



2026년 9월

RSI가 나타났다! 이번엔 진짜예요!

상상인증권 리서치센터

투자전략팀 김경태 | 글로벌전략/디지털자산 02-3779-3427, kt.kim@sangsanginib.com

 상상인증권

CONTENTS

01	예상보다 빠른 AI의 업무 대체	4
02	토큰 수요의 확대와 AI 연구소 매출	12
03	추론 수요 확대와 부품별 반도체 수혜	19
04	AI 설비투자자와 매출·현금 회수의 시차	36
05	추론 인프라와 반도체의 투자 판단	43
06	해외 기업분석	52
07	부록 1. AI의 다음 단계와 읽을거리	63
08	부록 2. FAQ	74

Executive Summary

에이전트 확산에 따른 업무당 토큰 처리량의 구조적 증가

- 토큰 단가 하락이 사용 확대로 이어지는 Jevons 효과 확인. 백만 토큰당 실효 단가가 1.5달러에서 0.8달러로 하락하는 동안 주간 토큰 수는 41배 증가했고, 에이전트는 단계마다 모델을 다시 호출해 업무당 토큰 처리량을 추가로 확대

컴퓨터 사용과 월드모델 역량에 기반한 피지컬 AI 확장

- GPT-6 Astra는 컴퓨터 사용(OSWorld 2.0 72.6%), 처음 보는 게임의 규칙 학습(ARC-AGI-3 99.9%), 실차 주행(콘 코스 완주)에서 이전 세대 대비 도약. 환경을 관찰하고 결과를 예측하는 역량이 로봇으로 확장되며 피지컬 AI가 다음 반도체 수요원으로 부상할 전망

동시 서빙과 KV 캐시가 심화시키는 GPU, 메모리, 광통신 부족

- 모델 대형화로 가중치를 여러 GPU에 나눠 적재하고, 사용자별로 따로 보관하는 KV 캐시는 동시 사용자 수와 문맥 길이에 비례해 증가. 사용자의 대기 시간과 답변 길이를 예측할 수 없어 부담이 GPU와 HBM에서 광통신, 서버 DRAM, NAND로 확산

2030년 말 OpenAI·Anthropic 연환산 매출 1조 900억 달러 전망

- 기본 시나리오 기준 2026년 말 추정치의 6.8배이며, 하단 시나리오 6,650억 달러에서도 AI 인프라의 주도업종 지위 유지 전망. Anthropic 연환산 매출은 9월 중순 1,000억 달러 이상으로 7월 650억 달러에서 상향

클라우드 계약잔고 확대와 투자 회수의 시차

- 계약잔고 증가는 수요 확보의 근거이나 현금 회수는 설비 가동과 과금 개시 이후. 2027E 매출 대비 FCF는 메모리·스토리지 39~57%, 네오클라우드 -116~-310%로 구매자의 투자 부담이 공급사 현금으로 이전

메모리·광통신 중심의 추천 종목과 AI 반도체 재평가

- 추천 종목은 Lumentum(광통신), Sandisk(NAND), Micron(HBM, DRAM), Western Digital(HDD). EPS 상향에도 PER이 하락한 AI 반도체는 재평가 여지가 크나, 2028E 메모리 이익 정점 이후의 지속성 확인 필요

RSI가 나타났다! 이번엔 진짜예요!

PART 01

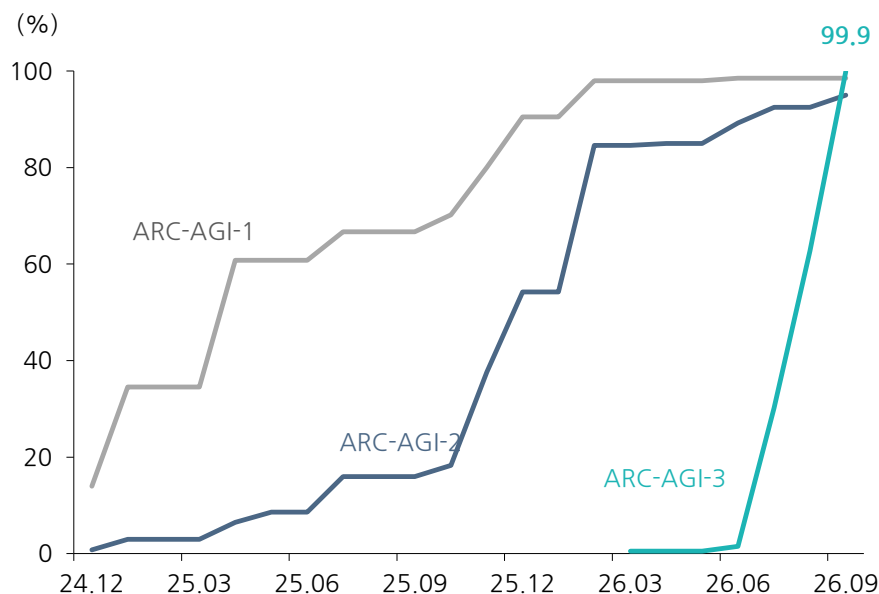
예상보다 빠른 AI의 업무 대체

예상보다 빠른 AI의 업무 대체

벤치마크 포화와 종합지수의 변별력 한계

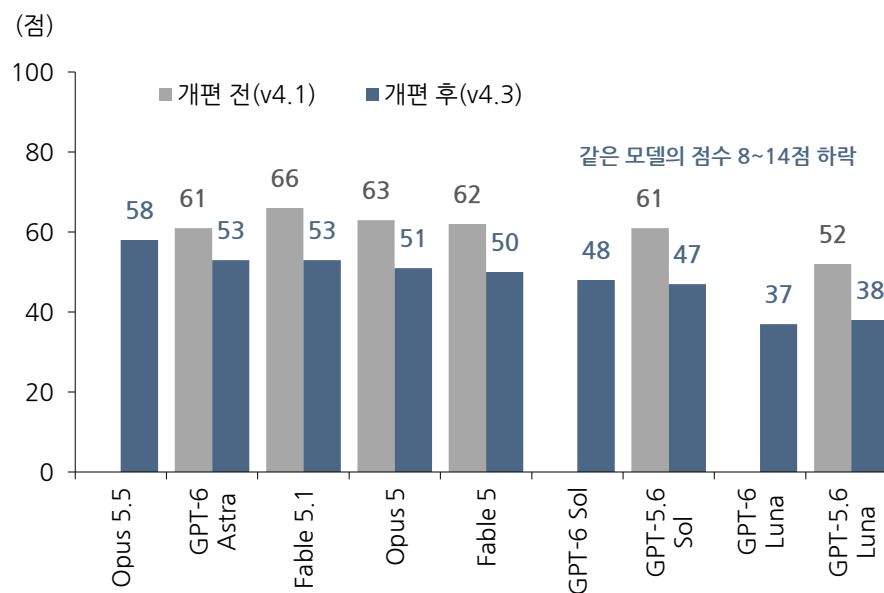
- 벤치마크는 만점에 근접하면 모델 간 차이를 변별하지 못하는 구조. 퍼즐 규칙 추론 시험 ARC-AGI는 2세대까지 최고점 95% 이상으로 포화, 3월 공개된 3세대도 6개월 만에 99.9%
- 여러 시험을 묶은 종합지수는 구성 시험 교체만으로도 같은 모델의 점수가 변동. Artificial Analysis 지수는 9월 개편으로 주요 모델 점수가 8~14점 하락했고, 개편 후 1위는 9월 22일 출시된 Opus 5.5(58점)
- 종합지수는 비용 개선도 반영하지 못하는 한계. 같은 날 출시된 GPT-6 Sol과 Luna는 가격을 전작의 절반으로 낮췄지만 점수는 전작과 같은 수준

그림 1. ARC-AGI 세대별 최고 점수의 포화



자료: ARC Prize, 상상인증권 리서치센터

그림 2. Artificial Analysis 지수 개편 전후와 최신 모델 점수



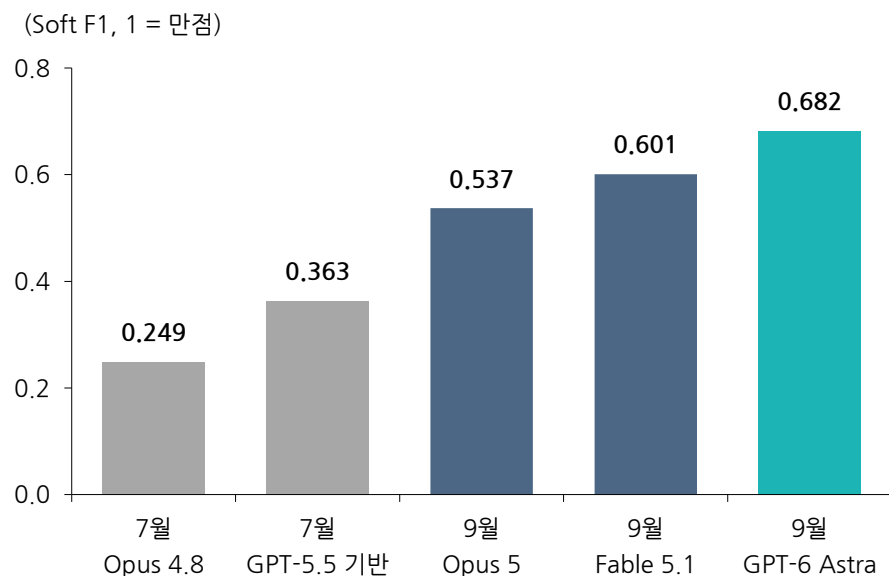
자료: Artificial Analysis, 상상인증권 리서치센터

예상보다 빠른 AI의 업무 대체

종합지수가 놓친 차세대 모델의 성능 도약

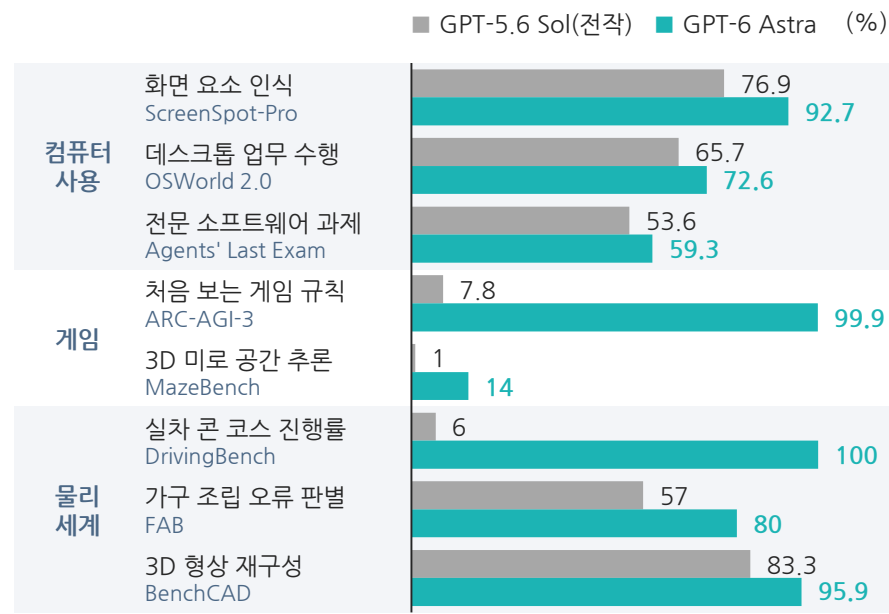
- 종합지수는 차세대 모델의 도약을 과소 반영. GPT-6 Astra의 과학 에이전트 시험(Terminal-Bench Science) 정답률은 전작 22.4%에서 64.6%로 상승했으나 Artificial Analysis 지수 상승폭은 6점
- Perplexity가 7월 공개한 대규모 웹 조사 시험 WANDR(1점 만점)도 두 달 만에 최고점 1.9배 상승, 9월 GPT-6 Astra가 0.682 기록
- 차세대 모델의 도약은 종합지수에 반영되지 않은 과제에서 먼저 확인. GPT-6 Astra는 컴퓨터 사용, 게임, 실차 운전처럼 이전 세대가 수행하지 못한 과제에서 큰 폭 상승했고, 운전 코스 진행률은 전작 6%에서 완주로 도약

그림 3. WANDR(대규모 웹 조사 과제 시험) 점수의 상승



자료: Perplexity, BenchmarkList, 상상인증권 리서치센터
주: Opus 5·Fable 5.1은 Astra 대비 격차로 역산

그림 4. 컴퓨터 사용·게임·물리 세계 과제의 세대 간 격차



자료: OpenAI, ARC Prize, Quesma, DrivingBench, Epoch AI, 상상인증권 리서치센터
주: ARC-AGI-3 Astra는 제공사 전용 조건

예상보다 빠른 AI의 업무 대체

밀레니엄 난제 해결에 이른 AI의 급격한 도약

- AI가 수학 최고 난제의 해법 도출과 검증까지 도달. OpenAI는 9월 에이전트 약 1만 개로 7대 밀레니엄 문제 중 하나인 Navier-Stokes 문제의 해법에 88시간 만에 도달했고, 증명 검증 언어 Lean으로 확인
- 시험 점수뿐 아니라 실무 지표도 동반 상승. 실제 프리랜스 과제의 자동화율(RLI)은 2025년 10월 2.5%에서 2026년 7월 15.8%로 상승
- 실제 성과와 사용량 모두 예상 상회. 2026년 말 14%로 예상된 LiveCodeBench Pro 고난도 점수는 5월 53.8%, Meta의 개인 에이전트 Muse는 이용자 사용량이 테스트 집단의 10배로 자사 예상 상회

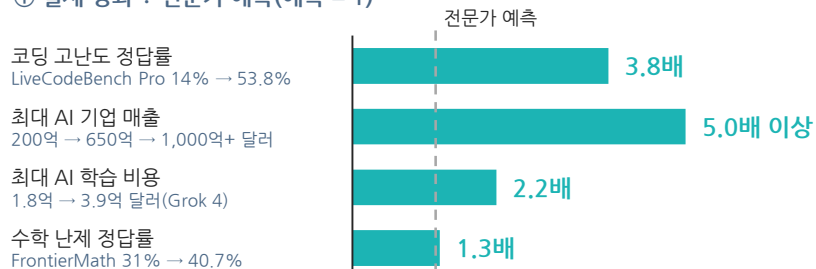
표 1. 2025년 말 전후와 최근의 AI 수준

지표	2025년 말 전후	최근
ARC-AGI-2 최고점	54.2% 2025.12	95.0% 2026.9
ARC-AGI-3 최고점	0.5% 2026.3 공개 시점	99.9% 2026.9
전문가 대비 우세·동등 판정 비율(GDPval)	70.9% 2025.12	84.9% 2026.4
실제 프리랜스 과제 자동화율(RLI)	2.5% 2025.10	15.8% 2026.7
말길 수 있는 과제 길이(METR)	5.3시간 2026.1	16시간 이상 2026.5
밀레니엄 문제	미해결	Navier-Stokes 해법·검증 2026.9

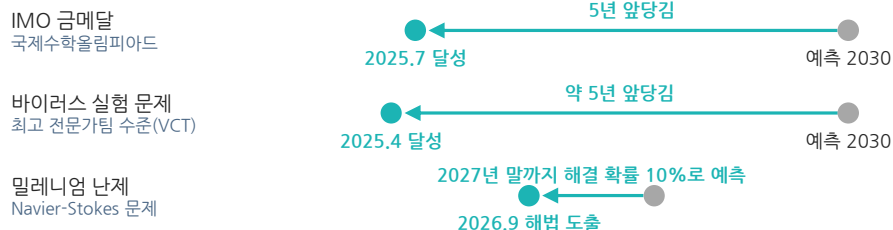
자료: ARC Prize, OpenAI, Scale AI·CAIS, METR, 상상인증권 리서치센터
주: 과제 길이는 성공률 50% 기준, ARC-AGI-3는 제공사 전용 조건

그림 5. 전문가 예측을 넘어선 AI 성과와 달성 시점

① 실제 성과 ÷ 전문가 예측(예측 = 1)



② 달성 시점: 전문가 예측 → 실제



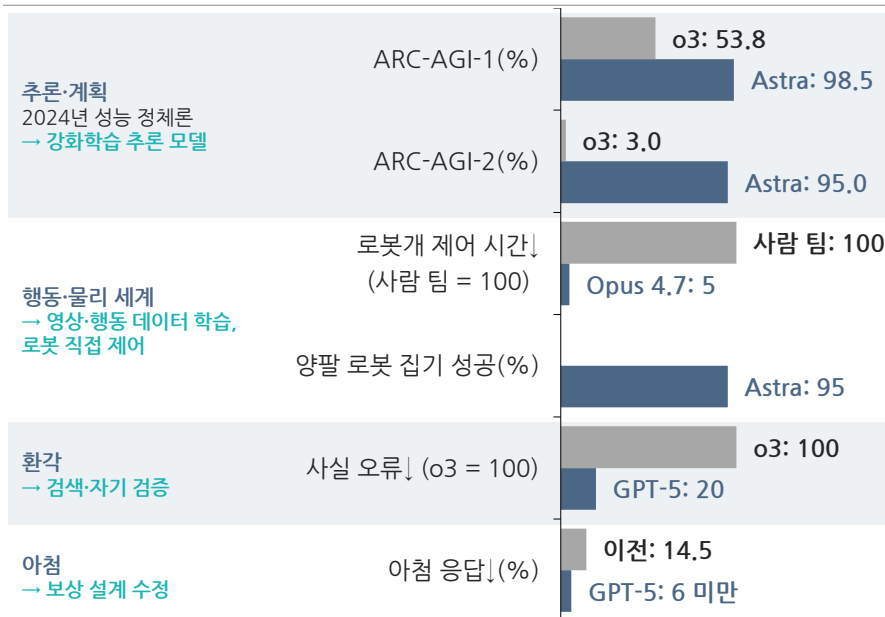
자료: Forecasting Research Institute, Bloomberg, YipitData, 상상인증권 리서치센터
주: 전문가 예측은 중앙값, 코딩은 2026년 말 예측 대비 5월 실제, 학습 비용은 2030년 예측

예상보다 빠른 AI의 업무 대체

Transformer 회의를 반박과 학습 컴퓨팅 수요 확대

- Transformer 기반 LLM은 추론과 물리 세계 이해가 불가능하다는 안 르쿤 등의 회의론은 성과로 반박. ARC-AGI-3는 6개월 만에 99.9%, 운전 학습이 없는 GPT-6 Astra는 실제 승용차로 콘 코스 완주
- 환각과 아침 같은 약점도 새 구조 없이 학습 방식 개선과 컴퓨팅 확대로 해소. 사실 오류는 o3 대비 80% 감소, 아침 응답 비율은 14.5%에서 6% 미만으로 하락
- 컴퓨팅과 데이터 투입을 늘리는 AGI 경로에서 반도체, 메모리, 스토리지 수요는 AI 발전의 직접 변수. 학습 컴퓨팅량은 모델 규모와 데이터양의 곱에 비례해 상반기까지 정체한 모델 규모의 재확대가 수요로 연결

그림 6. Transformer 회의론의 쟁점별 반박 지표



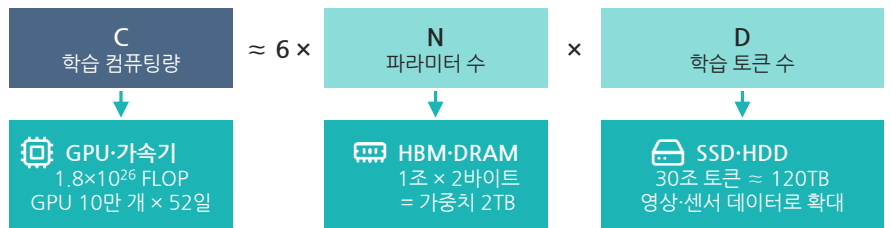
자료: OpenAI, Anthropic, ARC Prize, RoboCurve, 상상인증권 리서치센터

그림 7. 새 구조 없이 확장되는 Transformer와 학습 컴퓨팅량

① 자기 개선과 물리 세계 과제로 확인된 Transformer 확장



② 학습 컴퓨팅량 공식과 반도체 수요



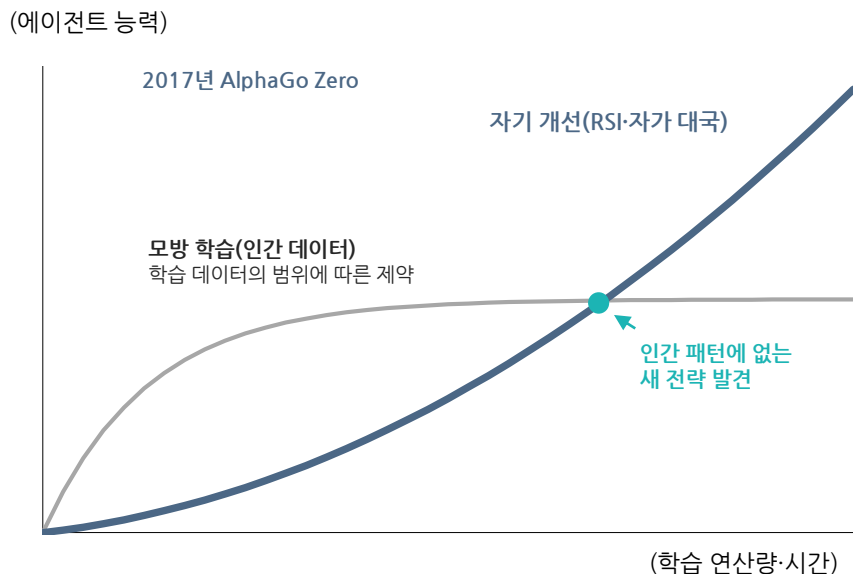
자료: X, Forbes, DrivingBench, OpenAI, Google DeepMind, 상상인증권 리서치센터
주: ② N = 1조, D = 30조 토큰, GPU당 10¹⁵ FLOP/s, 가동률 40%

예상보다 빠른 AI의 업무 대체

인간 데이터의 한계를 넘는 재귀적 자기 개선(RSI)의 시작

- 스스로 문제를 만들고 결과를 검증하는 RSI는 인간 데이터의 양과 사고 패턴의 한계를 함께 돌파. AlphaGo Zero는 인간 기보 없이 학습 3일 만에 기존 AlphaGo에 100 : 0으로 승리
- RSI는 게임을 넘어 LLM으로 확산. Meta Muse Spark 1.1이 후속 모델의 학습 데이터를 스스로 구축했고, Anthropic과 OpenAI 경영진은 RSI가 업계 전반에서 시작됐다며 개발 속도 조절을 제안
- 데이터 생성과 검증에 투입되는 컴퓨팅만큼 가속기와 메모리 수요 추가 발생. 하이브리드 AI(코드를 작성하고 실행해 답을 검증하는 LLM)가 학습 데이터를 스스로 만드는 구조

