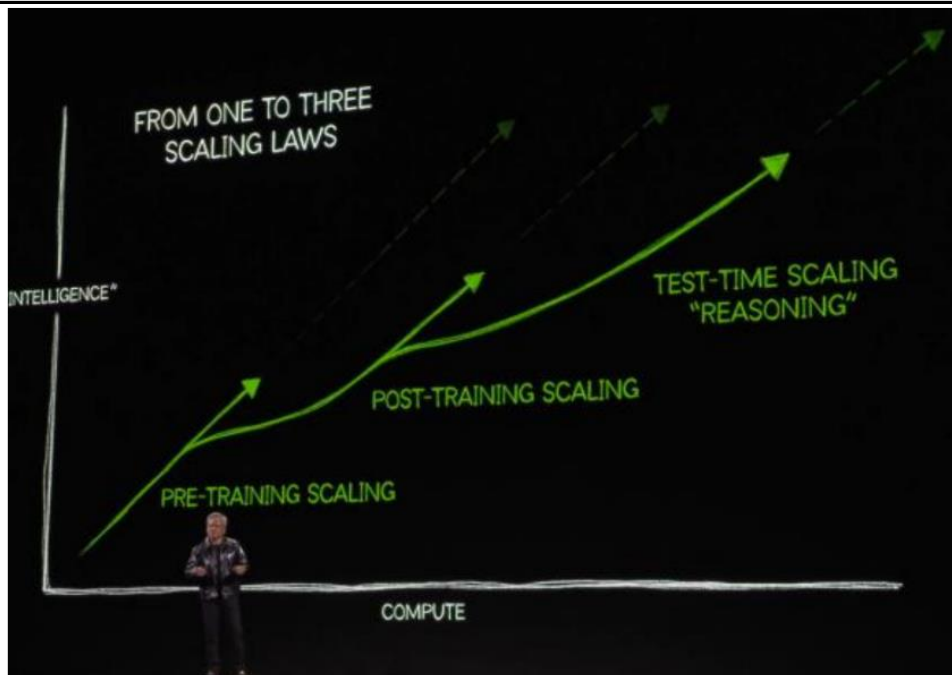
자료: DeepSeek-R1, SK 증권

DeepSeek-R1-Zero: 더 많은 사고 시간 할애를 통한 자연적인 추론 업무 학습



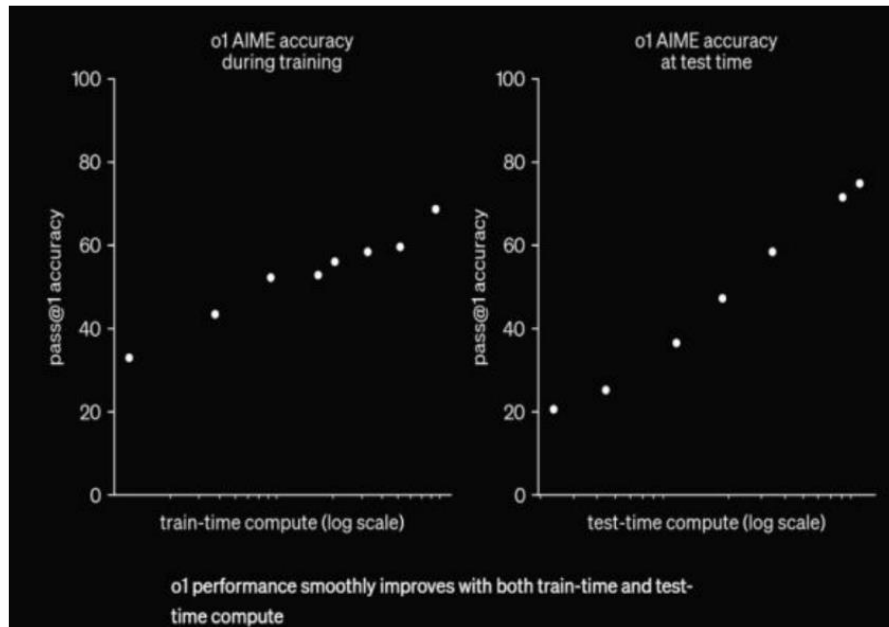
자료: DeepSeek-R1-Zero, SK 증권

CES2025에서 제시된 Scaling law의 지속: Test-Time Scaling



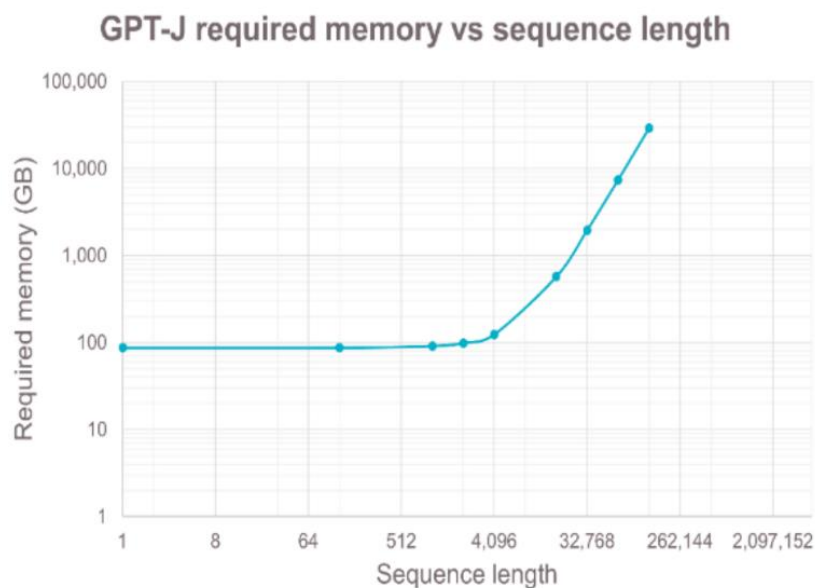
자료: CES2025 NVIDIA Keynote, SK 증권

Test time (Computing time)을 늘릴수록 모델 정확도 증가



자료: OpenAI, SK 증권

Sequence length의 증가는 요구 메모리의 증가를 견인



자료: Cerabras Inference, SK 증권

Compliance Notice

작성자(관리자)는 본 조사분석자료에 게재된 내용들이 본인의 의견을 정확하게 반영하고 있으며, 외부의 부당한 압력이나 간섭없이 신의성실하게 작성되었음을 확인합니다.

본 보고서에 언급된 종목의 경우 당사 조사분석담당자는 본인의 담당종목을 보유하고 있지 않습니다.

본 보고서는 기관투자가 또는 제 3자에게 사전 제공된 사실이 없습니다.

당사는 자료공표일 현재 해당기업과 관련하여 특별한 이해 관계가 없습니다.

종목별 투자의견은 다음과 같습니다.

투자판단 3 단계(6개월기준) 15%이상 -> 매수 / -15%~15% -> 중립 / -15%미만 -> 매도

SK 증권 유니버스 투자등급 비율 (2025년 01월 31일 기준)

매수	97.47%	중립	2.53%	매도	0.00%
----	--------	----	-------	----	-------