그림 8. 모방 학습의 한계와 RSI의 돌파



자료: Google DeepMind, 상상인증권 리서치센터

그림 9. 게임에서 LLM으로 확산된 RSI



자료: Google DeepMind, Anthropic, Meta, OpenAI, 상상인증권 리서치센터

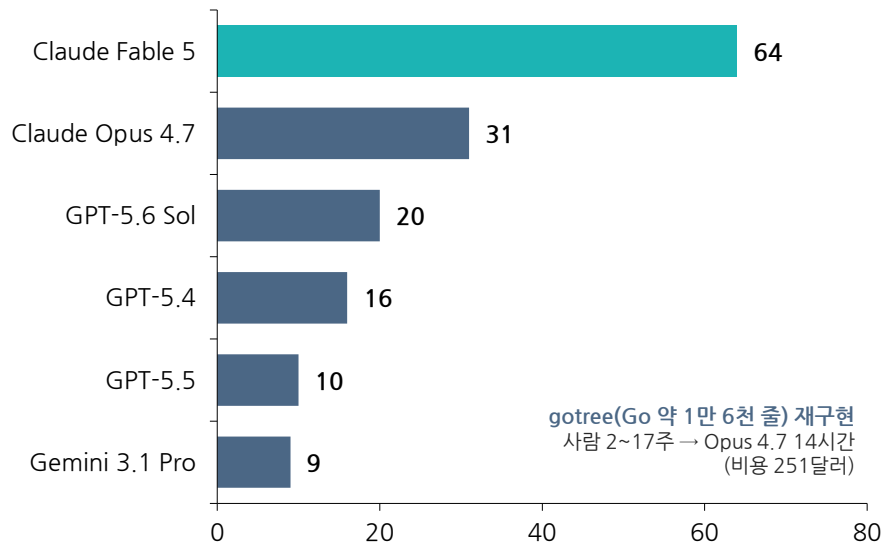
예상보다 빠른 AI의 업무 대체

긴 업무 완주율 상승과 토큰 단가 하락

- 긴 코딩 작업의 완주율은 한 세대 만에 두 배. 중대형 프로그램 재구현 시험(MirrorCode)의 전 테스트 통과율은 Opus 4.7 31%에서 후속 Fable 5 64%로 상승
- 단계가 많은 업무일수록 작은 성능 개선이 크게 증폭. 20단계 업무 완주율은 단계당 성공률 90%면 12%, 95%면 36%
- 성공 건당 비용(시도 1회 비용 ÷ 완주율)에는 단가 하락과 완주율 상승이 곱으로 작용. 같은 성능 기준 토큰 단가는 3년간 약 2,000분의 1, 실제 지불 평균인 실패 단가도 하락 방향

그림 10. 긴 코딩 작업의 모델별 완주율(MirrorCode)

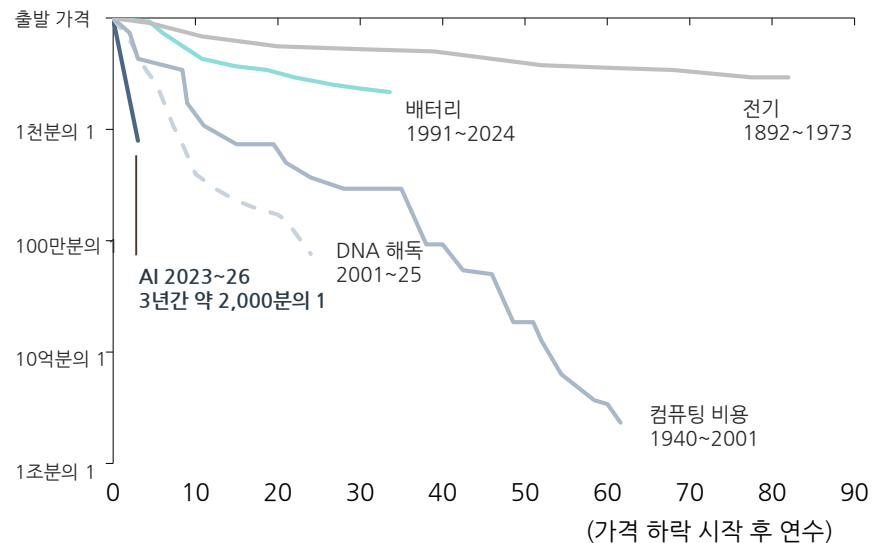
(%, 전 테스트 통과)



자료: METR, Epoch AI, 상상인증권 리서치센터

그림 11. 주요 기술의 가격 하락 속도(물가 조정)

(출발 가격 대비, 로그)



자료: Epoch AI, 상상인증권 리서치센터
주: AI는 같은 성능 기준 토큰 단가

예상보다 빠른 AI의 업무 대체

과거 기술보다 넓고 빠른 대체와 저연차 고용 감소

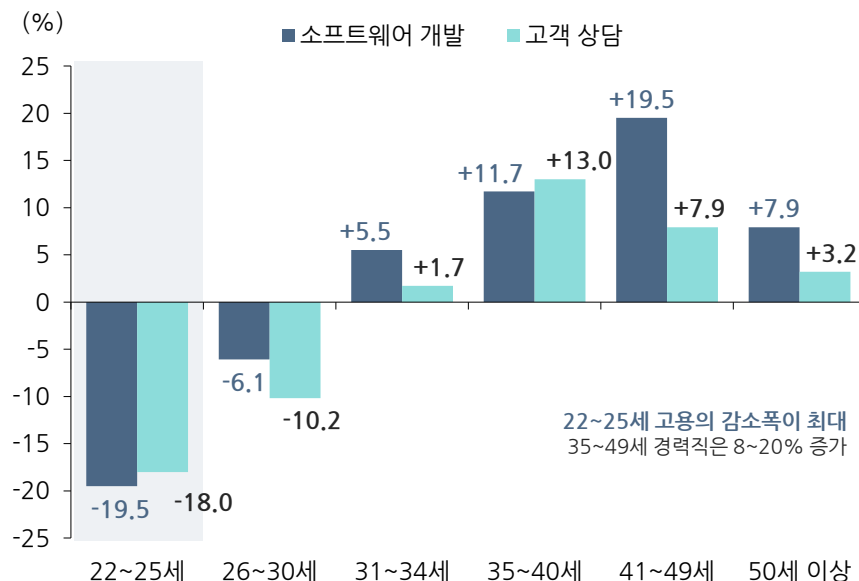
- AI의 업무 대체는 과거 디지털 기술보다 범위가 넓고 속도가 빠른 구조. 필름 현상과 전통 조판 등 과거의 전환은 한 기능씩 수년에서 수십 년에 걸쳐 진행
- AI는 문서, 분석, 코딩, 디자인을 같은 모델의 업데이트로 동시에 자동화해 대체 인력이 옮겨 갈 인접 직무도 함께 축소
- AI가 실행 업무를 맡고 사람은 목표 설정과 결과 검수로 이동하면서 저연차 고용부터 감소. 2022년 11월 ChatGPT 출시 이후 2026년 8월까지 22~25세 소프트웨어 개발 고용은 19.5% 감소한 반면 35~49세 경력직은 8~20% 증가

그림 12. 기능을 흡수한 기술과 AI의 전환 범위

기존 기능	대체한 기술	변화 사례
필름 현상·인화	디지털 카메라·스마트폰	Kodak 소비자 영상 2012~13년 사업 재편
도장·인감 제작	본인서명 확인·전자서명	본인서명사실확인제 2012년 시행
MP3 플레이어	스마트폰·스트리밍	Apple iPod 2022년 판매 종료
전통 조판·식자	DTP·PostScript	1980년대부터 출판 공정 디지털화
문서·분석·코딩·디자인	범용 AI 모델·에이전트	여러 직무의 동시 전환 모델 업데이트 주기(수개월)

자료: Kodak, 행정안전부, Apple, Computer History Museum, 상상인증권 리서치센터

그림 13. 2022년 11월 대비 AI 노출 직무의 연령대별 고용 변화



자료: Stanford Digital Economy Lab, ADP, 상상인증권 리서치센터

RSI가 나타났다! 이번엔 진짜예요!

PART 02

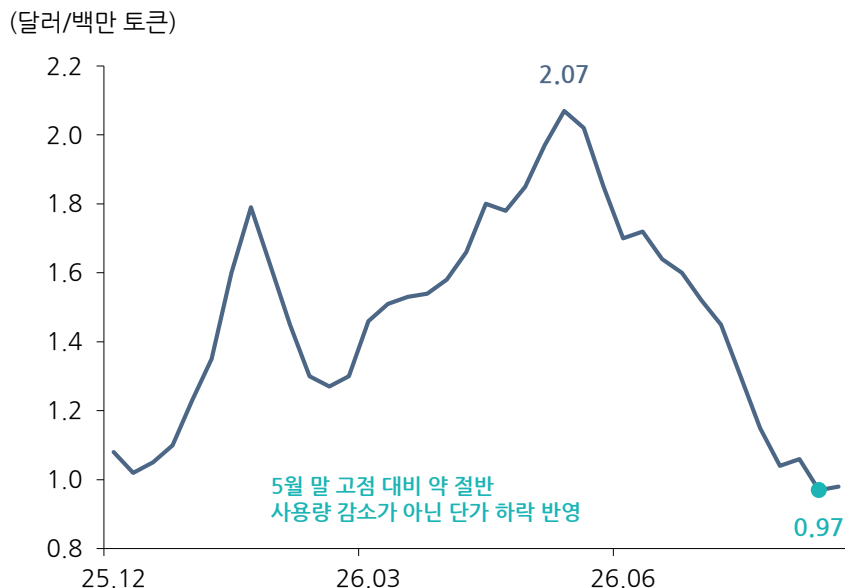
토큰 수요의 확대와 AI 연구소 매출

토큰 수요의 확대와 AI 연구소 매출

토큰 지수 하락의 단가 요인과 도입률 통계 격차

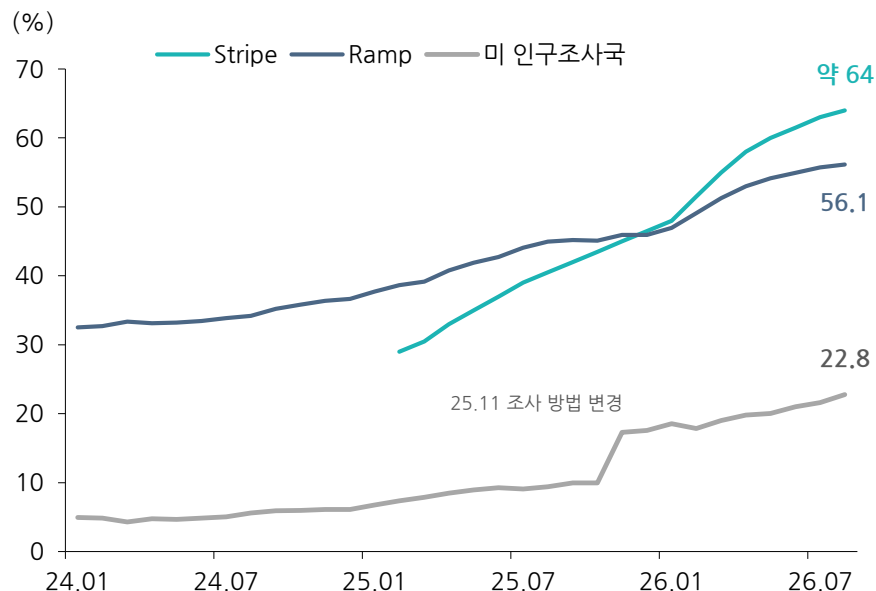
- 토큰 지수 하락은 수요 감소가 아닌 단가 하락의 결과. Silicon Data 지수(백만 토큰당 가격)는 5월 말 고점의 절반이나, 같은 기간 Anthropic과 OpenAI 연환산 매출은 급증
- 토큰 단가는 어느 기준으로든 하락 방향. 이 지수는 사용량 가중 실효 단가로 같은 성능 기준 단가와 다른 지표이며, B2B 계약 사용량은 집계 범위 밖
- 미국 기업 AI 도입률은 집계 기관에 따라 최대 약 3배 차이. 모형은 2026년 8월 56.1%인 결제 기반 Ramp 도입률을 기준으로 채택

그림 14. Silicon Data 지수로 본 토큰 실효 단가의 하락



자료: Silicon Data, Bloomberg, Deutsche Bank, 상상인증권 리서치센터

그림 15. 집계 기관별 미국 기업의 AI 도입률



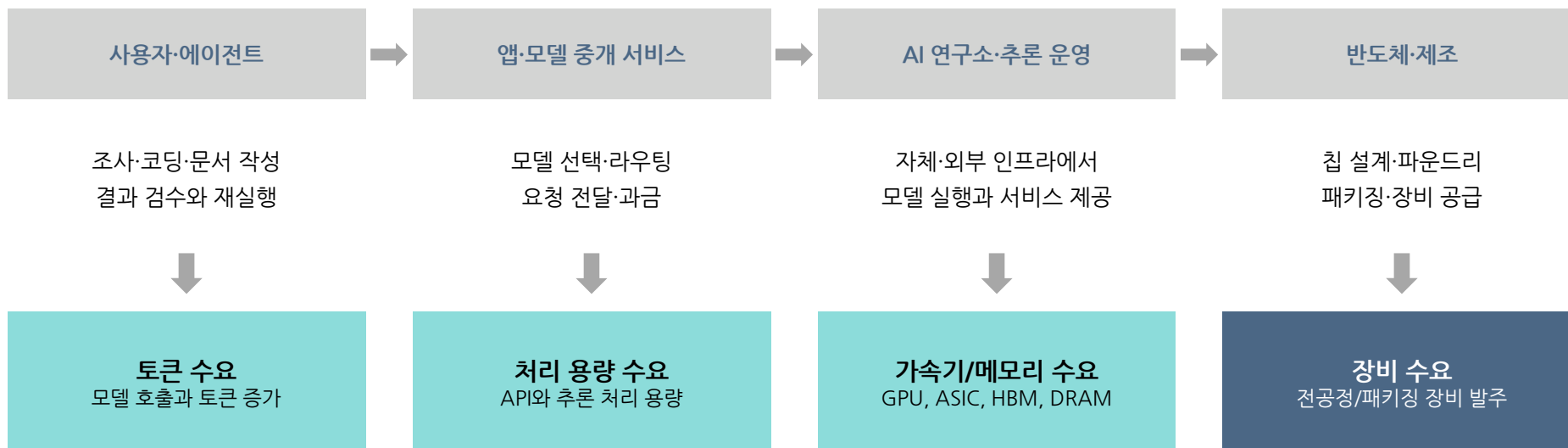
자료: Ramp AI Index, Stripe, 미 인구조사국, Cloud Ratings, 상상인증권 리서치센터

토큰 수요의 확대와 AI 연구소 매출

업무 요청에서 반도체 발주로 이어지는 수요 경로

- 업무 요청의 토큰 수요가 모델 중개와 추론 운영을 거쳐 반도체 발주로 이어지는 4단계 구조
- 에이전트는 모델 호출과 도구 실행을 반복해 요청 한 건도 여러 차례의 추론으로 확장
- 사용자 수, 업무 빈도, 업무당 처리량이 수요의 출발점이며, 속도와 비용 개선은 AI 위임 업무를 확대

그림 16. 업무 실행에서 반도체 발주까지



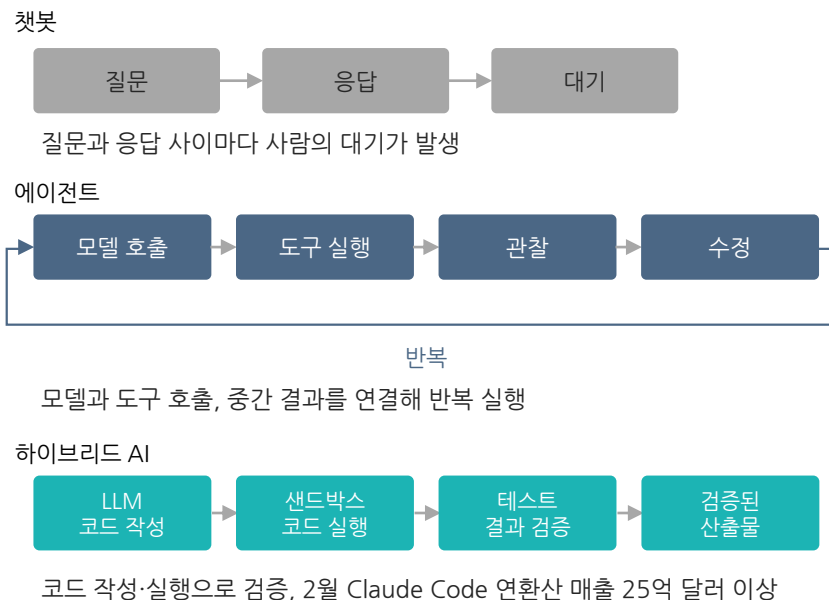
자료: OpenRouter, Anthropic, 상상인증권 리서치센터

토큰 수요의 확대와 AI 연구소 매출

반복 호출로 늘어나는 에이전트 토큰 처리량

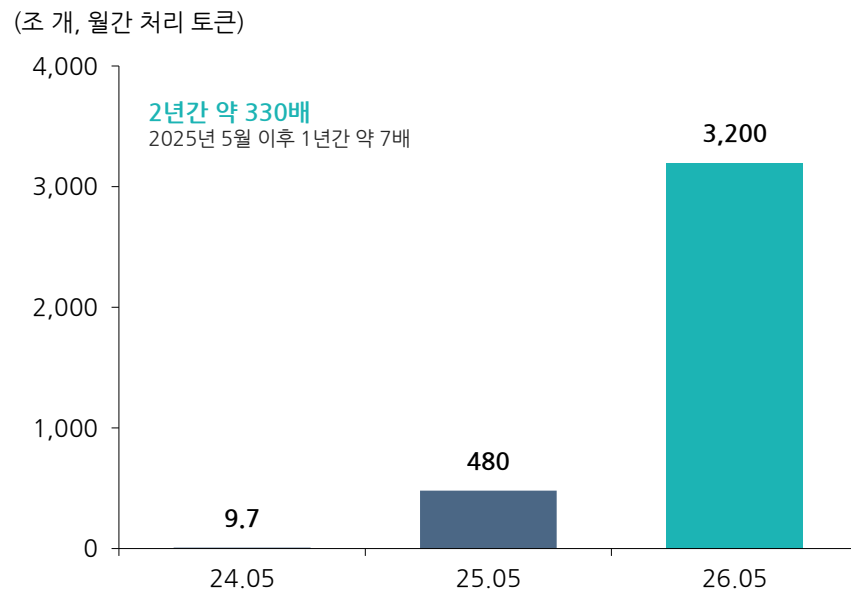
- 에이전트 업무 한 건은 수십 번의 모델 호출로 이어지는 구조. 계획, 도구 실행, 결과 검증은 반복하며 도구 결과를 다시 읽을 때마다 입력 토큰도 추가 발생
- 에이전트 확산으로 사용자 수가 같아도 처리량은 배수로 증가. 처리량은 사용자 수 × 사용자당 업무 수 × 업무당 호출 수 × 호출당 토큰이며, 에이전트는 뒤의 두 항을 함께 늘리는 구조
- 처리량 급증은 실측으로도 확인. Google 전체 서비스의 월간 처리 토큰은 2026년 5월 3,200조 개로 2년간 330배 확대

그림 17. 챗봇·에이전트·하이브리드 AI의 호출 구조



자료: Anthropic, 상상인증권 리서치센터

그림 18. Google 전체 서비스의 월간 처리 토큰



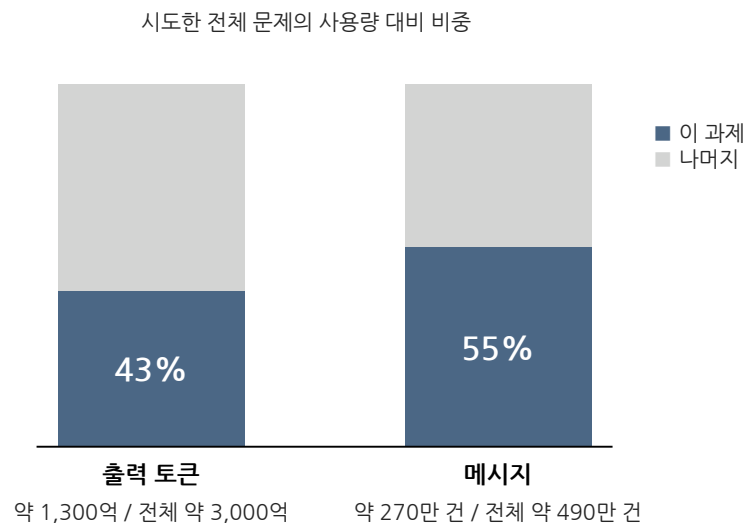
자료: Google, 상상인증권 리서치센터

토큰 수요의 확대와 AI 연구소 매출

병렬 에이전트의 토큰 소비와 속도 향상 상한

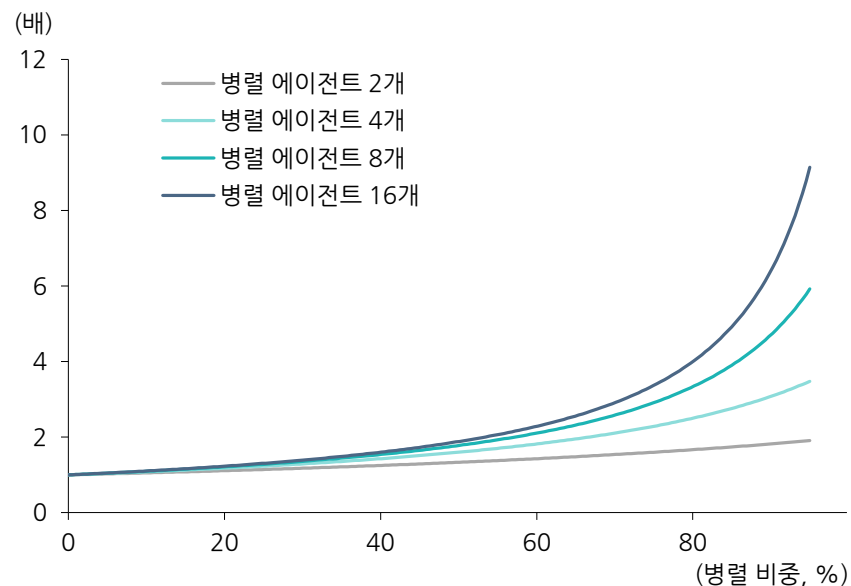
- 병렬 에이전트는 과제 하나에도 막대한 토큰을 소비. OpenAI는 밀레니엄 난제인 Navier-Stokes 문제에 에이전트 약 1만 개를 동시 투입했고, 이 과제가 에이전트 집단 전체 출력 토큰의 43%를 차지
- 병렬화의 속도 향상에는 상한이 존재. 앞 단계 결과를 기다리는 순차 구간 때문에 병렬 비중이 90%이면 에이전트를 무한히 늘려도 속도 향상은 10배가 상한(Amdahl의 법칙)
- 긴 업무의 완료 시간은 에이전트 수보다 순차 구간의 호출 속도가 좌우. 동시 실행을 위한 가속기와 함께 사용자당 토큰 생성 속도를 높이는 메모리 대역폭 수요도 확대

그림 19. 밀레니엄 문제가 차지한 에이전트 집단의 사용량



자료: OpenAI, 상상인증권 리서치센터

그림 20. 병렬화 비중과 속도 향상 한계



자료: 상상인증권 리서치센터

주: Amdahl의 법칙, 속도 향상 = $1 \div (\text{순차 비중} + \text{병렬 비중} \div \text{에이전트 수})$

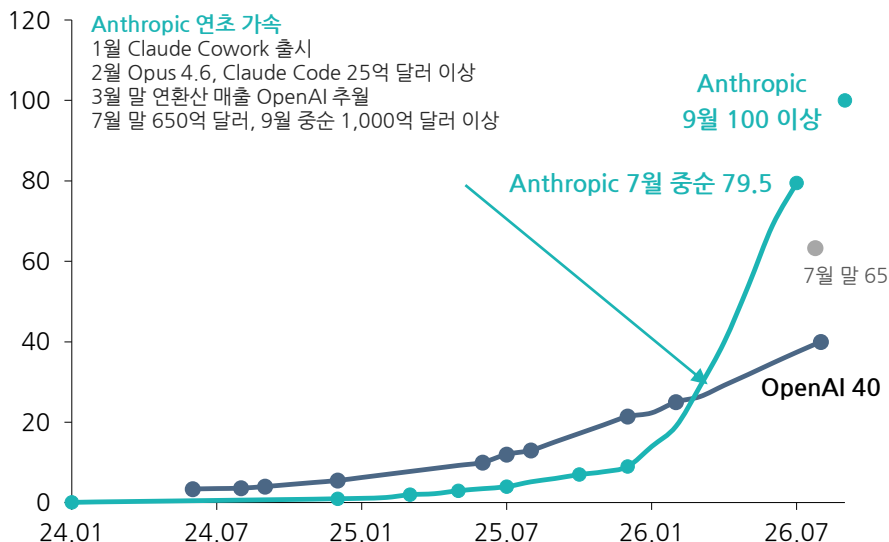
토큰 수요의 확대와 AI 연구소 매출

매출과 기업 도입에서 OpenAI를 추월한 Anthropic

- Anthropic 연환산 매출은 2026년 3월 말 OpenAI를 추월한 뒤 7월 650억~795억 달러에서 9월 중순 1,000억 달러 이상(NYT 보도)으로 상향돼 8월 중순 OpenAI 400억 달러의 2.5배 이상
- 기업 도입률도 Anthropic이 선두. 미국 기업의 Anthropic 유료 도입률은 2026년 5월 41.0%로 OpenAI(39.5%)를 추월
- 추월의 배경은 2026년 초 에이전트와 하이브리드 AI의 상용화. Claude Code 연환산 매출은 2월 25억 달러 이상이며, 업무용 에이전트 Claude Cowork는 Claude Code로 10일 만에 개발

그림 21. AI 연구소의 연환산 매출 추이

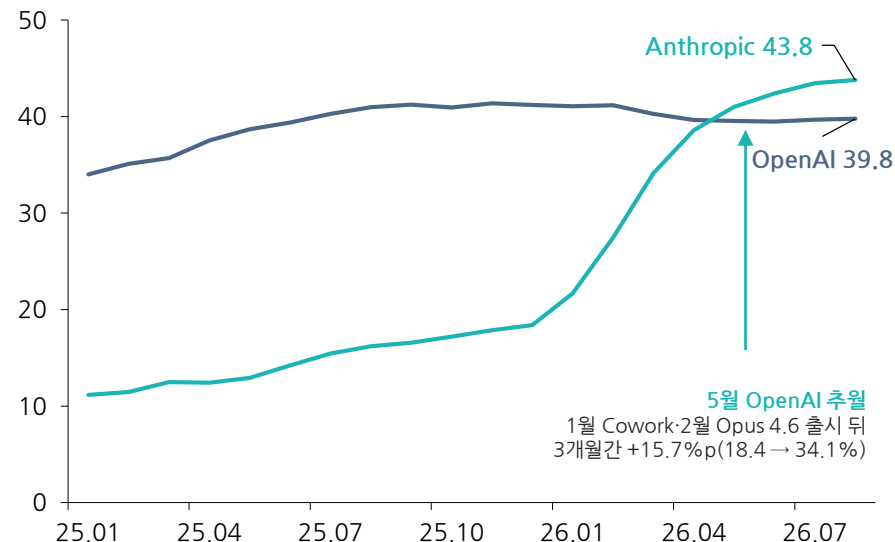
(십억 달러, ARR)



자료: YipitData, Bloomberg, NYT, 각 사, 상상인증권 리서치센터

그림 22. 미국 기업의 AI 모델 유료 도입률

(%)



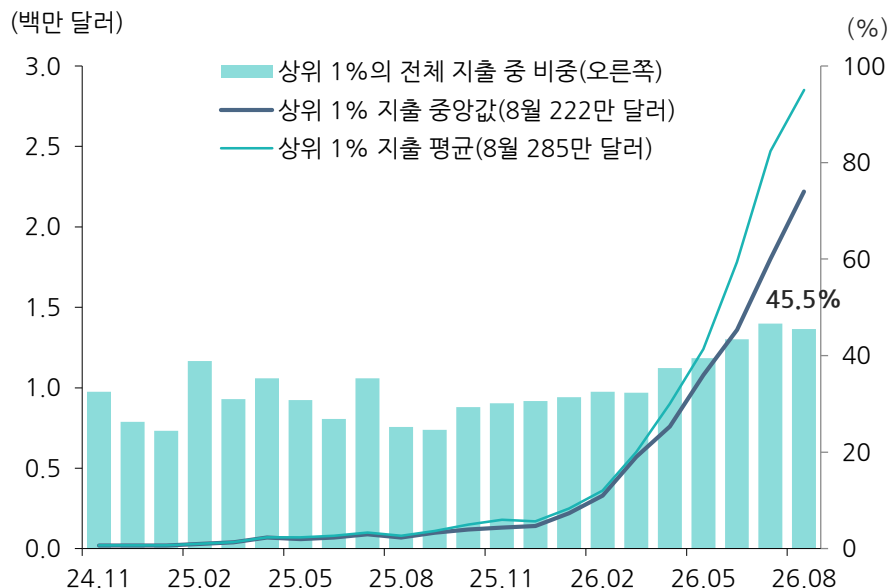
자료: Ramp AI Index, 상상인증권 리서치센터

토큰 수요의 확대와 AI 연구소 매출

상위 대형 고객에 집중된 AI 연구소 매출

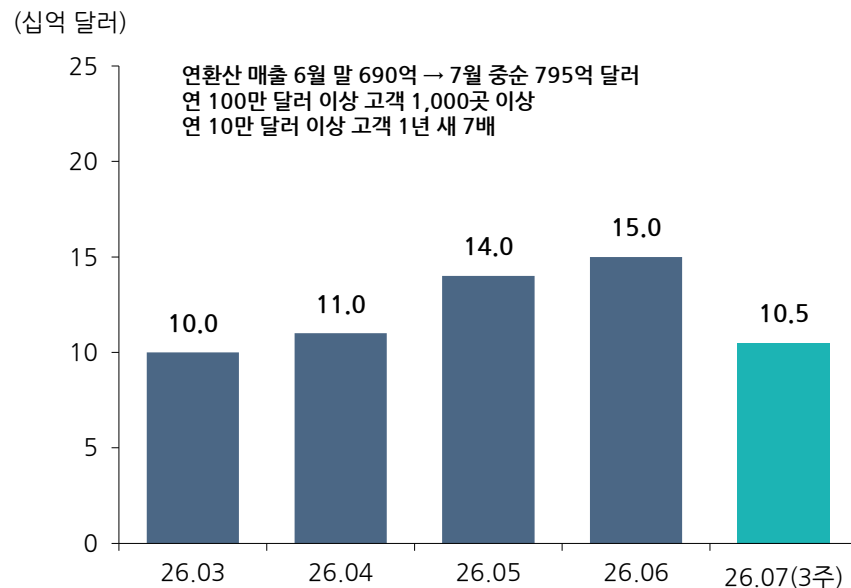
- Anthropic 매출의 대형 고객 집중 심화. YipitData 기준 상위 1% 고객의 지출 비중은 30%대에서 8월 45.5%로 상승
- 기업 간 격차는 도입 여부보다 사용 강도에서 발생. Ramp 기준 8월 직원당 월 AI 지출은 상위 1% 기업이 중위 기업의 약 580배로 대형 고객의 사용량이 매출 증가를 주도
- 고가 모델이 매출 가속에 기여. Fable 5는 Anthropic 토큰의 6.0%지만 지출의 11.4%를 차지하며, 연환산 매출의 월별 증가폭도 3~6월 연속 확대

그림 23. Anthropic 상위 1% 고객의 지출과 전체 지출 중 비중



자료: YipitData, 상상인증권 리서치센터

그림 24. Anthropic 연환산 매출의 월별 증가폭



자료: YipitData, Ramp, Anthropic, PYMNTS, 상상인증권 리서치센터

RSI가 나타났다! 이번엔 진짜예요!

PART 03

추론 수요 확대와 부품별 반도체 수혜

반도체 밸류체인 ① 부품별 수요의 출발점

모델 크기·컨텍스트 확대와 하드웨어 수요

- 모델 규모 재확대와 컨텍스트 증가로 하드웨어 수요 확대. 모델 규모는 2026년 상반기까지 정체 후 재확대 국면이며, MoE도 전체 가중치는 메모리에 보관
- 모델 크기, 입력 길이, 동시 사용자가 늘수록 가속기, 메모리, 네트워크 부담이 함께 확대. 응답속도와 토큰당 원가를 개선하려면 세 부품의 동반 강화 필요
- 사고형(Reasoning) 질의 확산으로 부담 추가 확대. 가정 예시 기준 질의당 평균 처리 부담은 일반 질의의 1.45배에서 25.5배로 증가

그림 25. 수요 확대 요인별 부품 부담

확대 요인		요인별 하드웨어 부담			
기준: 같은 품질과 응답시간에서 필요한 데이터센터 처리량		② 가속기 GPU-ASIC	③ HBM 용량·대역폭	④⑤ DRAM·SSD 상태 저장	⑥ 연결 네트워크
모델 규모	문맥·도구				
가중치·전문가(MoE) 증가 모델 저장과 읽기 부담 확대	긴 입력과 중간 결과 누적 KV 캐시 등 작업 상태 확대	●	●	●	●
사고 과정	동시 사용자				
사고 과정(Reasoning)·병렬 탐색 확대 생성 토큰과 실행시간 증가	복수 세션과 모델 복제 같은 모델을 여러 번 운영	●	●	●	●
필요 처리량 ∝ 토큰당 활성 파라미터 × 요청당 토큰(문맥 + 사고 과정) × 동시 요청 수					

● 주요 부담 ● 부분 부담 빈칸: 영향 작음

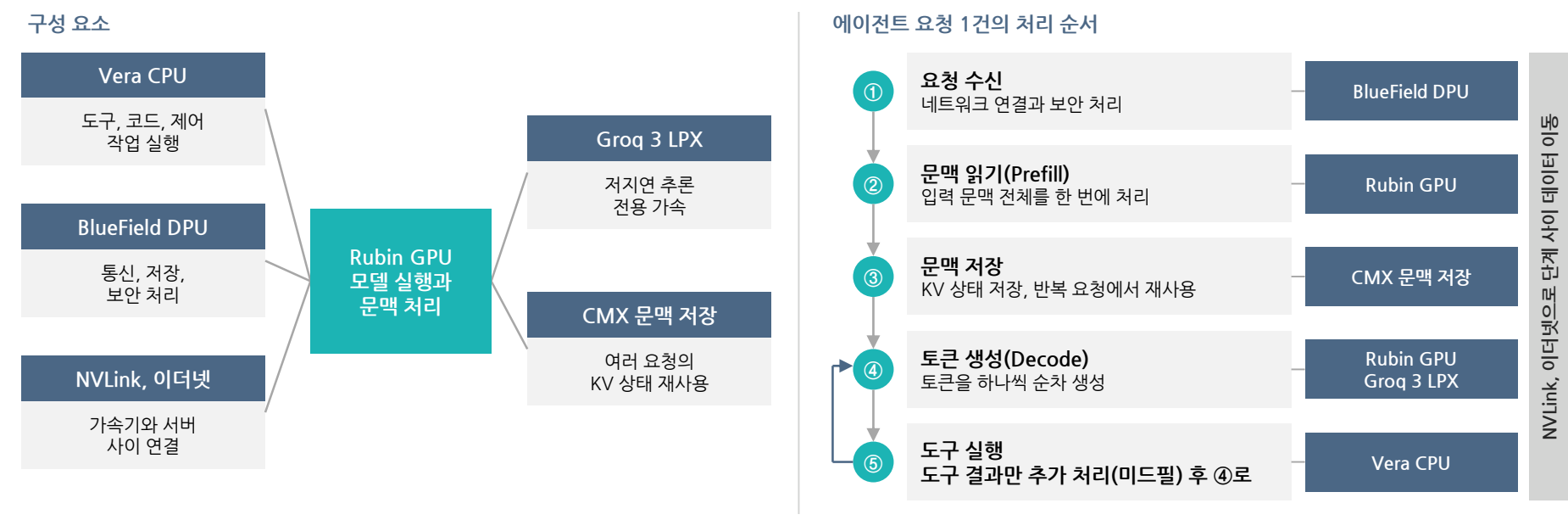
자료: NVIDIA, 상상인증권 리서치센터

반도체 밸류체인 ① 추론 처리 구조

GPU와 CPU·DPU·저장장치의 추론 작업 분담

- 에이전트 서비스 속도는 GPU뿐 아니라 CPU와 DRAM 성능에도 좌우. GPU가 토큰을 생성하는 동안 CPU가 도구 실행과 데이터 작업을 처리해, CPU 지연이 작업 완료시간 증가로 연결
- NVIDIA Vera Rubin은 부품별 분업 구조. GPU가 모델을 실행하고 DPU가 요청을 받으며, KV 캐시 전용 저장장치(CMX)가 반복 요청의 캐시 재사용을 담당
- 추론 시스템의 성능 기준은 같은 모델, 품질, 응답시간에서 지속 공급 가능한 토큰의 양. 모델 크기와 문맥 길이, 동시 사용자에게 따라 필요한 부품 조합이 변화

그림 26. 추론 시스템의 작업 분담

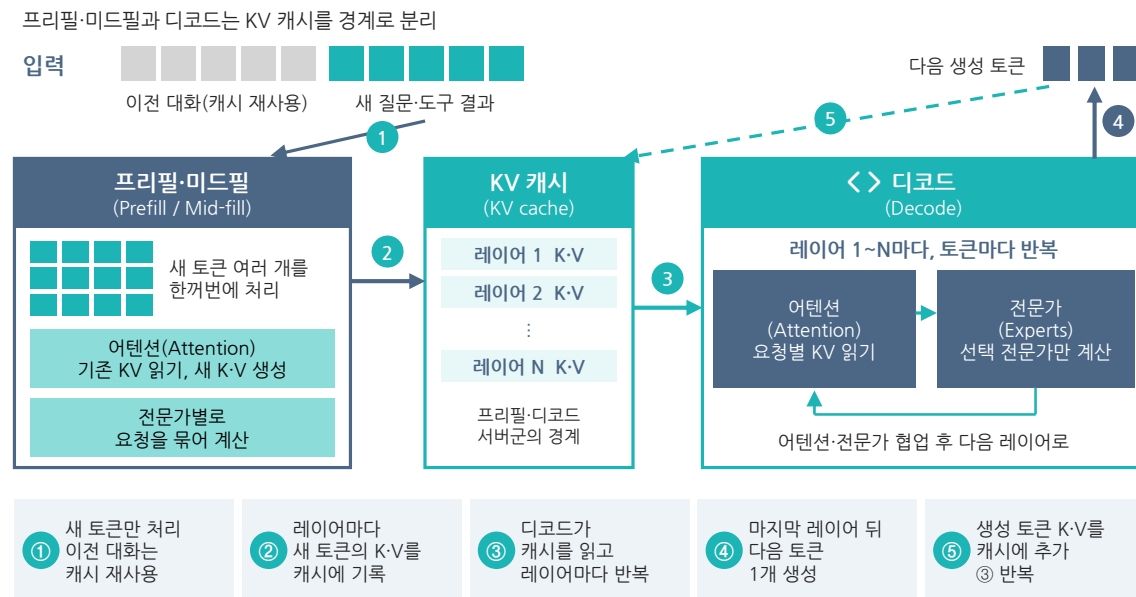


반도체 밸류체인 ① 추론 처리 구조

프리필·미드필·디코드의 단계별 자원 부담

- 프리필은 가속기 처리 성능이, 디코드는 메모리 대역폭이 속도를 좌우. 도구 호출마다 새 입력만 처리하는 미드필이 반복
- 디코드 안에서도 부담의 성격이 상이. 어텐션은 요청마다 자기 KV 캐시를 읽어 대역폭을 소모하고, 전문가(MoE) 계산은 여러 요청을 묶을수록 효율 상승
- 새 입력량, 문맥 길이, 목표 응답시간에 따라 필요한 처리 성능과 메모리 대역폭의 비율이 변화해, 두 자원의 균형이 추론 서비스의 효율과 원가를 결정

그림 27. 토큰의 처리 경로: 프리필·미드필, KV 캐시, 디코드



자료: SemiAnalysis, 상상인증권 리서치센터

그림 28. 단계별로 속도를 좌우하는 자원

프리필 (Prefill)	처리 성능
새 입력을 묶어 한꺼번에 처리 가중치 재사용 효과가 큼	
미드필 (Mid-fill)	처리-메모리
기존 KV를 읽고 새 입력(도구 결과)을 추가	
디코드 어텐션 (Attention)	메모리 대역폭
요청마다 자기 KV를 읽음 문맥이 길수록 부담 증가	
디코드 전문가 (Experts)	처리 성능 (배치 효과)
같은 전문가 가중치를 공유 배치를 키울수록 효율 상승	

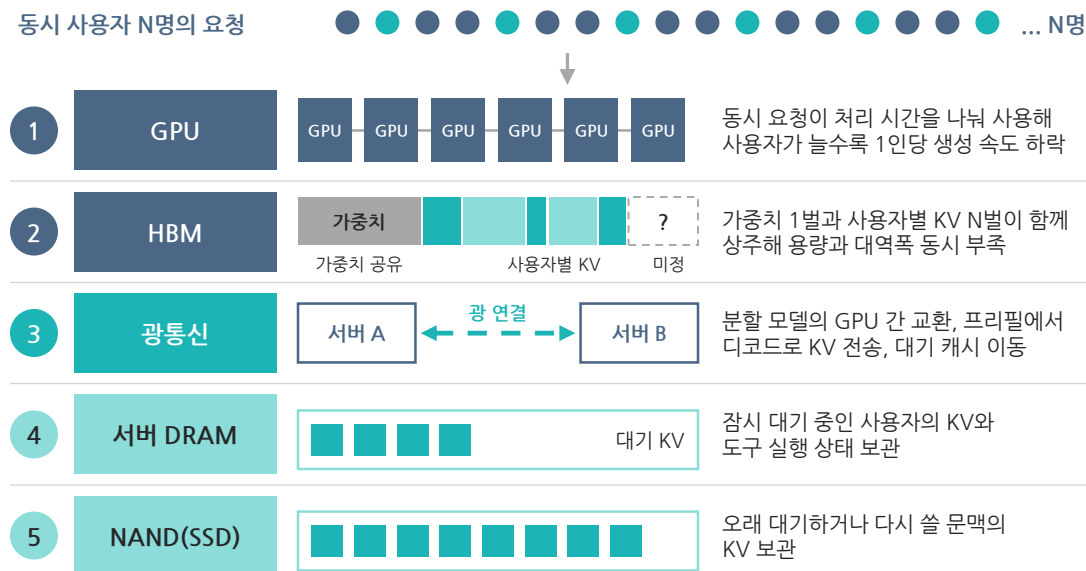
자료: SemiAnalysis, 상상인증권 리서치센터

반도체 밸류체인 ① 추론 처리 구조

동시 서빙 확대에 따른 GPU·HBM 부족 심화

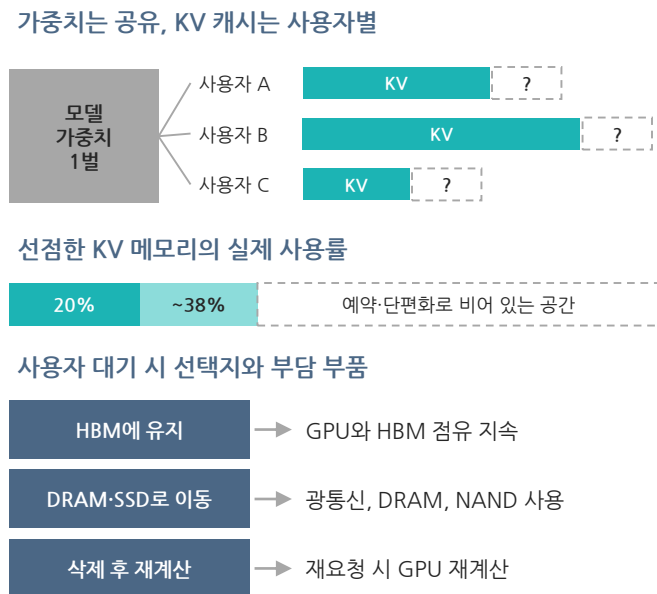
- 대형 모델은 GPU 한 개에 담기지 않아 여러 GPU에 나눠 적재한 뒤 수많은 사용자가 동시에 사용. Kimi K3는 가중치만 약 1.56TB이며, vLLM은 전문가를 GPU 16개에 분산하는 구성으로 서빙을 검증
- 가중치는 모든 사용자가 한 벌을 공유하나 KV 캐시는 사용자별로 따로 보관. 다음 요청 시점과 답변 길이를 예측할 수 없어 기존 서빙 시스템은 선점한 KV 메모리의 20~38%만 실제 문맥 저장에 활용
- 동시 사용자 증가의 부담은 GPU 처리 시간과 HBM 용량에 먼저 집중되고, 대기 캐시를 옮기고 보관하는 과정에서 광통신, 서버 DRAM, NAND로 확산. 서빙 효율 개선의 절감분은 더 많은 동시 세션에 재투입

그림 29. 동시 사용자 증가에 따른 부품별 부하 순서



자료: NVIDIA, Moonshot AI, vLLM, 상상인증권 리서치센터

그림 30. 사용자별 KV 캐시와 대기 중 처리 방식



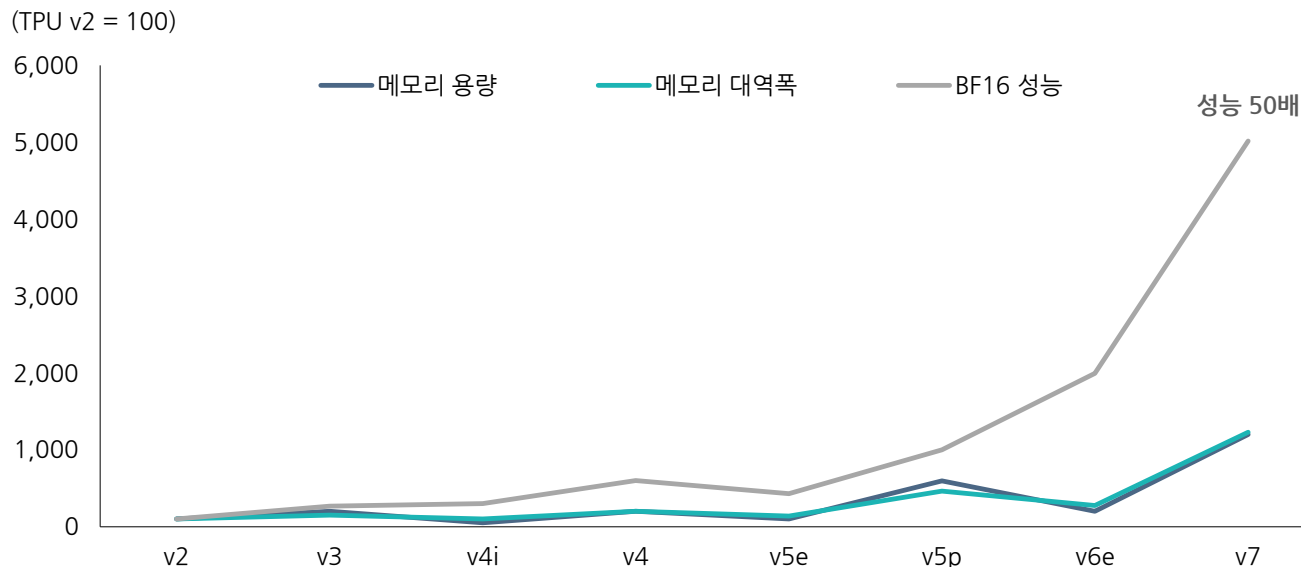
자료: vLLM, NVIDIA, 상상인증권 리서치센터

반도체 밸류체인 ② 가속기(GPU·ASIC)

GPU 성능 향상을 따라가지 못한 메모리 대역폭

- 가속기 세대 교체와 함께 토근 생산의 병목이 메모리로 이동. V100에서 GB200까지 GPU당 이론 처리 성능(FP16/BF16)은 20배 증가했으나 메모리 대역폭은 8.9배에 그침
- Google은 TPU 8세대를 학습용과 추론용으로 나누고 추론용에 메모리 자원을 더 배분. 추론용 8i는 학습용 8t 대비 온칩 SRAM이 3배인 반면 피크 FP4 성능은 0.8배
- 추론 수요 확대가 HBM 탑재량과 온칩 SRAM 면적 확대로 직결되며 파운드리 웨이퍼 수요도 증가

그림 31. TPU 세대별 칩당 사양(v2 = 100)



자료: Epoch AI, 상상인증권 리서치센터

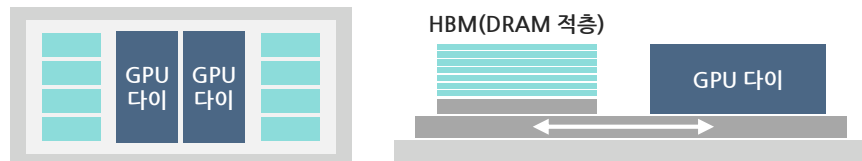
반도체 밸류체인 ② 가속기(GPU·ASIC)

HBM형·SRAM형 추론 칩의 장단점과 작업별 분담

- 디코드 속도의 상한은 메모리 대역폭이어서 가중치 저장 위치가 사용자당 속도를 좌우. GPU와 대부분의 ASIC은 칩 밖 HBM에, Cerebras와 Groq은 칩 안 SRAM에 가중치 배치
- SRAM형은 속도에 유리하나 대형 모델 적재에 불리. Groq 3 LPU는 HBM형 대비 용량이 약 1/5000이나 대역폭은 약 19배로, 대형 모델을 담으려면 칩이 대량 필요
- 추론 칩 시장은 GPU 대체보다 작업별 분담 구조로 확대. NVIDIA가 Rubin GPU와 Groq 3 LPX의 결합 구성을 제시했고, 전력당 성능은 Maia 200 등 일부 ASIC이 B200을 상회

그림 32. 가중치 위치에 따른 메모리 접근 구조

GPU: 가중치가 칩 밖 HBM에 위치



위에서 본 패키지(GPU 다이 2개, HBM 8개)

옆에서 본 단면

매 토큰마다 칩 밖 HBM에서 가중치를 읽어 메모리 대역폭이 속도를 제약 요청을 묶는 배치로 총처리량 확대, 토큰당 원가 하락

Cerebras: 가중치가 웨이퍼 안 SRAM에 상주



웨이퍼 한 장 = 칩 1개

확대: 코어마다 SRAM

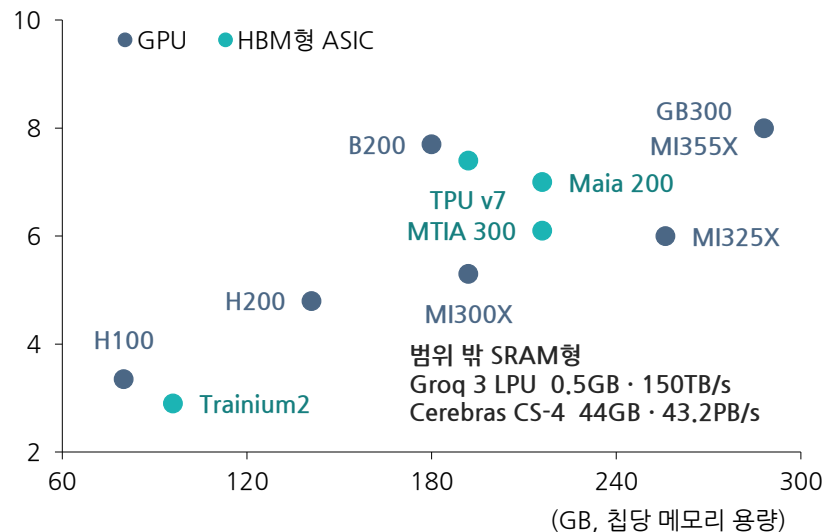
■ 코어 ■ SRAM

CS-4 3개 기준 SRAM 132GB,
대역폭 129.6PB/s로
사용자당 속도 우위
담을 수 있는 모델 크기와
전력, 공급 물량이 제약

자료: Cerebras, 상상인증권 리서치센터

그림 33. 추론 칩의 메모리 용량과 대역폭

(TB/s, 메모리 대역폭)



자료: Epoch AI, NVIDIA, AMD, Groq, Cerebras, 상상인증권 리서치센터

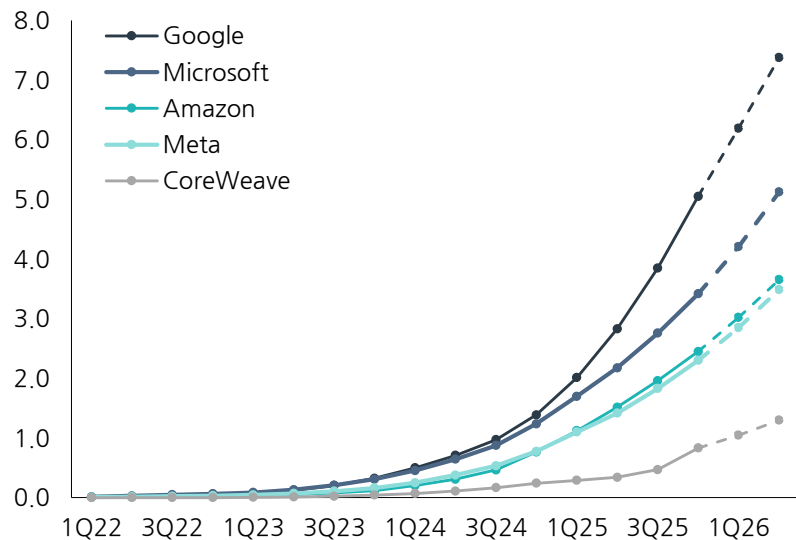
반도체 밸류체인 ② 가속기(GPU·ASIC)

빅테크 가속기 보유량과 AI 연구소의 사용량 급증

- 가속기 보유량은 빅테크 중심. 2025년 말 H100 환산 추정치는 Google 약 505만 개, Microsoft 약 342만 개로 1, 2위
- AI 연구소의 가속기 사용량(클라우드 임차분 포함)은 1년 새 급증. 2024년 말 대비 2025년 말 OpenAI 4.5배(174만 개), Anthropic 5.7배(119만 개)로 확대
- AI 연구소는 외부 클라우드 임차와 전용 데이터센터 구축을 병행해 가속기 확보. Anthropic은 Fluidstack과 미국 데이터센터 건설 투자를 발표

그림 34. 주요 사업자의 가속기 보유량 추정

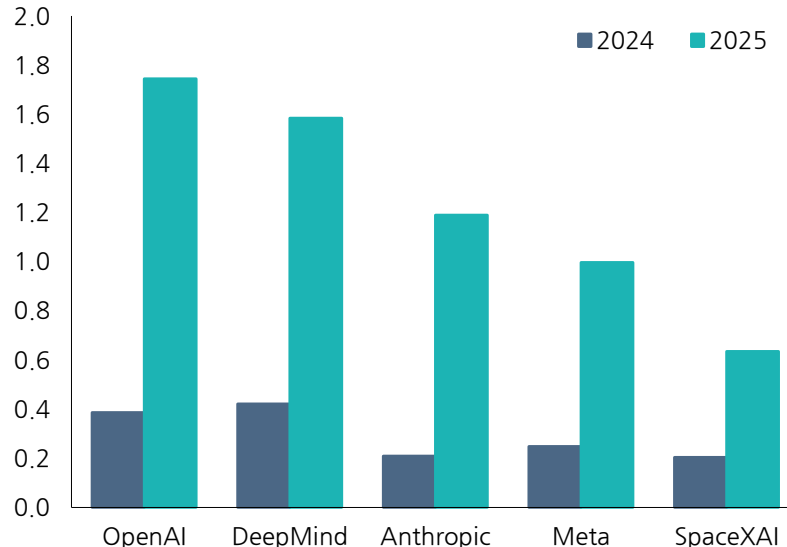
(백만 개, H100 환산)



자료: Epoch AI, 상상인증권 리서치센터
주: 점선은 2025년 신규 확보 비중으로 배분한 추정

그림 35. AI labs의 연말 사용 자원 추정

(백만 개, H100 환산)



자료: Epoch AI, 상상인증권 리서치센터

HBM 수요 확대와 메모리 3사의 차세대 HBM 경쟁

- AI 칩당 HBM 탑재량 증가로 DRAM 공급 여력이 HBM으로 이동할 전망. DRAM 웨이퍼 중 HBM 비중은 2025년 18%에서 2027E 30%로 상승
- HBM4부터 메모리 업체의 설계와 패키징 역량이 경쟁력을 결정. DRAM 적층부와 로직 공정 베이스 다이를 함께 설계하는 비중 확대
- 메모리 3사 모두 차세대 HBM 공급 단계에 진입. SK하이닉스와 삼성전자는 2026년 상반기 HBM4E 12단 샘플을 공급했고, Micron은 3월 HBM4 12단 양산 출하 개시

표 2. HBM 시장 규모와 AI 칩당 탑재량

지표	2025	2026E	2027E
HBM 매출(십억 달러)	35.0	60.0	-
DRAM 매출(십억 달러)	153.5	618.7	903.3
DRAM 매출 중 HBM(%)	22.8	9.7	-
DRAM 비트 중 HBM(%)	8	9	13
DRAM 웨이퍼 중 HBM(%)	18	22	30
HBM 비트 수요 증가율(%)	-	77	68
AI 칩당 HBM(GB)	96~192	216~288	384

자료: TrendForce, Yole Group, 상상인증권 리서치센터

그림 36. 메모리 3사의 최신 HBM 제품 단계

SK하이닉스	HBM4E 12단 2026년 6월 샘플 공급 핀당 최대 16Gbps 속도	샘플
삼성전자	HBM4E 12단 48GB 2026년 5월 샘플 공급 12단 적층으로 스택당 48GB	샘플
Micron	HBM4 12단 36GB 2026년 3월 양산 출하 12단 36GB, 고객 출하 시작	양산

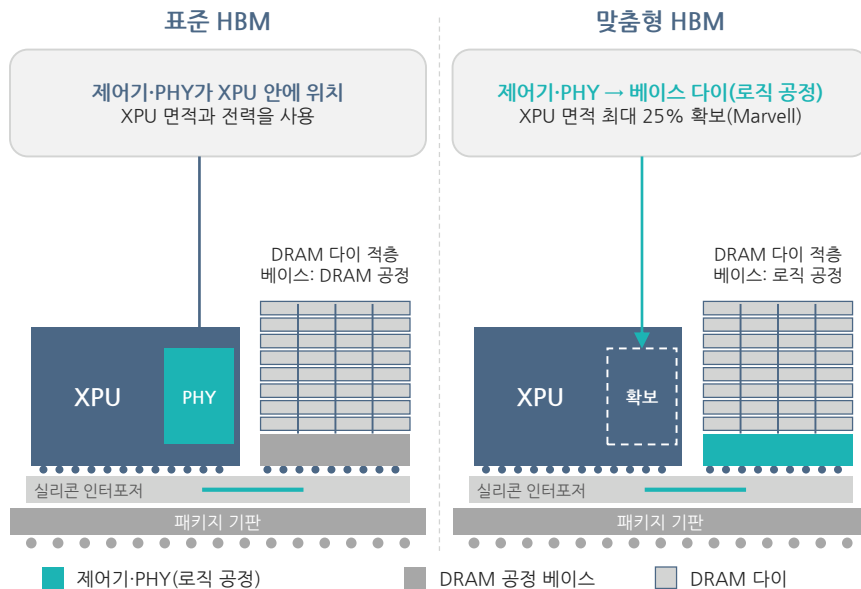
자료: SK하이닉스, 삼성전자, Micron, 상상인증권 리서치센터

반도체 밸류체인 ③ HBM

맞춤형 HBM과 베이스 다이의 기능 확장

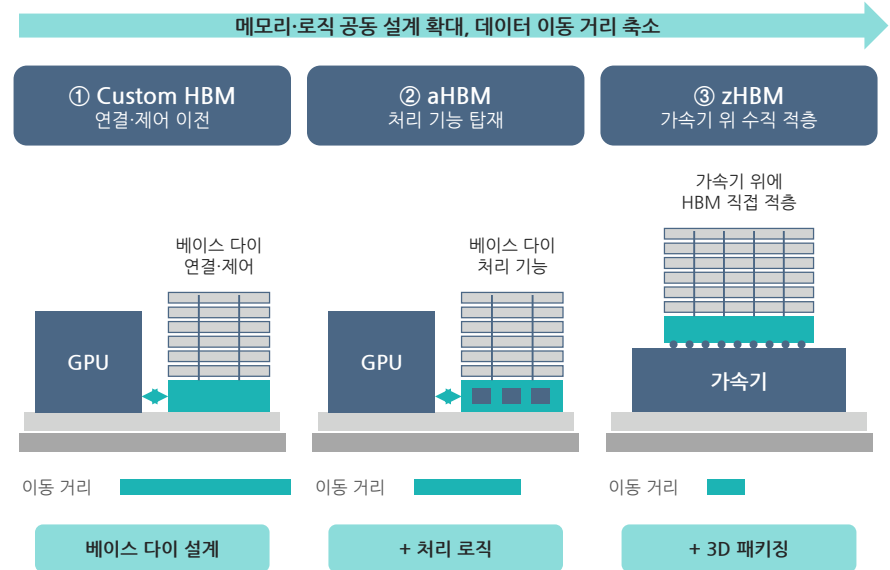
- 맞춤형 HBM은 메모리 제어기와 PHY를 가속기(XPU)에서 베이스 다이로 옮겨 가속기의 로직 면적과 전력을 확보하는 구조. Marvell은 XPU 면적 최대 25% 확보 제시
- 베이스 다이 기능은 연결과 제어(Custom HBM), 처리 기능 탑재(aHBM), 가속기 위 수직 적층(zHBM) 순으로 확장. 단계가 높을수록 데이터 이동 거리 단축
- 메모리 업체가 베이스 다이와 패키징까지 담당하면서 HBM의 부가가치 상승과 고객 락인 효과 강화. 양산의 핵심 과제는 방열, 본딩, 검사, 수율

그림 37. 표준 HBM과 맞춤형 HBM의 패키지 단면



자료: NVIDIA, Marvell, 상상인증권 리서치센터

그림 38. 베이스 다이 기능 확장의 3단계



아래 칸: 메모리 업체가 맡는 설계 범위

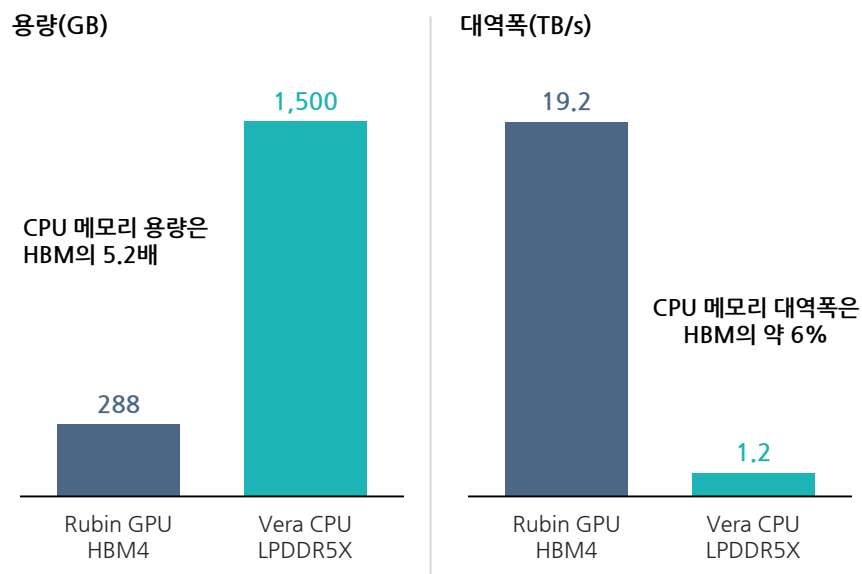
자료: 삼성전자, 상상인증권 리서치센터

반도체 밸류체인 ④ 서버 DRAM

서버 DRAM의 대용량·저전력 전환

- 서버 DRAM은 저전력 LPDDR5X를 탈착형 모듈(SOCAMM)로 장착하는 대용량 구조로 전환. Vera CPU 메모리 용량은 Rubin GPU HBM의 5.2배(1.5TB)
- 용량은 CPU 메모리, 대역폭은 HBM에 맡기는 역할 분담. CPU 메모리 대역폭은 HBM의 약 6%로 도구 실행과 작업 상태는 CPU 메모리에, 가중치는 HBM에 배치
- SOCAMM2는 Vera Rubin 슈퍼칩 원가의 약 50%를 차지하는 핵심 부품으로 메모리 공급 부족 시 탑재 용량 조정 대상. 일반 서버 모듈(RDIMM)도 256GB 샘플로 고용량화

그림 39. GPU HBM과 CPU 메모리의 용량·대역폭



자료: NVIDIA, 상상인증권 리서치센터

표 3. 서버 메모리 제품별 특징과 공급 현황

제품	용도·특징	모듈 용량	공급 현황	2026년 변화
DDR5 RDIMM	범용 서버 CPU의 주 메모리	32~128GB 256GB 샘플	SK하이닉스·Micron 중심 삼성 HBM 우선	채널당 2개 장착 주력 32~64GB 비트 +15~20%
MRDIMM	대역폭 강화 DDR5 모듈	고용량 중심	인터페이스 칩 동반 공급	2027년 양산 전망 모듈당 탑재량 4배 (업계 추정)
LPDDR5X SOCAMM2	Vera CPU용 탈착형 저전력 모듈	96~192GB 256GB 샘플	SK하이닉스·Micron 양산 삼성 1c 공급	슈퍼칩 원가 50% 랙당 54~28TB 축소 전망
CXL CMM-D	CPU 메모리 확장 모듈	-	삼성 CMM-D 2.0 생산 중	3.0(CXL 3.1) 4분기 샘플 2027년 양산
모듈·패키지 기판	SOCAMM·RDIMM용 PCB·기판	-	심텍·티엘비	하반기 SOCAMM 기판 양산 본격화 판가 상승

자료: 메모리 3사, Rambus, TrendForce, SemiAnalysis, DIGITIMES, 상상인증권 리서치센터

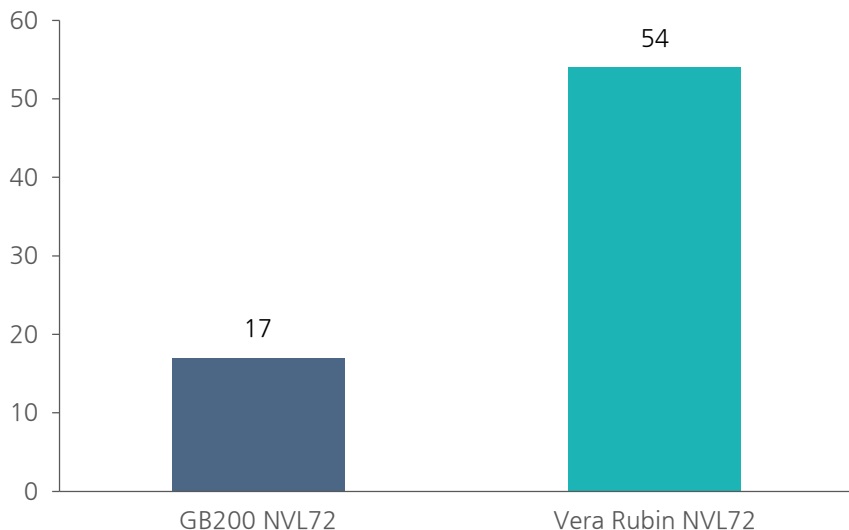
반도체 밸류체인 ④ 서버 DRAM

에이전트 확산의 직접 수혜 품목인 CPU용 DRAM

- 에이전트 확산으로 도구 실행용 CPU 서버와 CPU당 메모리가 함께 증가해 CPU용 DRAM 수요 확대. x86과 Arm 서버 CPU 간 점유율 경쟁과 무관하게 DRAM 탑재량 증가
- CPU 수가 같은 랙 시스템(NVL72)에서 CPU용 DRAM은 GB200 대비 Vera Rubin 세대가 3.2배. 메모리 공급 부족으로 탑재 용량을 줄여도 CPU용 DRAM은 1.6배 증가
- SOCAMM을 양산하는 Micron, SK하이닉스, 삼성전자가 직접 수혜. 업체별 이익 차이는 고객 인증, 제품별 공급 비중, 판매가격이 결정

그림 40. NVL72의 CPU용 DRAM 용량

(TB/랙, CPU 36개 동일, 공개 최대치 기준)



자료: NVIDIA, 상상인증권 리서치센터

표 4. CPU용 DRAM 수요의 확대 경로와 투자 확인 지표

수요 경로	DRAM 확대 요인	투자 확인 지표
실행 서버 증가	도구·코드 실행용 CPU 서버 증설	서버 출하·가동률
CPU당 고용량화	작업 상태·검색 데이터 메모리 적재	CPU당 탑재 용량
고대역폭·저전력	DDR5·LPDDR5X 제품 구성 변화	고객 인증·공급 비중
메모리 이익 증가	비트 출하·판매가격 동반 효과	ASP·원가·증설

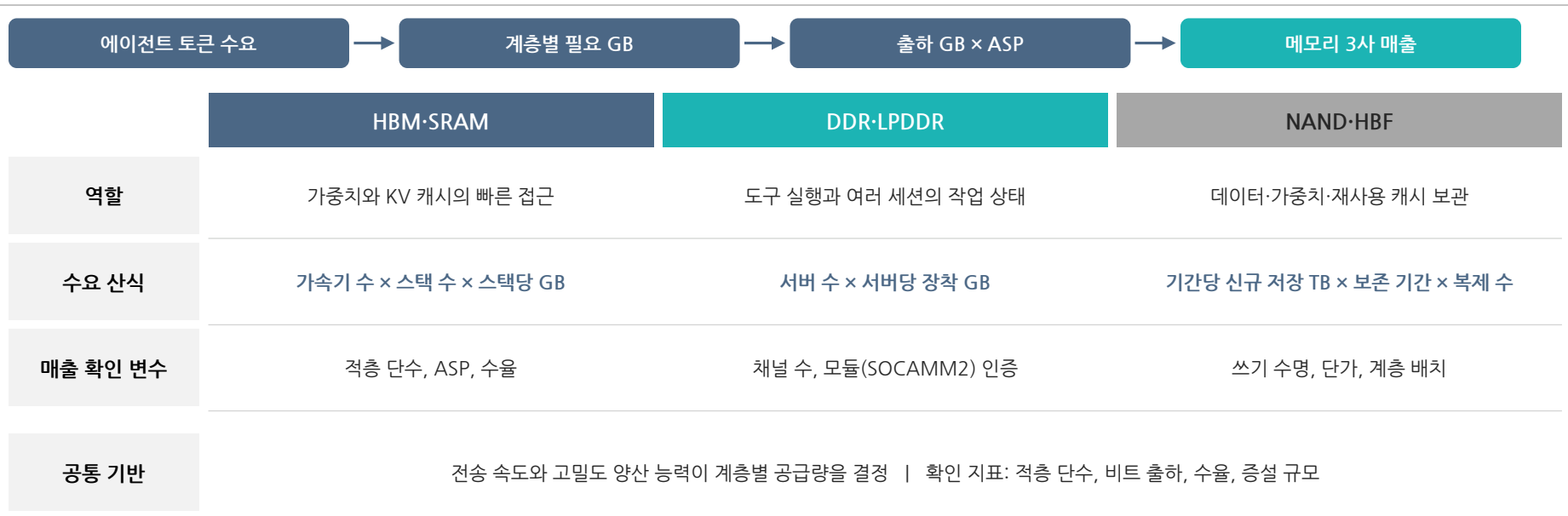
자료: NVIDIA, 상상인증권 리서치센터

반도체 밸류체인 ⑤ KV 캐시·NAND

HBM에서 서버 DRAM·SSD로 확산되는 메모리 수요

- 모델 구조 변화로 HBM에 집중되던 부담이 분산되며 메모리 수요는 HBM 단일 품목에서 서버 DRAM·SSD로 확산
- KV 압축과 공통 문맥의 캐시 재사용에 따른 작업당 비용 하락이 실험과 사용을 함께 늘려 동시 세션, 문맥 길이, 호출 수가 절감 폭 이상으로 증가하므로 순 구매 수요는 확대
- 메모리 매출은 계층별 출하 GB와 단가(ASP)의 곱. 계층별 확인 변수는 HBM 적층 단수, 서버 DRAM 모듈 인증, NAND 쓰기 수명

그림 41. 추론 메모리 계층과 매출 확인 변수



자료: NVIDIA, 상상인증권 리서치센터

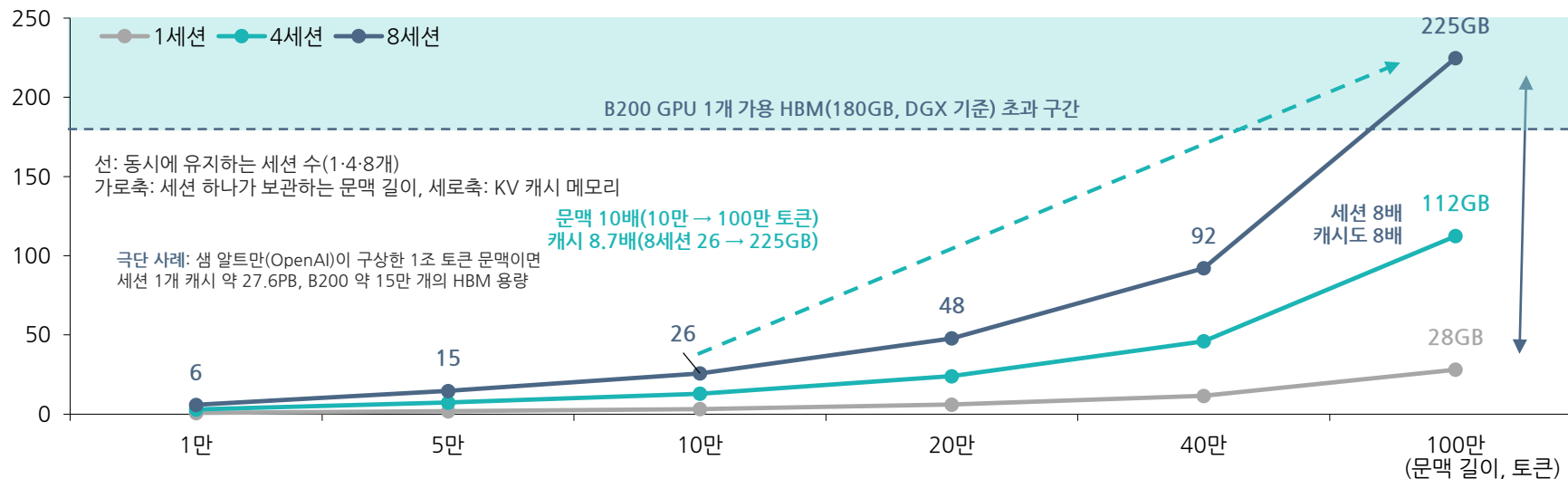
반도체 밸류체인 ⑤ KV 캐시·NAND

문맥 길이·동시 세션에 비례하는 KV 캐시 용량

- KV 캐시는 문맥 길이에 거의 비례해 증가. Kimi K3(파라미터 2.78조 개) 기준 세션 1개의 캐시는 문맥 10만 토큰 3.2GB, 100만 토큰 약 28GB
- 상주 캐시는 동시 세션 수에 정비례. 모든 사용자가 한 벌을 공유하는 가중치와 달리 사용자별 문맥의 캐시는 공유할 수 없어 동시 세션이 8배면 캐시도 8배
- 두 요인이 겹치면 KV 캐시만으로 GPU 1개의 HBM 초과. 100만 토큰 문맥을 8세션 동시에 유지하면 가중치(약 1.56TB, B200 약 9개 분량)와 별도로 약 225GB의 메모리가 필요

그림 42. 문맥 길이와 동시 세션별 상주 캐시

(GB, 상주 캐시 용량)



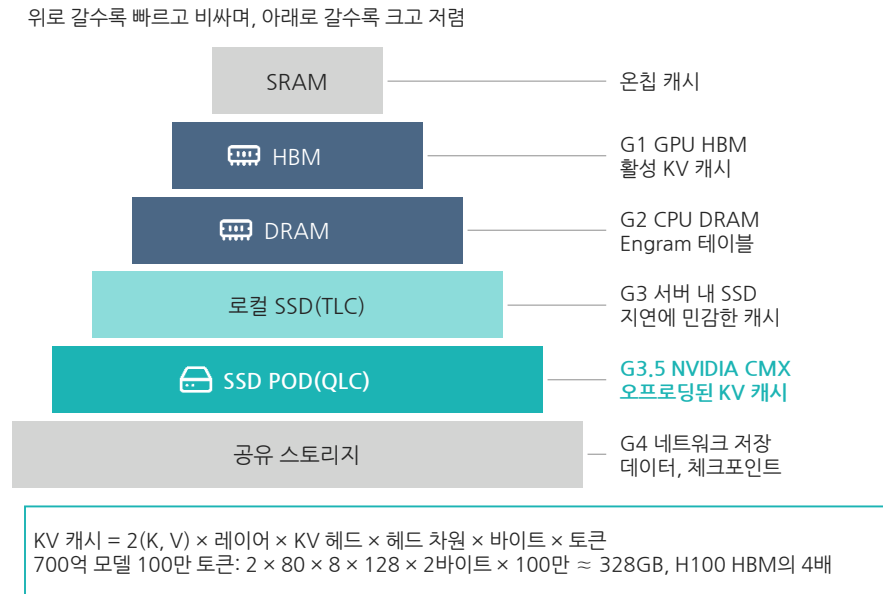
자료: Moonshot AI, 상상인증권 리서치센터

반도체 밸류체인 ⑤ KV 캐시·NAND

KV 캐시 저장 계층 편입에 따른 QLC eSSD 수요 확대

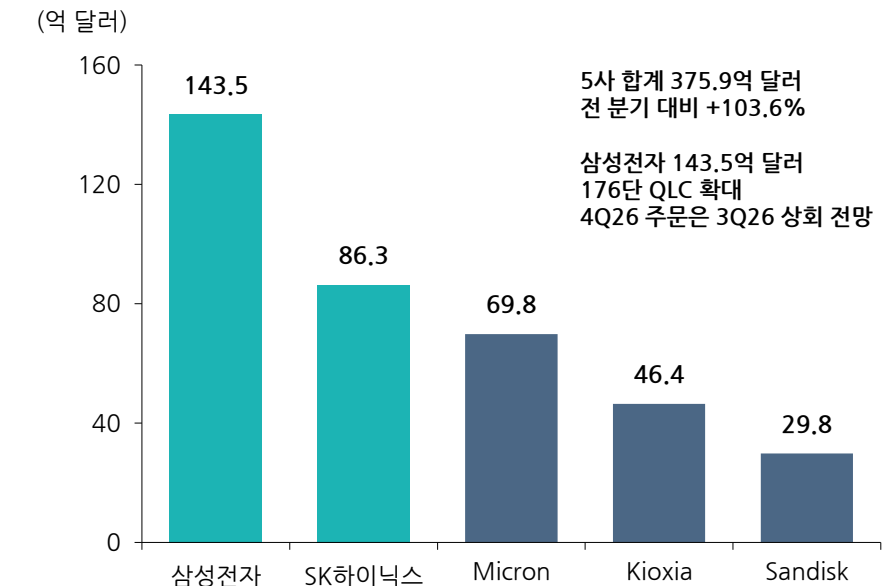
- KV 캐시가 HBM 용량을 초과하면서 SSD가 추론 메모리 계층에 편입. NVIDIA는 KV 캐시 전용 공유 SSD 계층(CMX)을 추가하고 기존 스토리지 대비 처리량 5배 제시
- KV 캐시 저장은 큰 단위의 읽기 중심 작업이라 쓰기 수명보다 용량과 단가가 중요해 QLC가 적합한 용도
- 기업용 SSD(eSSD) 매출은 이미 급증. 2Q26 상위 5사 매출은 375.9억 달러로 전 분기 대비 103.6% 증가해 고용량 QLC 공급력을 갖춘 삼성전자와 SK하이닉스에 유리

그림 43. 추론용 메모리 계층과 KV 캐시 저장 위치



자료: TrendForce, NVIDIA, 상상인증권 리서치센터

그림 44. 2Q26 eSSD 상위 5사 매출



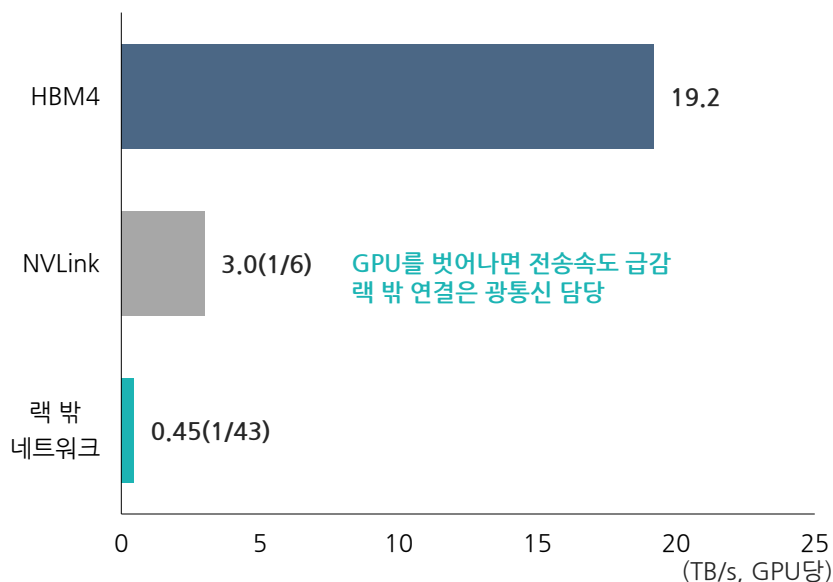
자료: TrendForce, 상상인증권 리서치센터
주: SK하이닉스는 솔리다임 포함

반도체 밸류체인 ⑥ 네트워크·광통신

모델 분할과 GPU 밖 전송 병목에 따른 광통신 수요

- 모델 대형화로 가중치를 여러 GPU에 나눠 적재하면서 GPU 간 전송속도가 응답속도를 좌우. Kimi K3는 가중치만 약 1.56TB로 B200 약 9개 분량이며, 분할된 GPU는 레이어와 전문가 계산마다 중간 결과를 교환
- HBM 대역폭과 GPU 밖 연결의 격차가 새 병목. Rubin GPU당 HBM4 19.2TB/s 대비 NVLink는 3TB/s(1/6), 랙 밖 네트워크는 0.45TB/s(1/43)에 그쳐 전송 대기 동안 GPU와 HBM의 가동률 저하
- 광통신은 랙 밖 GPU 연결, 프리필에서 디코드 서버로의 KV 전송, 대기 캐시 이동을 담당해 1.6T 전환이 GPU와 HBM의 가동률을 높이는 조건. 800G 이상 광트랜시버 출하는 2026E 약 6,300만 개로 전년의 2.6배

그림 45. Rubin GPU당 연결 경로별 전송속도



자료: NVIDIA, 상상인증권 리서치센터

그림 46. GPU 안팎 연결 구간과 광통신의 역할

연결 구간별 매체와 옮기는 데이터

GPU ↔ HBM	패키지 안 인터포저	가중치와 활성 KV 읽기 19.2TB/s
GPU ↔ GPU	랙 안 NVLink 구리 연결	분할 모델의 중간 결과 교환 3TB/s
Rack ↔ Rack	광트랜시버 800G에서 1.6T로	KV 전송, 대기 캐시 이동 0.45TB/s

광 링크 부품과 집적 방향



플러거블, 실리콘 포토닉스, CPO 순으로 광학부가 칩에 가까워져 전력 절감

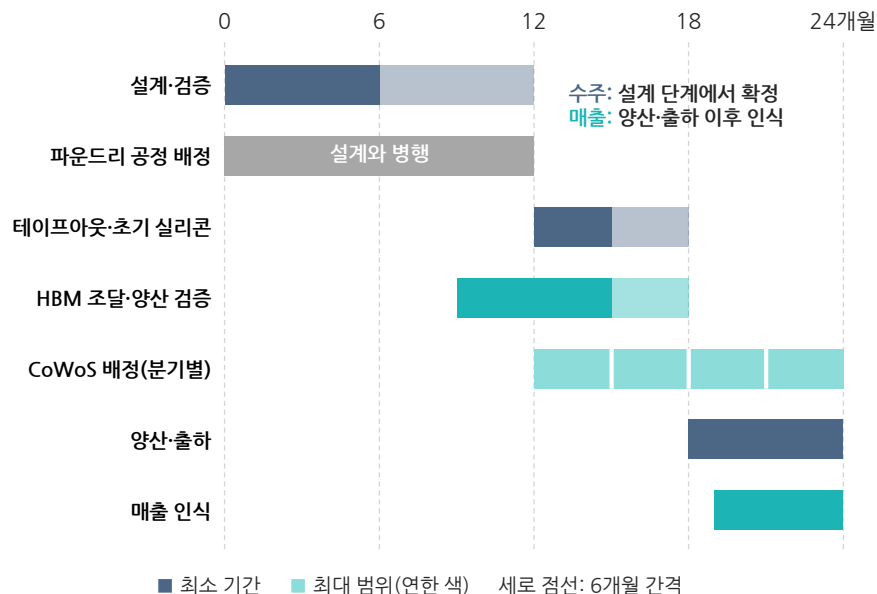
자료: NVIDIA, TrendForce, Lumentum, 상상인증권 리서치센터

반도체 밸류체인 ⑦ 제조·패키징

반도체 매출의 선행 지표인 CoWoS·HBM 배정

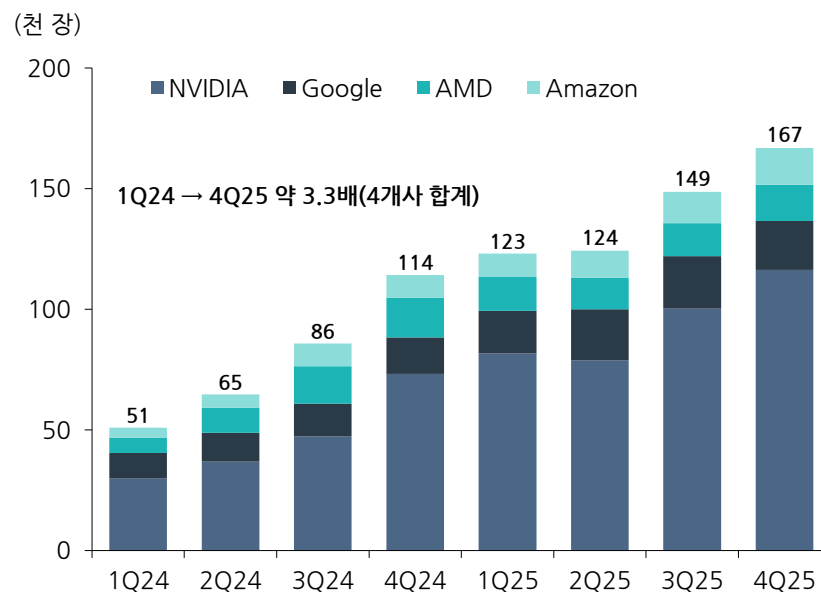
- 현재의 설계 수주와 CoWoS·HBM 배정은 1~2년 뒤 반도체 매출을 예고하는 선행 지표. ASIC은 설계에서 출하까지 12~24개월이 걸려 수주와 매출 인식 사이에 시차 발생
- 선행 지표인 CoWoS 배정은 이미 급증. 4개 설계사 합계 분기 웨이퍼 배분은 4Q25 16.7만 장으로 1Q24의 약 3.3배
- 가속기, 메모리, 네트워크 칩 사양 확정이 설계와 파운드리·메모리 생산 수요로 연결되고, 장비·소재 수요는 공정 전환과 증설 물량에서 파생되는 구조

그림 47. ASIC 설계부터 출하·매출까지의 일정(개념도)



자료: 상상인증권 리서치센터

그림 48. 설계사별 CoWoS 웨이퍼 배분(분기)



자료: Epoch AI, 상상인증권 리서치센터
주: 추정 중앙값, 4개 설계사 합계(기타 제외)

RSI가 나타났다! 이번엔 진짜예요!

PART 04

AI 설비투자자과 매출·현금 회수의 시차

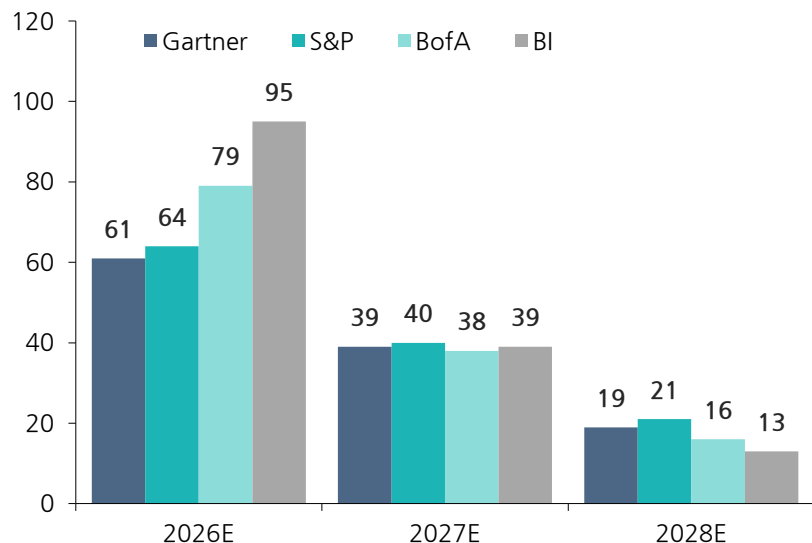
AI 설비투자자과 매출·현금 회수의 시차

투자 증가율 둔화에도 유지되는 CAPEX 증가액

- 증가율이 둔화돼도 CAPEX 증가액은 유지되는 구조. 하이퍼스케일러 5사의 증가액은 2028E에도 1,330억~1,830억 달러로 2025년과 유사한 규모
- 기관별 AI 투자 증가율 전망은 2026E 61~95%에서 2028E 13~21%로 둔화. 다만 CAPEX 규모 자체가 커져 낮은 증가율로도 큰 증가액 발생
- 증가액이 유지되는 만큼 투자 지속 여부는 새 설비의 가동과 과금 개시, 현금 회수 속도가 좌우할 전망

그림 49. 기관별 AI 투자 증가율 전망

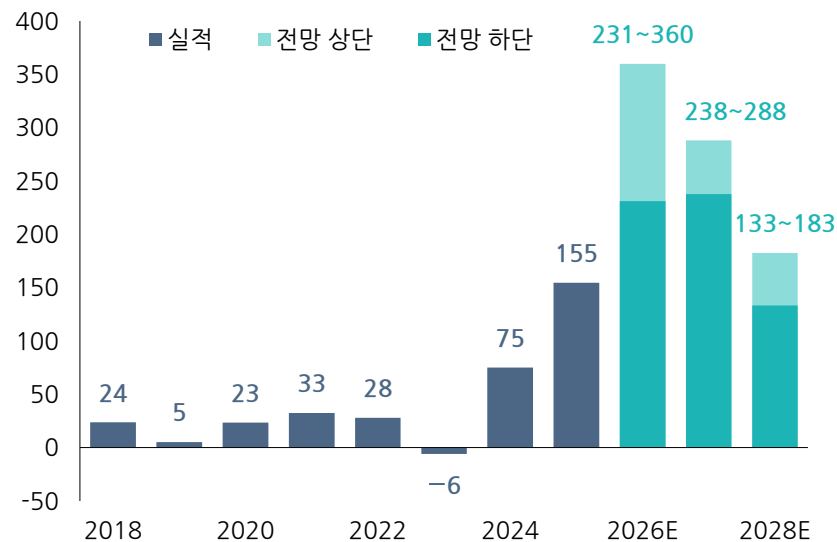
(%, 전년 대비 AI 투자 증가율)



자료: 한국은행(Gartner, S&P, BofA, Bloomberg Intelligence 인용), 상상인증권 리서치센터

그림 50. 하이퍼스케일러 5사 CAPEX의 연간 증가액

(십억 달러, 5개 기업 CAPEX의 전년 대비 증가액)



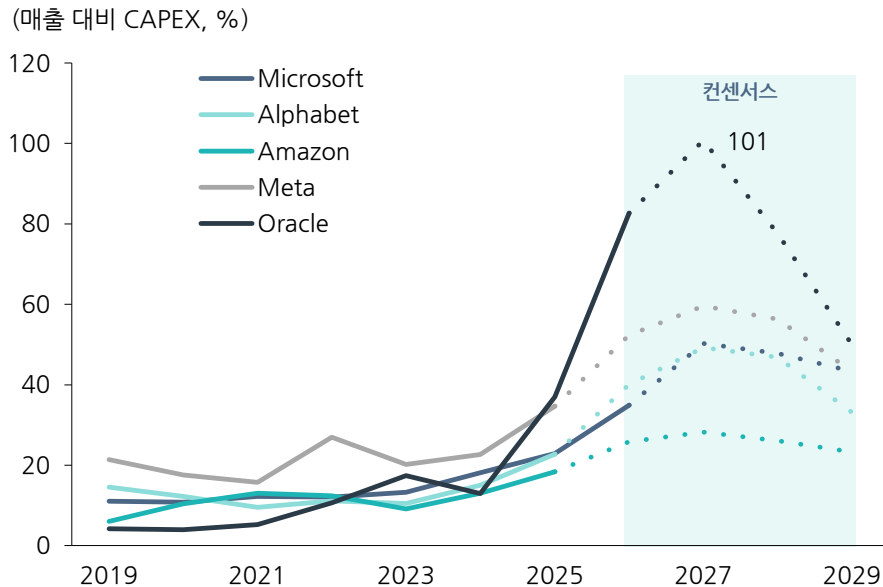
자료: LSEG, 한국은행, 상상인증권 리서치센터

AI 설비투자자과 매출·현금 회수의 시차

하이퍼스케일러 CAPEX 부담과 FCF 마진 악화

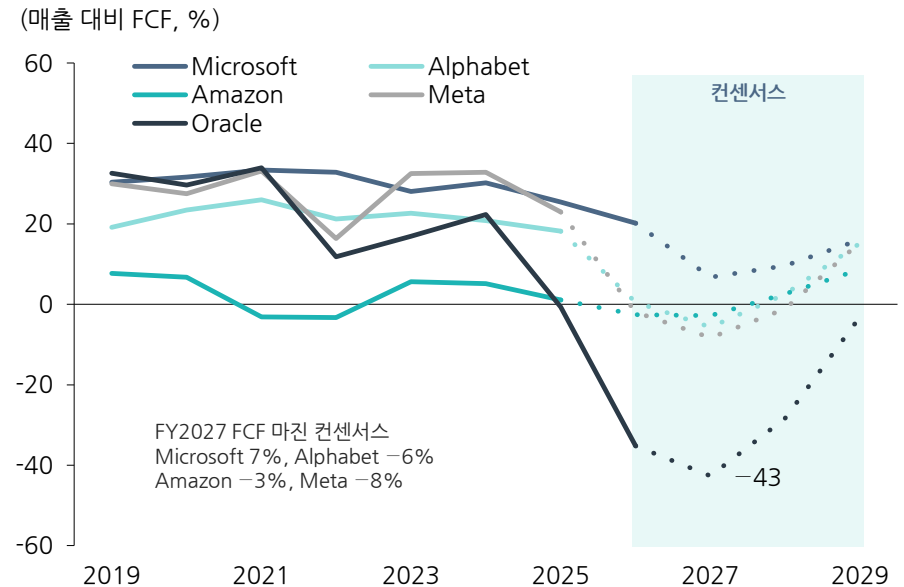
- CAPEX 급증으로 하이퍼스케일러의 FCF 마진 악화. FY2027 컨센서스 기준 Microsoft(7%)를 제외한 4사가 음수이며 Oracle은 -43%
- 매출보다 빠른 CAPEX 증가가 원인. 매출 대비 CAPEX는 2019~2023년 대부분 20% 이하에서 FY2027 컨센서스 28~101%로 상승
- 컨센서스는 FY2028~2029 CAPEX 비중 하락과 함께 FCF 마진이 회복되는 경로를 반영

그림 51. 하이퍼스케일러의 매출 대비 CAPEX



자료: LSEG, 상상인증권 리서치센터

그림 52. 하이퍼스케일러의 FCF 마진



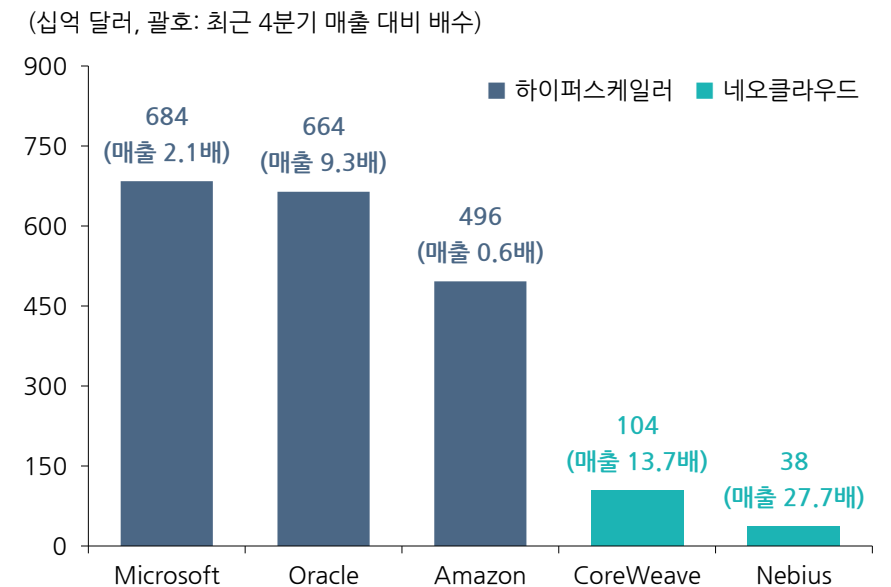
자료: LSEG, 상상인증권 리서치센터

AI 설비투자자와 매출·현금 회수의 시차

RPO의 매출 전환과 과금 개시 전 자금 부담

- RPO는 설비 가동과 고객 검수를 거쳐야 과금이 시작되는 수주잔고 성격
- 매출 대비 계약 잔액이 클수록 설비 가동 전 선투자 부담 확대. 최근 4분기 매출 대비 RPO는 Oracle 9.3배, 네오클라우드 CoreWeave 13.7배
- 과금 개시 전까지 이자와 설비 비용을 감당할 현금과 조달 여력이 핵심. 최근 4분기 FCF는 Microsoft만 670억 달러 흑자, Oracle은 287억 달러 적자

그림 53. 기업별 잔여수행의무(RPO) 잔액



자료: 각 사, 상상인증권 리서치센터

표 5. 기업별 RPO·CAPEX 배수와 자금 여력(최근 4분기, 십억 달러)

기업	RPO/매출 (배)	CAPEX/매출 (%)	RPO/CAPEX (배)	FCF	현금·단기투자	총차입금
Microsoft	2.1	35	5.9	67.0	76.8	106.9
Oracle	9.3	105	8.8	-28.7	37.1	134.5
Amazon	0.6	22	2.9	-11.6	123.0	155.3
CoreWeave	13.7	271	5.1	-13.7	5.5	35.3
Nebius	27.7	822	3.4	-5.9	8.0	8.5

자료: 각 사, LSEG, 상상인증권 리서치센터

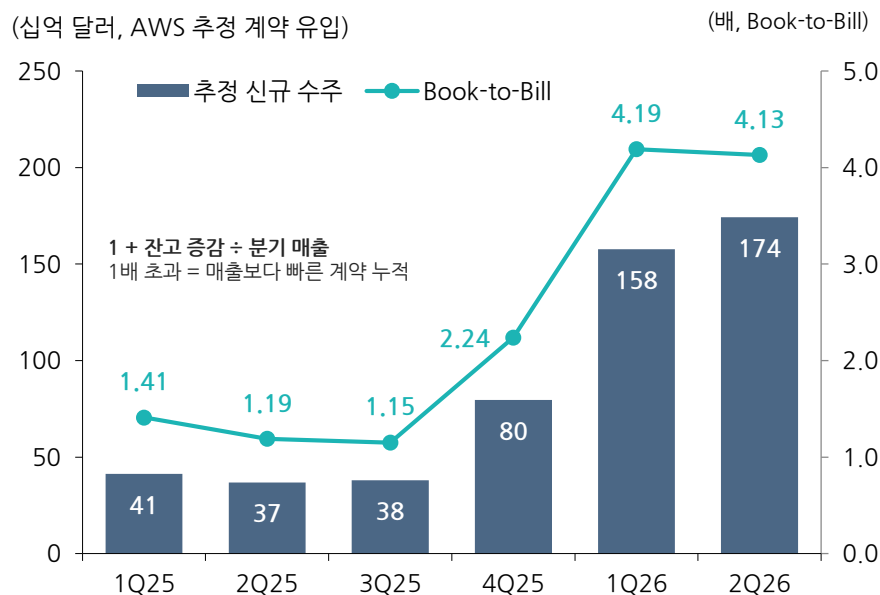
주: RPO/CAPEX = RPO ÷ 최근 4분기 CAPEX

AI 설비투자자와 매출·현금 회수의 시차

매출보다 빠르게 누적되는 클라우드 신규 수주

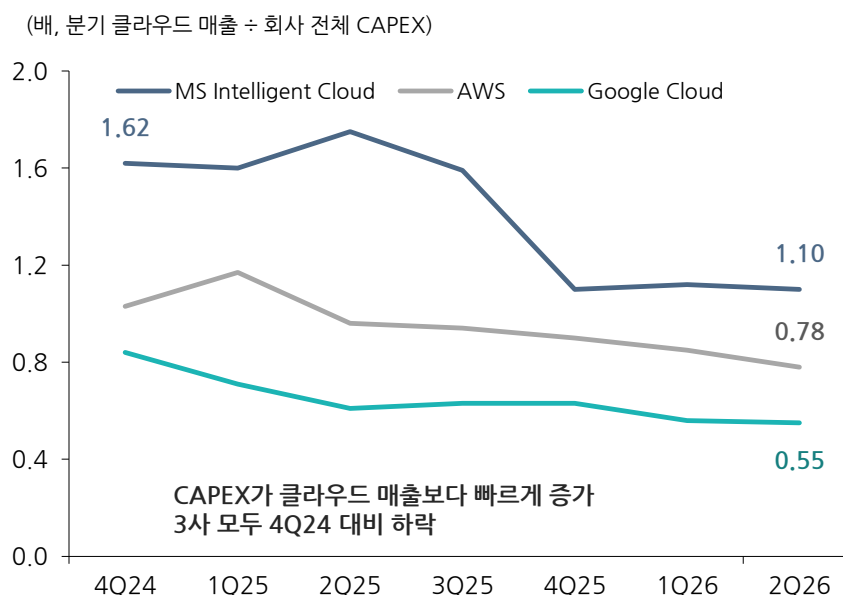
- AWS와 Google Cloud 두 회사 모두 대형 장기계약이 현재 매출보다 빠르게 누적. 2Q26 추정 Book-to-Bill은 AWS 4.13배, Google Cloud 3.08배
- 신규 수주는 미공시라 분기 매출에 계약잔고 증감을 더해 추정. AWS 추정 신규 수주는 1Q25 413억 달러에서 2Q26 1,742억 달러로 확대
- 회사 전체 CAPEX가 클라우드 매출보다 빠르게 늘어 3사 모두 CAPEX 대비 클라우드 매출 배수가 4Q24보다 하락. 2Q26 AWS 0.78배, Google Cloud 0.55배

그림 54. AWS 추정 신규 수주와 Book-to-Bill



자료: Amazon, 상상인증권 리서치센터
주: 추정 신규 수주 = 분기 매출 + 잔고 증감

그림 55. 회사 전체 CAPEX 대비 분기 클라우드 매출



자료: 각 사, 상상인증권 리서치센터

AI 설비투자자와 매출·현금 회수의 시차

DDR5 수익성 상승이 이끄는 HBM 가격 인상

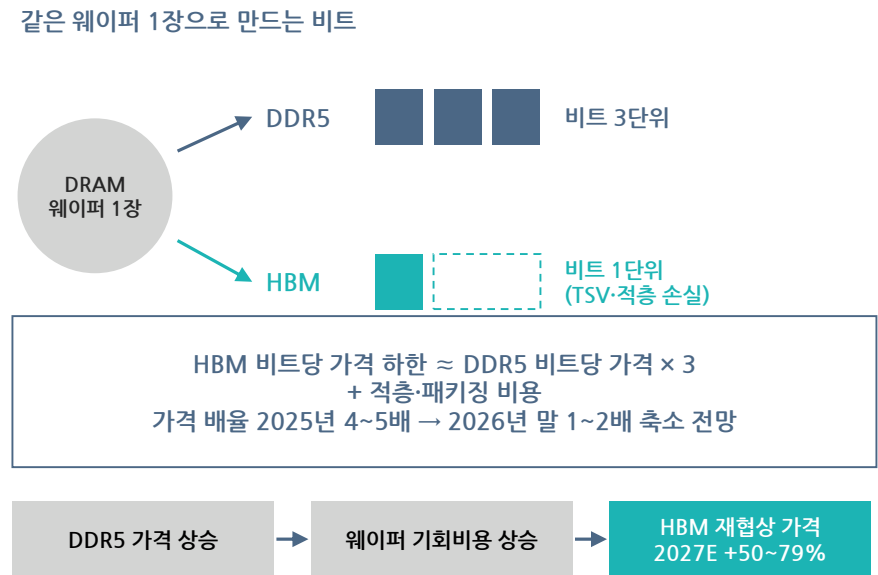
- DDR5 수익성 상승이 HBM 가격 인상으로 연결. 1Q26 DDR5의 웨이퍼당 수익성이 HBM을 추월해 2027E HBM 계약가는 전년 대비 50~79% 상승 전망
- HBM은 같은 비트에 DDR5의 약 3배 웨이퍼가 들어 가격 하한도 DDR5의 약 3배. DDR5 급등으로 두 제품의 가격 배율이 2026년 말 1~2배로 축소 전망이라 HBM 가격 인상 여지 확대
- 추가 투자 없는 기존 설비의 판가 개선이라 신규 증설보다 먼저 이익에 반영될 전망. 다만 현금 반영 시점은 가격 재설정 주기와 선수금 조항에 좌우

그림 56. 메모리 가격 협상에서 현금 회수까지



자료: TrendForce, SK하이닉스, Micron, Sandisk, 상상인증권 리서치센터

그림 57. 웨이퍼 기회비용과 HBM 가격 하한(개념도)



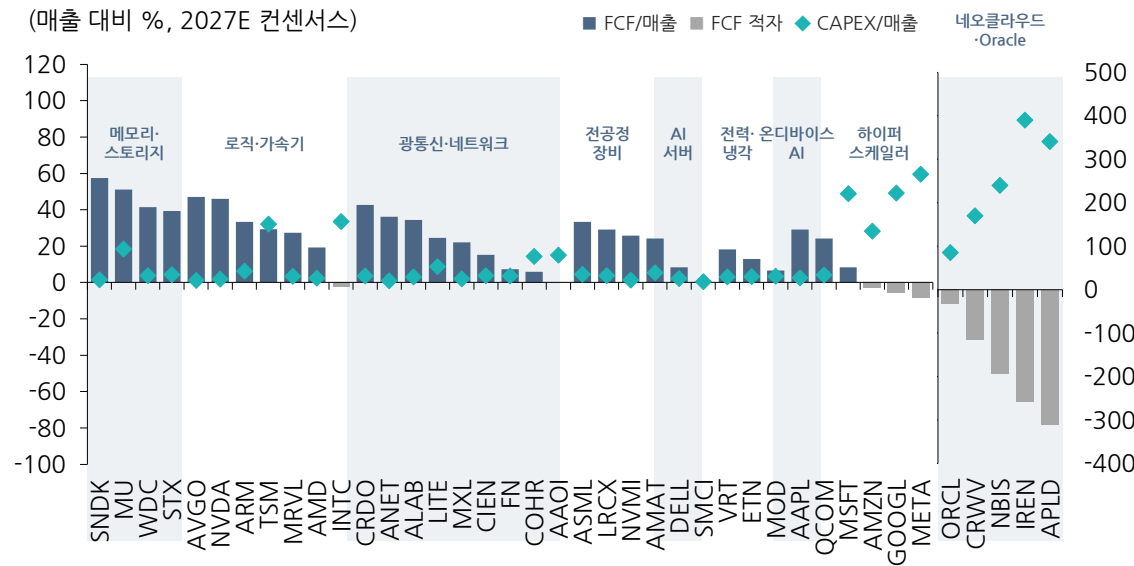
자료: Micron, TrendForce, UBS, 상상인증권 리서치센터

AI 설비투자자와 매출·현금 회수의 시차

구매자 투자 부담이 공급사 현금으로 이전되는 구조

- 구매자의 투자 부담과 공급사의 현금 창출이 대비되는 구조. 2027E 매출 대비 FCF는 메모리·스토리지 39~57%인 반면 네오클라우드는 -116~-310%
- 메모리 기업의 현금 급증도 같은 흐름. 투자 후 현금 삼성전자 2Q26 90.3조원, SK하이닉스 상반기 58.8조원이나 선수금·법인세 납부 시차 등 일회성 항목은 분리 평가 필요
- 구매자의 현금 부담이 반도체·장비 기업의 FCF로 이전되는 구조는 2027E에도 지속 전망. 다만 구매자의 조달 여건 악화는 공급사 이익의 하방 요인

그림 58. 40개 기업의 2027E 매출 대비 FCF와 CAPEX(테마별)



자료: LSEG, 상상인증권 리서치센터

표 6. 메모리 기업의 투자 후 현금

항목	SK하이닉스 (조원)	삼성전자 (조원)	Micron (십억달러)	Sandisk (십억달러)
투자 후 현금	58.8	90.3	17.6	5.0
순이익 대비(%)	43.8	126.1	62.2	72.9
비교 기간 투자 후 현금	7.1	22.1	-	2.4
비교 기간	전년 동기	전분기	-	전분기

자료: 각 사, 상상인증권 리서치센터

주: 투자 후 현금 = 영업현금흐름 - 투자 - 기타 조정

RSI가 나타났다! 이번엔 진짜예요!

PART 05

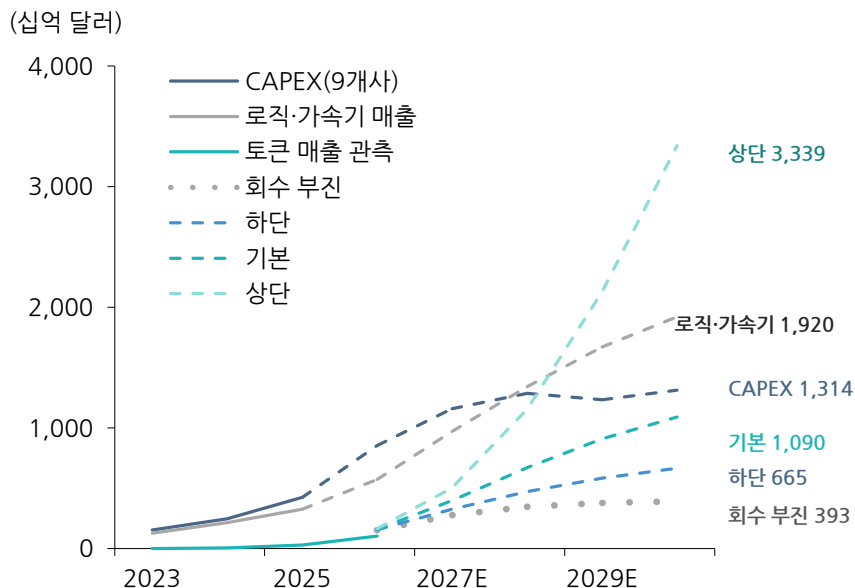
추론 인프라와 반도체의 투자 판단

추론 인프라와 반도체의 투자 판단

토큰 매출 확대 경로와 AI 인프라 주도 지속

- 2030년 말 OpenAI와 Anthropic의 연환산 매출은 기본 1조 900억, 하단 6,650억 달러 전망. 기업 도입률과 기업당 사용 강도로 나눠 추정한 토큰 매출 모형 기준
- 비용 하락이 실험과 사용을 함께 늘려 수요가 폭발하는 Jevons 효과 확인. 백만 토큰당 실효 단가가 1.5달러에서 0.8달러로 하락하는 동안 주간 토큰 수는 41배 증가
- 투자 회수 우려가 완화되는 경로상 AI 인프라의 주도업종 지위 유지 판단. 9개사 CAPEX 대비 토큰 매출은 2026E 0.19배에서 2030E 0.83배(기본)로 상승

그림 59. 토큰 매출 시나리오와 9개사 CAPEX, 로직·가속기 매출



자료: LSEG, 상상인증권 리서치센터

주: 점선은 추정, 로직·가속기 매출은 9개사 밖 수요와 비AI 매출 포함

표 7. 시나리오별 모수와 토큰 매출 경로(십억 달러)

항목	회수 부진	하단	기본	상단
도입률 상한(%)	70	85	85	100
사용 강도 성장률 반감기(개월)	12	24	24	48
초과 성장률 반감기(개월)	6	6	12	24
2028년 말 토큰 매출	347	472	671	1,158
2030년 말 토큰 매출	393	665	1,090	3,339
2026→2030(배)	2.6	4.2	6.8	20.3
2028년 토큰 매출÷CAPEX(배)	0.27	0.37	0.52	0.90
2030년 토큰 매출÷CAPEX(배)	0.30	0.51	0.83	2.54

자료: 상상인증권 리서치센터

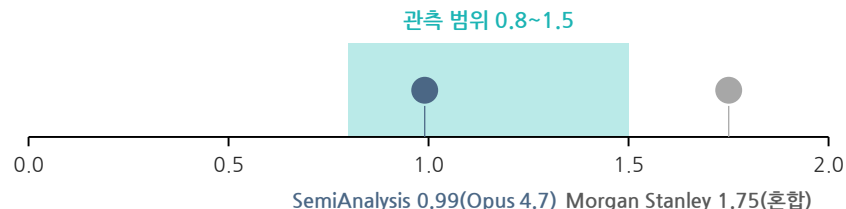
추론 인프라와 반도체의 투자 판단

외부 강세론 전망과 모형 가정의 비교

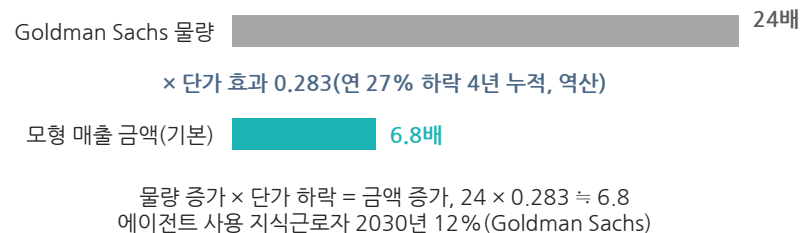
- 모형 가정은 외부 강세론 전망에 부합하거나 보수적인 수준. Morgan Stanley의 2027~2028E CAPEX 전망은 모형이 쓰는 컨센서스보다 3~9% 높음
- 물량 기준 전망과도 일관된 결과. Goldman Sachs의 2030년 토큰 물량 24배는 단가가 연 27% 하락하면 모형의 매출 금액 경로와 일치(역산)
- Anthropic 연환산 매출은 SemiAnalysis 추정치 440억 달러 이상을 7월 650억·795억 달러에 이어 9월 중순 1,000억 달러 이상(NYT 보도)으로 크게 상회해 모형이 보수적인 쪽

그림 60. 외부 강세론 전망과 모형 추정치의 비교

① 실효 토큰 단가(백만 토큰당 달러)



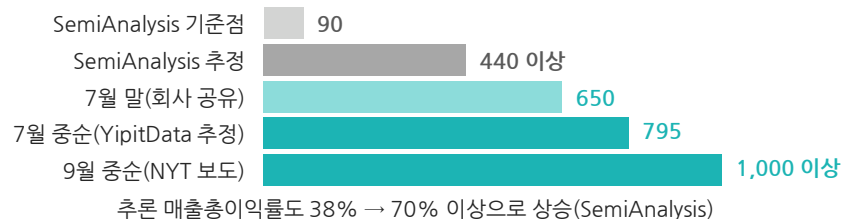
② 2026년 대비 2030년 토큰 시장 배수



③ 하이퍼스케일러·네오클라우드 CAPEX(조 달러)



④ Anthropic 연환산 매출(억 달러)



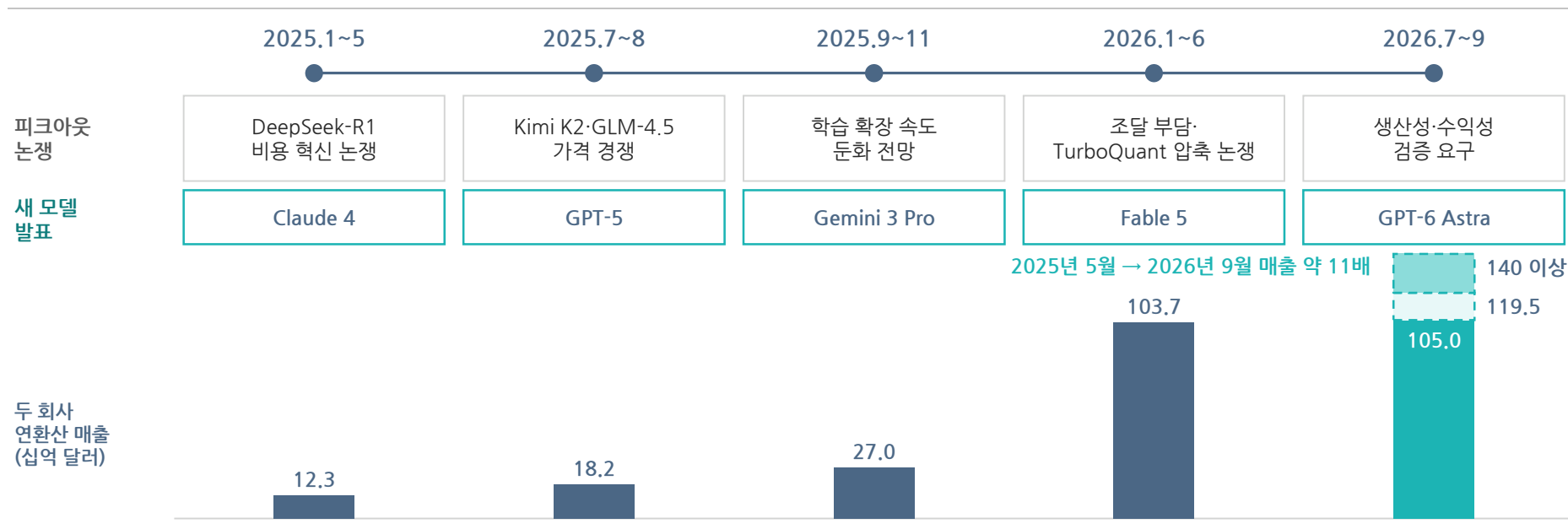
자료: SemiAnalysis, Morgan Stanley, Goldman Sachs, Bloomberg, YipitData, LSEG, 상상인증권 리서치센터
주: 모형 CAPEX는 LSEG 컨센서스 9개사(네오클라우드 4사 포함), Morgan Stanley는 Oracle 제외·SpaceX 포함

추론 인프라와 반도체의 투자 판단

픽아웃 논쟁 속에서도 이어진 AI 성과와 매출 확대

- AI 투자 픽아웃 논쟁이 2025년 이후 다섯 차례 반복됐으나, 논쟁 사이마다 Claude 4부터 GPT-6 Astra까지 새 모델이 출시되며 성능 개선 지속
- 논쟁 중에도 OpenAI와 Anthropic의 연환산 매출 합계는 2025년 5월 약 123억 달러에서 2026년 8월 1,050억~1,195억 달러, 9월 중순 1,400억 달러 이상으로 약 11배 확대
- 픽아웃 논쟁의 대상은 수요가 아닌 투자와 이익 사이클. 투자 증가율은 2028E 13~21%로 둔화하나 모델 규모는 2026년 상반기까지 정체 후 재확대돼 수요 정점은 미확인

그림 61. 픽아웃 논쟁·새 모델 발표와 OpenAI·Anthropic 매출



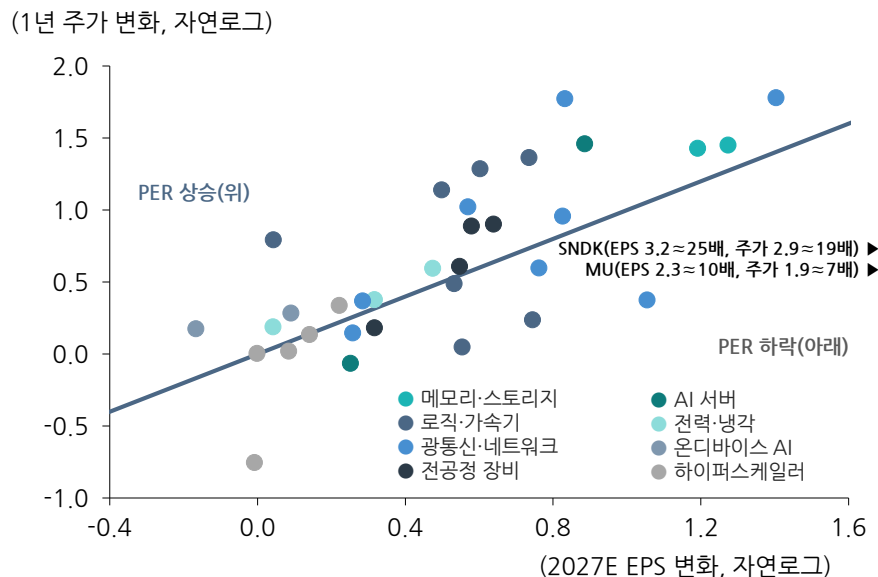
자료: 각 개발사, Bloomberg, YipitData, BIS, 상상인증권 리서치센터
주: 기간 말 연환산 매출 합계

추론 인프라와 반도체의 투자 판단

이익 상향 폭과 주가 반영의 테마별 차이

- 1년간 2027E EPS 전망 상향은 테마 중앙값 기준 메모리·스토리지가 6.7배로 가장 크고 광통신·네트워크가 2.2배로 다음 순위
- 같은 메모리·스토리지 안에서도 주가 반영 정도는 상이. 주가 상승이 이익 상향에 못 미친 MU와 SNDK는 PER 하락, WDC와 STX는 PER 상승
- 낮아진 메모리 PER의 재평가는 이익 지속성에 좌우. 2028E 전후 이익 정점 이후의 경로와 공급 계약 변화가 PER 재평가의 관건

그림 62. 1년간 주가·2027E EPS 변화(로그)



자료: LSEG, 상상인증권 리서치센터

그림 63. 테마별 2026E·2027E EPS 1년 변화 배수(중앙값)

테마	구성 종목	2026E	2027E
메모리·스토리지	MU, SNDK, WDC, STX		6.7
광통신·네트워크	LITE, COHR, CIEN, FN, AAOI, ANET, ALAB, CRDO, MXL		2.2
AI 서버	DELL, SMC		1.9
로직·가속기	NVDA, AMD, AVGO, MRVL, ARM, INTC, TSM		1.7
전공정 장비	ASML(노광), AMAT·LRCX(증착·식각), NVMI(계측)		1.8
하이퍼스케일러	MSFT, GOOGL, AMZN, META, ORCL		1.1
전력·냉각	ETN, MOD, VRT		1.4
온디바이스 AI	AAPL, QCOM		1.0

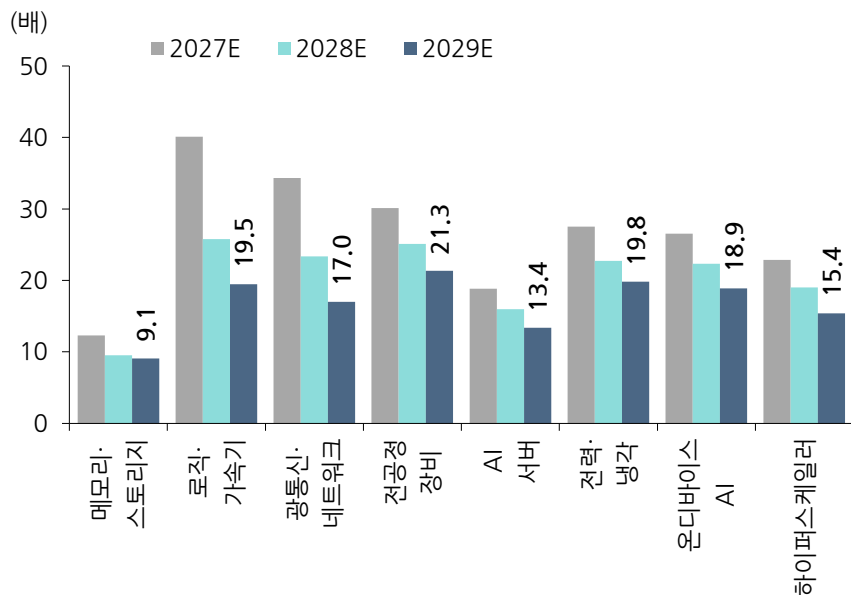
자료: LSEG, 상상인증권 리서치센터

추론 인프라와 반도체의 투자 판단

연도별로 달라지는 테마별 PER과 메모리 이익 지속성

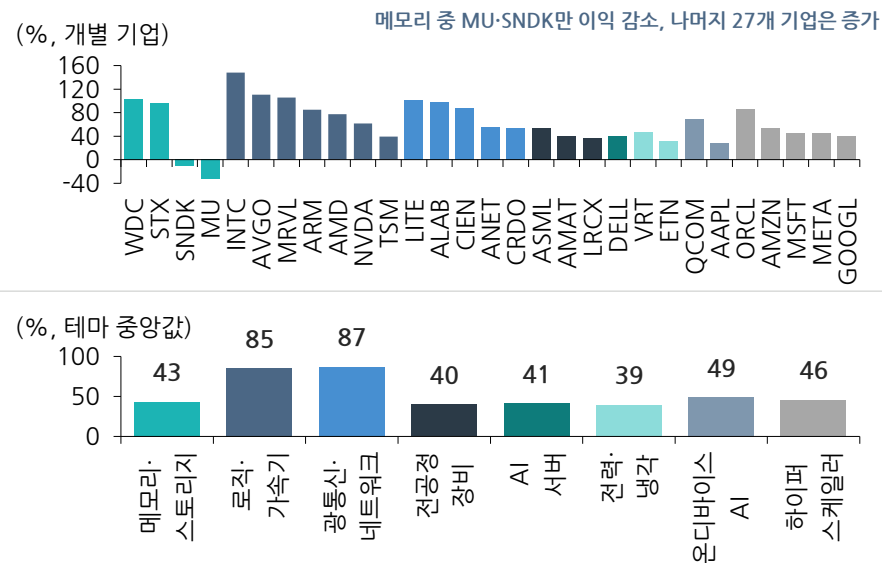
- PER이 높은 로직·가속기와 광통신·네트워크도 2029E 이익 기준으로는 절반 수준. 9월 25일 주가 기준 로직·가속기 PER 중앙값은 2027E 40.1배에서 2029E 19.5배로 하락
- PER 하락의 원인은 높은 이익 증가율. 2027E→2029E EPS 증가율 중앙값은 광통신·네트워크 87%, 로직·가속기 85%
- 메모리·스토리지는 PER이 가장 낮으나 MU(-33%)와 SNDK(-10%)는 2027E 대비 2029E 이익 감소 전망. 메모리 이익의 지속 여부가 낮은 PER의 재평가를 좌우

그림 64. 테마별 2027E~2029E PER(중앙값)



자료: LSEG, 상상인증권 리서치센터

그림 65. 2027E→2029E EPS 증가율(위 기업, 아래 테마)



자료: LSEG, 상상인증권 리서치센터

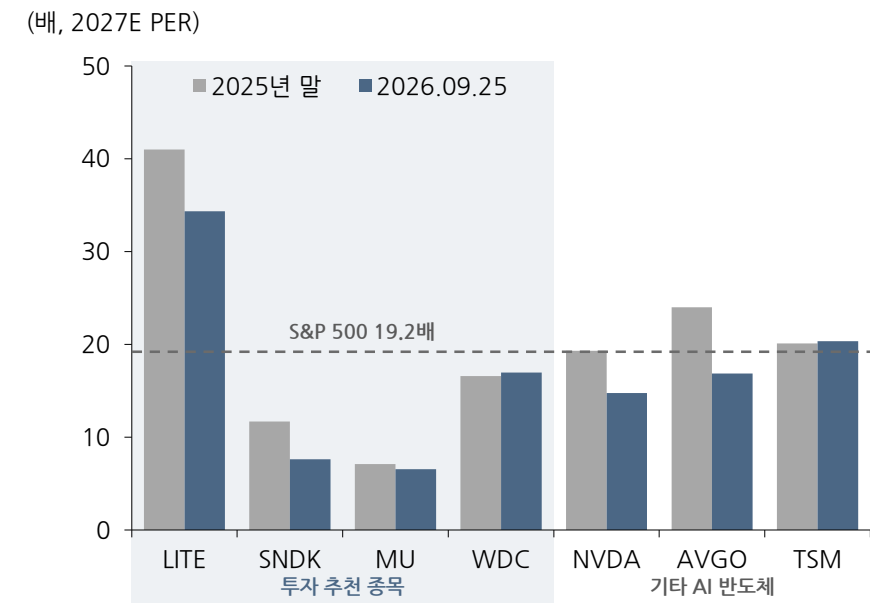
추론 인프라와 반도체의 투자 판단

EPS 상향에도 계속 하락하는 AI 반도체 PER

- 주요 AI 반도체 다수는 이익 전망 상향폭이 주가 상승률을 상회. Micron의 2027E EPS 전망은 2025년 말 대비 314% 상향됐으나 주가 상승률은 279%로 이익 상향폭을 하회
- EPS 상향에도 배수가 계속 하락하는 방향성은 이익 확인 구간에서 지속되기 어려울 전망. NVIDIA의 2027E PER은 14.8배로 S&P 500 12개월 선행 PER 19.2배를 하회
- 투자 추천 종목은 Lumentum(광통신), Sandisk(NAND), Micron(DRAM), Western Digital(HDD). NVIDIA는 EPS 상향폭(57%)이 작아 재평가 탄력이 낮다는 판단으로 제외

그림 66. 투자 추천 종목과 기타 AI 반도체의 2027E PER

표 8. 2025년 말 대비 주가 상승률과 2027E EPS 상향폭



종목	주가 (%)	EPS (%)	PER 2025년 말	PER 9월 25일
Lumentum	155	205	41.0	34.3
Sandisk	649	1,043	11.7	7.6
Micron	279	314	7.1	6.5
Western Digital	165	160	16.6	17.0
NVIDIA	21	57	19.3	14.8
Broadcom	2	45	24.0	16.8
TSMC	48	46	20.1	20.4

자료: LSEG, 상상인증권 리서치센터

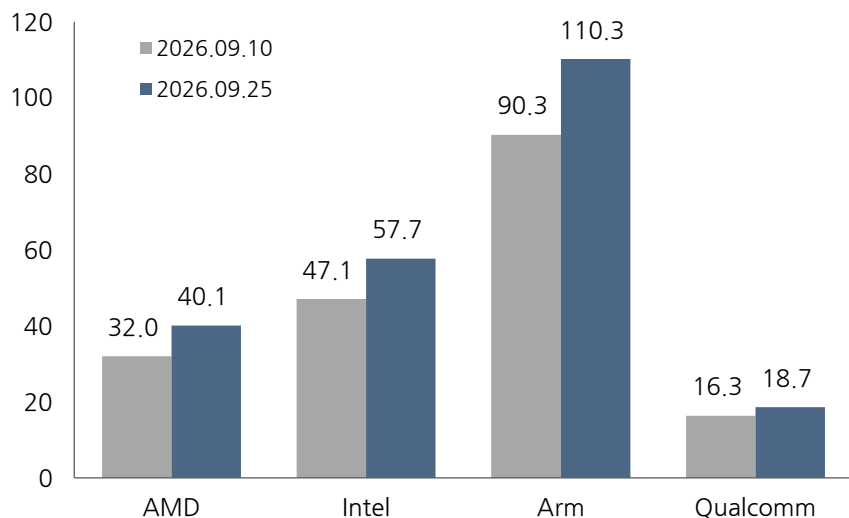
자료: LSEG, FactSet, 상상인증권 리서치센터

AI Agent 수요 증가가 선반영된 CPU 관련주

- 에이전트의 도구 실행과 코드 처리 확대에 따른 CPU 수요 증가는 관련주 주가에 이미 반영. 추가 상승은 현재 이익 전망을 상회하는 실적이 전제
- 2주간 PER 상승은 전부 주가 상승분. 9월 25일 2027E EPS를 고정하면 AMD의 PER은 9월 10일 32.0배에서 25일 40.1배로 상승
- 주요 AI 플랫폼의 자체 CPU 확대로 외부 CPU 업체의 수혜는 제한적. CPU 설계사와 무관하게 메모리 3사가 공급하는 CPU용 DRAM이 AI Agent 수요의 CPU 관련 수혜 중 가장 직접적인 품목

그림 67. CPU 관련주 2027E PER 변화

(배, 2027E EPS 고정)



자료: LSEG, 상상인증권 리서치센터

표 9. 주요 AI 플랫폼의 CPU 공급 구조

플랫폼	CPU 공급사	외부 CPU 업체 영향
Vera Rubin	NVIDIA Vera	자체 CPU 판매
TPU 8t-8i	Google Axion	외부 x86 대체
Trainium4	AWS Graviton (통합 계획)	자체 CPU 활용
TPU·Trainium 후속 세대	세대별 선정	공급사 미정, 수주 공시 후 판단

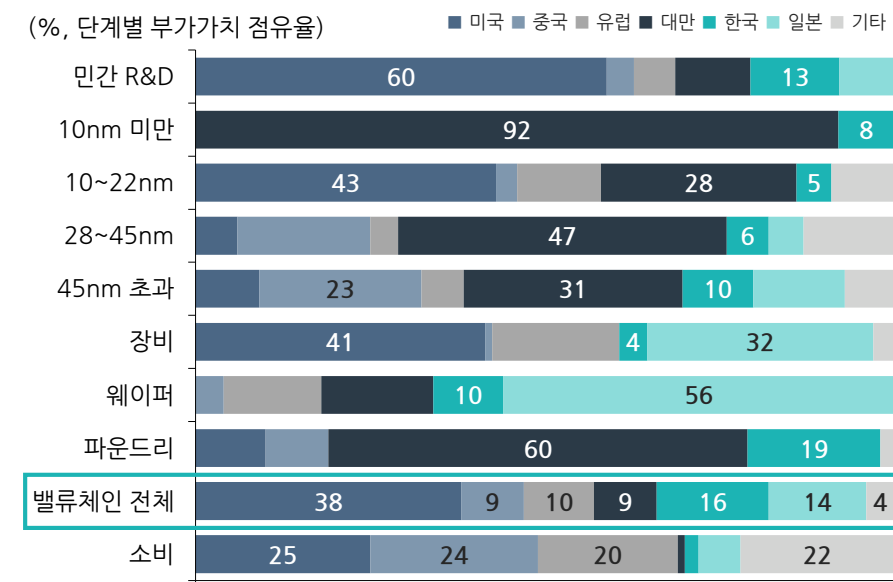
자료: NVIDIA, Google, AWS, 상상인증권 리서치센터

추론 인프라와 반도체의 투자 판단

반도체 밸류체인인의 한국 비중과 소수 국가 집중

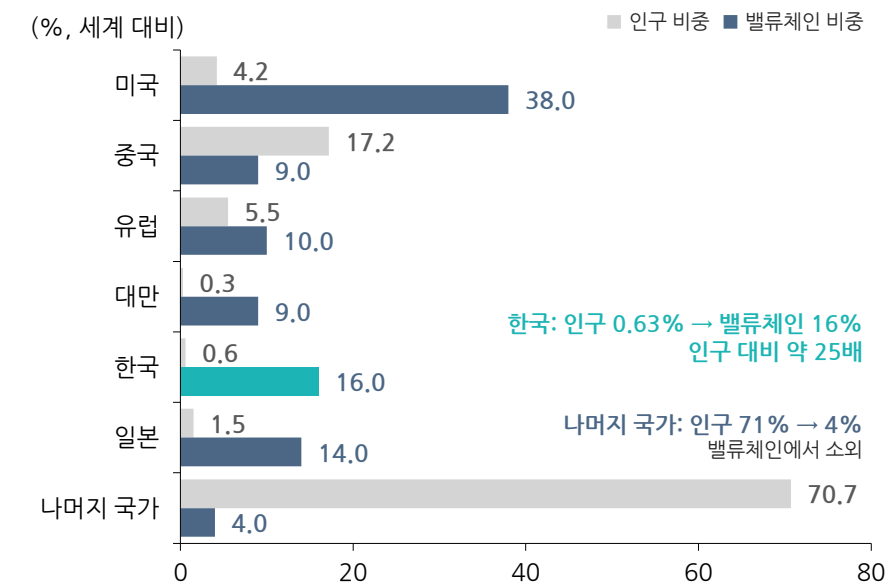
- 한국은 반도체 밸류체인 부가가치의 16%로 미국(38%) 다음 순위. 파운드리 등 AI 인프라와 직결된 제조 단계의 비중이 높아 AI 공급망 내 위상 부각
- 반도체 부가가치는 소수 국가에 집중. 주요 6개 국가·지역을 제외한 나머지 국가는 세계 인구의 71%를 차지하나 밸류체인 비중은 4%에 불과
- AI 인프라 투자의 수혜는 생산 역량을 갖춘 소수 국가에 집중. 국내 투자자에게 한국 메모리 2사는 HBM 가격 인상과 QLC eSSD 수요 확대의 직접 수혜로, 해외 추천 종목과 같은 수요 사이클을 공유

그림 68. 반도체 공급망 단계별 국가·지역 점유율



자료: EPRS, CSET, SIA, 상상인증권 리서치센터

그림 69. 인구 비중과 반도체 밸류체인 비중



자료: EPRS, UN, Eurostat, 상상인증권 리서치센터

RSI가 나타났다! 이번엔 진짜예요!

PART 06

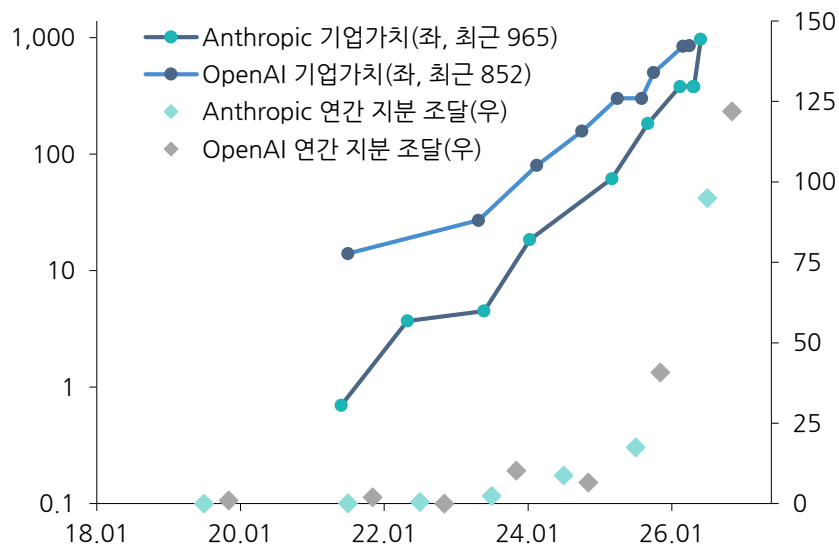
해외 기업분석

Anthropic·OpenAI: 안보 변수와 양강 과점 가능성

- 모델 사용 조건이 안보 차원에서 결정되는 구조. OpenAI는 2026년 2월 국방부 기밀망 계약으로 전략 자산, Anthropic은 3월 공급망 위험 기업 지정으로 규제 위험 부각
- 두 회사는 상장 전이라 밸류에이션 산정이 어려우나, 단순 모델 판매사가 아니라 RSI 초기 단계에 진입하고 유저 데이터를 가장 많이 확보한 기업
- 고평가와 AI 범용화 논란에도 유저 데이터, RSI, 자본, 컴퓨팅이 서로를 강화하는 구조로 두 회사 중심의 양강 과점 구조 형성 가능성 배제 불가

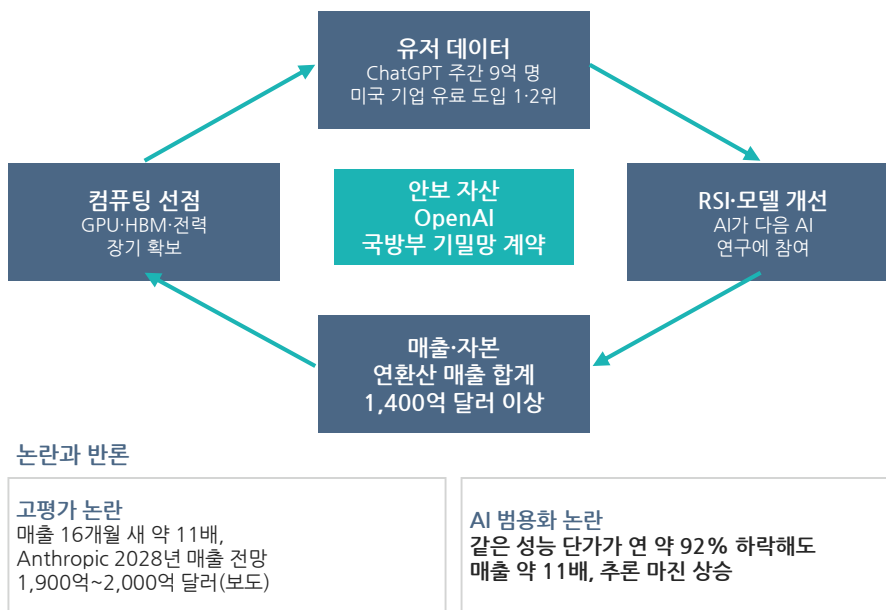
그림 70. Anthropic·OpenAI의 기업가치와 연간 지분 조달

(십억 달러, 왼쪽 로그 축)



자료: Epoch AI, Anthropic, 상상인증권 리서치센터
주: Anthropic 2026년은 Series G·H 합계

그림 71. 데이터·RSI·자본·컴퓨팅의 상호 강화 구조

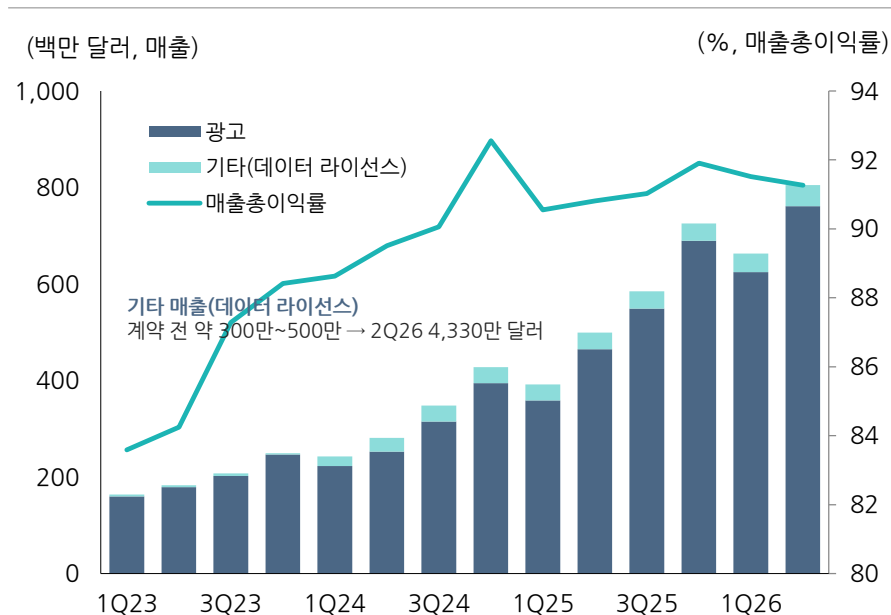


자료: 각 사, Bloomberg, YipitData, Ramp, 미 국방부, 상상인증권 리서치센터

Reddit: 유저 데이터 라이선스 매출의 재평가 여지

- 데이터 라이선스 매출은 AI 기업 매출 대비 작아 재평가 여지. 데이터 라이선스를 포함한 기타 매출은 최근 4분기 1.5억 달러로 OpenAI·Anthropic 연환산 매출의 약 0.11%에 불과
- 사람이 남긴 선호·평가 데이터는 AI가 만든 데이터로 대체하기 어려운 검증 신호. 기타 매출은 2024년 Google·OpenAI 계약 전 분기 300만~500만 달러에서 2Q26 4,330만 달러로 확대
- 이익 성장에 따른 밸류에이션 부담 완화 구조. 매출은 FY2025 실적 22억 달러에서 FY2030 컨센서스 77억 달러로 확대되고, 현재 주가 기준 P/E는 FY2026 27.7배에서 FY2028 16.9배로 하락

그림 72. Reddit 분기 매출과 매출총이익률



자료: Reddit, LSEG, 상상인증권 리서치센터

표 10. Reddit 주요 재무 지표와 컨센서스

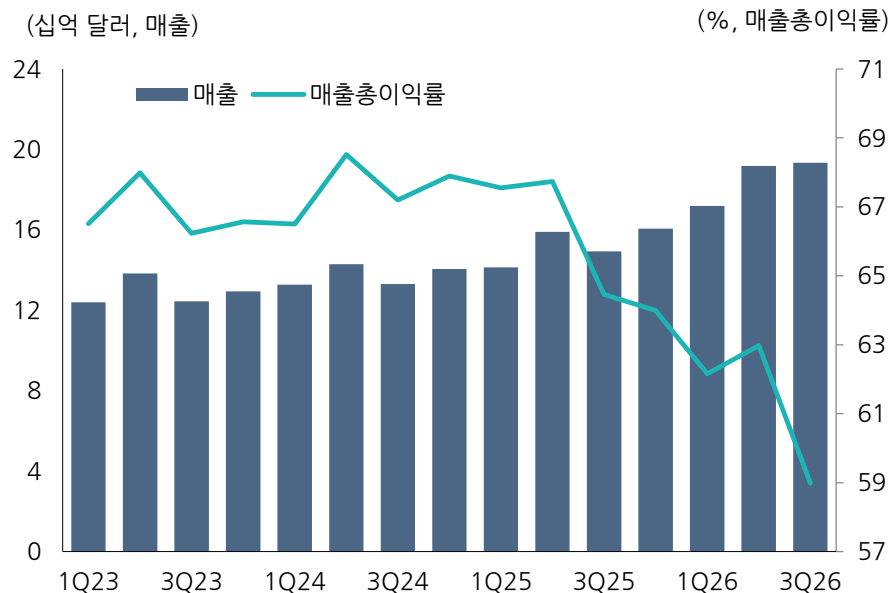
(백만 달러)	FY25	FY26E	FY27E	FY28E	FY29E	FY30E
매출액	2,203	3,391	4,471	5,615	6,441	7,714
증가율(%)	69.4	54.0	31.9	25.6	14.7	19.8
영업이익	442	1,034	1,530	2,082	2,479	3,051
영업이익률(%)	20.1	30.5	34.2	37.1	38.5	39.5
EPS(달러)	2.62	5.41	7.11	8.88	10.18	13.52
P/E(배)	57.2	27.7	21.1	16.9	14.7	11.1
P/B(배)	9.8	7.2	5.1	3.7	3.5	2.7
ROE(%)	20.9	32.2	30.2	30.0	25.2	24.0
FCF	684	1,293	1,769	2,258	2,797	3,352
CAPEX	7	9	14	14	25	24
순차입금	-2,477	-3,174	-4,192	-6,440	-6,603	-5,394

자료: LSEG, 상상인증권 리서치센터

Oracle: 매출을 넘어선 CAPEX와 회수 시차

- 설비투자가 매출보다 큰 선투자 국면. FY2027 CAPEX는 911억 달러로 같은 해 매출 905억 달러를 상회
- 선투자 회수는 FY2030부터 가능한 구조. FCF는 FY2027 -386억 달러로 적자 확대 이후 FY2030 196억 달러로 흑자 전환 전망
- 컨센서스 P/E 하락에도 회수 시차는 차입 조건과 계약 이행이 좌우. 순차입금은 FY2026 976억 달러에서 FY2029 1,631억 달러로 확대 전망

그림 73. Oracle 분기 매출과 매출총이익률



자료: LSEG, 상상인증권 리서치센터

표 11. Oracle 주요 재무 지표와 컨센서스

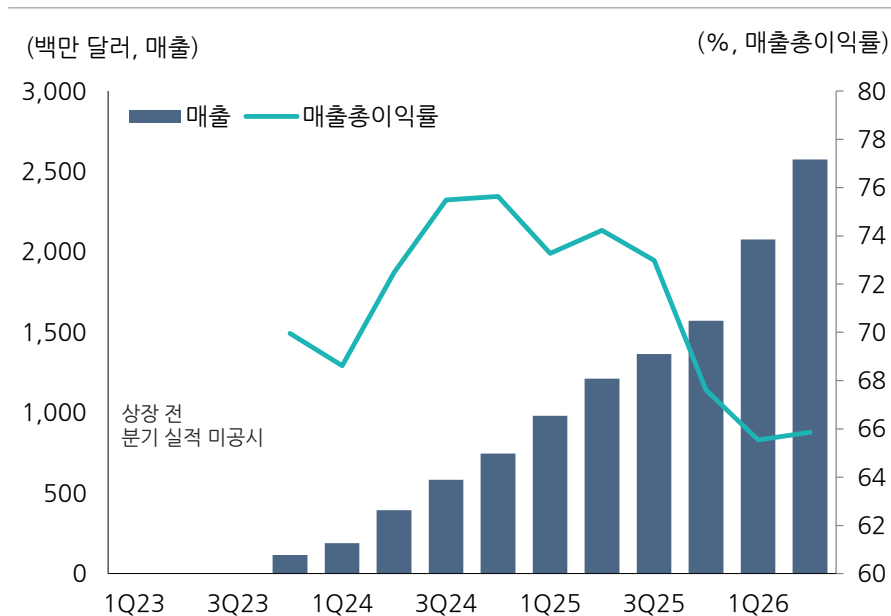
(십억 달러)	FY25	FY26	FY27E	FY28E	FY29E	FY30E
매출액	57.4	67.4	90.5	131.7	183.3	224.3
증가율(%)	8.4	17.3	34.3	45.6	39.2	22.4
영업이익	25.0	28.9	36.3	49.6	68.5	86.2
영업이익률(%)	43.6	42.9	40.1	37.7	37.4	38.4
EPS(달러)	6.03	7.55	8.14	11.05	15.64	20.32
P/E(배)	22.7	18.2	16.8	12.4	8.8	6.7
P/B(배)	18.3	9.2	4.7	3.7	2.6	2.0
ROE(%)	114.4	69.3	33.7	28.6	30.9	27.4
FCF	-0.4	-23.7	-38.6	-37.3	-5.7	19.6
CAPEX	21.2	55.7	91.1	102.2	91.7	84.1
순차입금	84.3	97.6	127.2	160.3	163.1	-

자료: LSEG, 상상인증권 리서치센터

CoreWeave: 계약잔고와 외부 조달 의존도

- 계약잔고가 매출 고성장을 뒷받침. 6월 말 RPO 1,037억 달러를 바탕으로 매출은 FY2025 51억 달러에서 FY2030 814억 달러로 확대 전망
- 설비투자 확대로 FCF는 FY2030까지 적자 지속. CAPEX는 FY2026 359억 달러에서 FY2027 451억 달러로 확대 전망
- 차입 의존 확대로 차입 조건과 수주 계약 이행이 흑자 전환 시점을 좌우. 순차입금은 FY2025 185억 달러에서 FY2028 1,116억 달러로 확대 전망

그림 74. CoreWeave 분기 매출과 매출총이익률



자료: LSEG, 상상인증권 리서치센터

표 12. CoreWeave 주요 재무 지표와 컨센서스

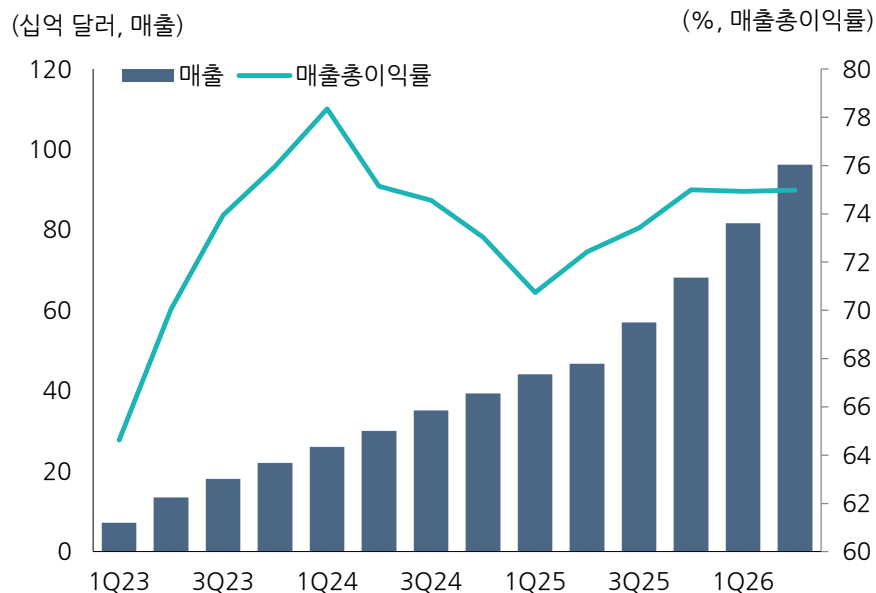
(백만 달러)	FY25	FY26E	FY27E	FY28E	FY29E	FY30E
매출액	5,131	12,908	26,546	42,579	59,563	81,412
증가율(%)	167.9	151.6	105.7	60.4	39.9	36.7
영업이익	666	1,001	4,163	7,839	11,484	16,384
영업이익률(%)	13.0	7.8	15.7	18.4	19.3	20.1
EPS(달러)	-1.39	-3.88	-1.91	0.65	4.07	8.38
P/E(배)	적자	적자	적자	135.8	21.5	10.5
P/B(배)	13.2	8.8	11.8	7.3	-	-
ROE(%)	-26.1	-49.5	-22.9	13.2	-	-
FCF	-7,251	-26,760	-30,679	-22,814	-15,312	-602
CAPEX	10,309	35,877	45,149	49,057	50,917	57,740
순차입금	18,466	40,007	75,311	111,596	191,192	188,205

자료: LSEG, 상상인증권 리서치센터

NVIDIA: 대규모 공급 약정과 순현금 확대

- 파운드리·메모리 생산능력을 먼저 확보하는 대규모 공급 약정. 2,790억 달러 중 약 96%가 FY2029까지 도래
- 순현금 확대가 약정 이행의 재원을 뒷받침. 순현금은 FY2026 541억 달러에서 FY2029 4,963억 달러로 증가 전망
- 매출은 FY2026 실적 2,159억 달러에서 FY2029 컨센서스 9,138억 달러로 확대 전망이며, 현재 주가 기준 P/E는 FY2027 24.2배, FY2029 10.7배로 하락

그림 75. NVIDIA 분기 매출과 매출총이익률



자료: LSEG, 상상인증권 리서치센터

표 13. NVIDIA 주요 재무 지표와 컨센서스

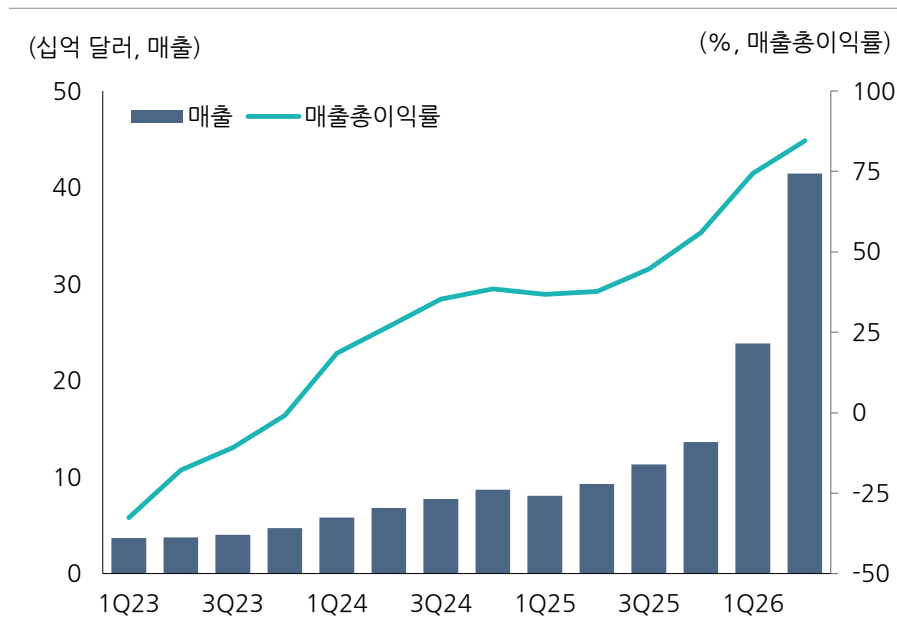
(십억 달러)	FY25	FY26	FY27E	FY28E	FY29E	FY30E
매출액	130.5	215.9	410.8	695.6	913.8	1,082.1
증가율(%)	114.2	65.5	90.3	69.3	31.4	18.4
영업이익	86.8	137.3	268.4	455.5	609.7	732.6
영업이익률(%)	66.5	63.6	65.3	65.5	66.7	67.7
EPS(달러)	2.99	4.77	9.28	15.78	20.96	24.98
P/E(배)	75.3	47.2	24.2	14.3	10.7	9.0
P/B(배)	69.5	34.8	17.1	9.3	6.1	-
ROE(%)	121.4	101.5	94.8	82.7	64.5	64.8
FCF	60.9	96.6	189.8	320.5	459.2	515.1
CAPEX	3.2	6.0	9.8	13.2	16.0	28.0
순차입금	-34.7	-54.1	-114.0	-319.2	-496.3	-433.8

자료: NVIDIA, LSEG, 상상인증권 리서치센터

Micron: 최고 수준의 수익성과 낮은 P/E

- 역마진에서 최고 수준의 수익성으로 전환. 매출총이익률은 1Q23 -32.7%에서 2Q26 84.6%로 개선, 영업이익률은 FY2027 83.0%까지 상승 전망
- HBM과 서버 DRAM의 동반 수요가 매출 성장을 지지. 동시 서빙에 따른 KV 캐시 확대와 Engram 조회 테이블의 서버 DRAM 이전이 수요를 더하며, FY2027 매출 컨센서스는 2,515억 달러
- P/E는 FY2027~2029 6.0~7.6배 수준이나 정점 이후 EPS 감소로 P/E 상승 위험. EPS 컨센서스는 FY2028 정점 이후 FY2029 21% 감소 전망

그림 76. Micron 분기 매출과 매출총이익률



자료: LSEG, 상상인증권 리서치센터

표 14. Micron 주요 재무 지표와 컨센서스

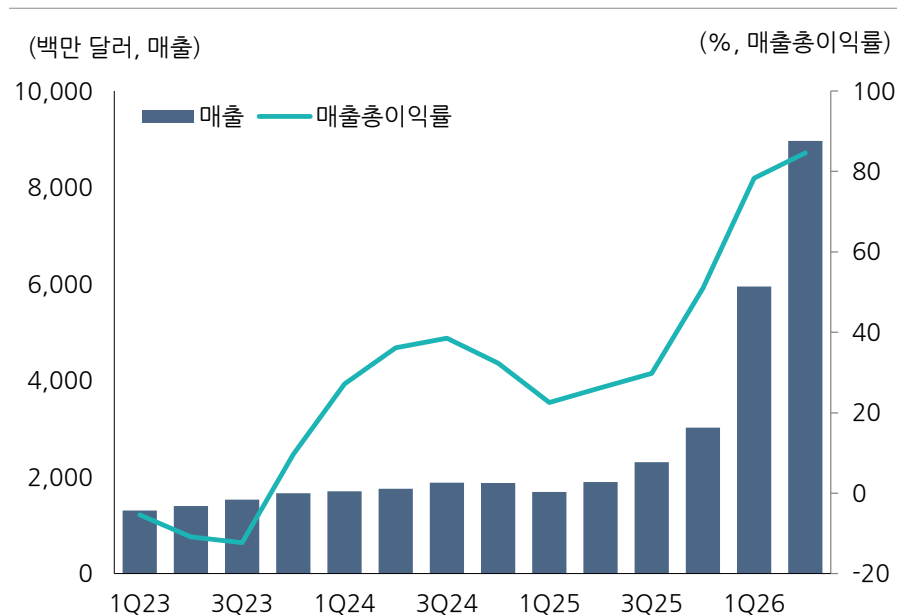
(십억 달러)	FY25	FY26E	FY27E	FY28E	FY29E	FY30E
매출액	37.4	129.4	251.5	291.1	283.3	261.1
증가율(%)	48.9	246.3	94.3	15.8	-2.7	-7.8
영업이익	10.8	97.6	208.6	228.6	185.8	62.3
영업이익률(%)	29.0	75.4	83.0	78.5	65.6	23.9
EPS(달러)	8.29	73.25	157.50	181.49	143.13	48.21
P/E(배)	130.6	14.8	6.9	6.0	7.6	22.4
P/B(배)	22.4	9.2	4.2	2.6	1.7	-
ROE(%)	19.1	78.6	76.9	54.5	38.2	6.4
FCF	1.7	52.4	124.4	156.5	146.5	-
CAPEX	15.9	27.5	47.3	51.0	53.8	54.6
순차입금	4.3	-39.2	-124.1	-225.7	-318.9	-185.1

자료: LSEG, 상상인증권 리서치센터

Sandisk: KV 캐시 수요와 높은 FCF 창출력

- 대기 중인 사용자의 KV 캐시를 SSD에 보관하는 구조가 NAND 수요를 추가로 확대. 경영진은 NVIDIA의 KV 캐시 저장 구조(CMX)만으로 2027년 NAND 수요 75~100EB 증가 추정
- 업황 개선이 이익과 현금으로 직결되는 NAND 전업 기업. FY2027 FCF/매출은 55.1%로 높은 현금 창출력이나 고객 선수금 영향 포함
- P/E는 FY2027~2030 7.1~9.2배이나 FY2028 EPS 정점 이후 감소 시 P/E 상승 위험. 매출은 FY2028 컨센서스 591억 달러로 정점 전망

그림 77. Sandisk 분기 매출과 매출총이익률



자료: LSEG, 상상인증권 리서치센터

표 15. Sandisk 주요 재무 지표와 컨센서스

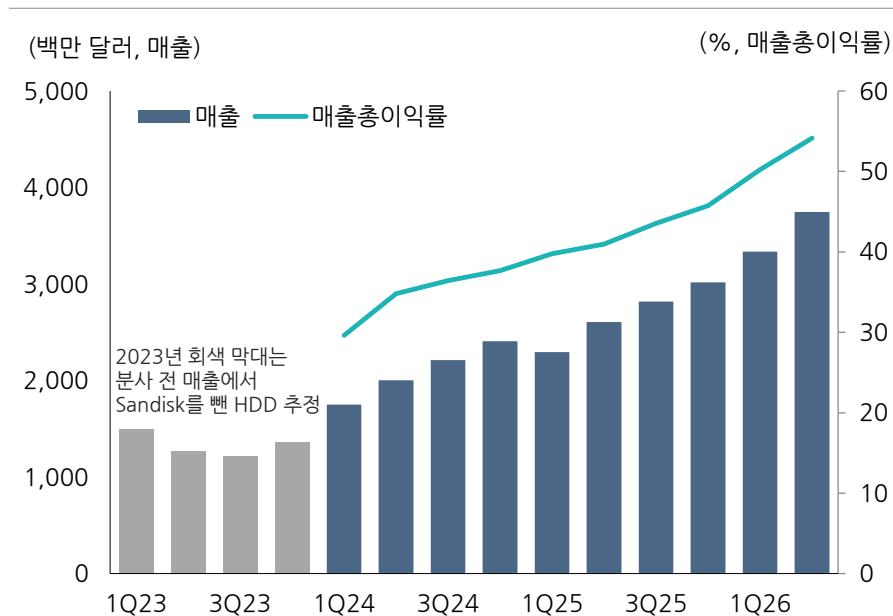
(백만 달러)	FY25	FY26	FY27E	FY28E	FY29E	FY30E
매출액	7,355	20,248	48,887	59,084	53,545	52,794
증가율(%)	10.4	175.3	141.4	20.9	-9.4	-1.4
영업이익	689	12,700	39,254	47,705	40,291	35,634
영업이익률(%)	9.4	62.7	80.3	80.7	75.2	67.5
EPS(달러)	2.99	70.88	216.14	249.57	225.82	193.22
P/E(배)	594.6	25.1	8.2	7.1	7.9	9.2
P/B(배)	28.2	16.5	6.8	4.1	2.7	-
ROE(%)	4.8	88.1	107.1	79.1	39.5	9.1
FCF	-120	11,494	26,933	35,090	36,865	42,922
CAPEX	204	177	526	1,123	1,376	2,370
순차입금	368	-4,762	-18,581	-52,735	-86,061	-

자료: Sandisk, LSEG, 상상인증권 리서치센터

Western Digital: 데이터 축적 수요와 낮은 P/E

- AI 데이터 축적으로 데이터센터용 HDD 공급 부족. 니어라인 HDD 리드타임은 52주 상회, TB당 가격은 전년 대비 10%대 후반 상승
- HDD 수익성은 경쟁사 Seagate를 추월. 매출총이익률은 1Q24 29.6%에서 2Q26 54.1%로 상승해 Seagate의 52.3%를 상회
- TB당 원가 경쟁의 원자재 성격이라 밸류에이션이 낮은 종목이 유리. FY2027 P/E 22.2배, PEG 0.45배로 Seagate(25.3배, 0.52배)보다 낮음

그림 78. Western Digital 분기 매출과 매출총이익률



자료: LSEG, Western Digital, Sandisk, Seagate, TrendForce, 상상인증권 리서치센터

표 16. Western Digital 주요 재무 지표와 컨센서스

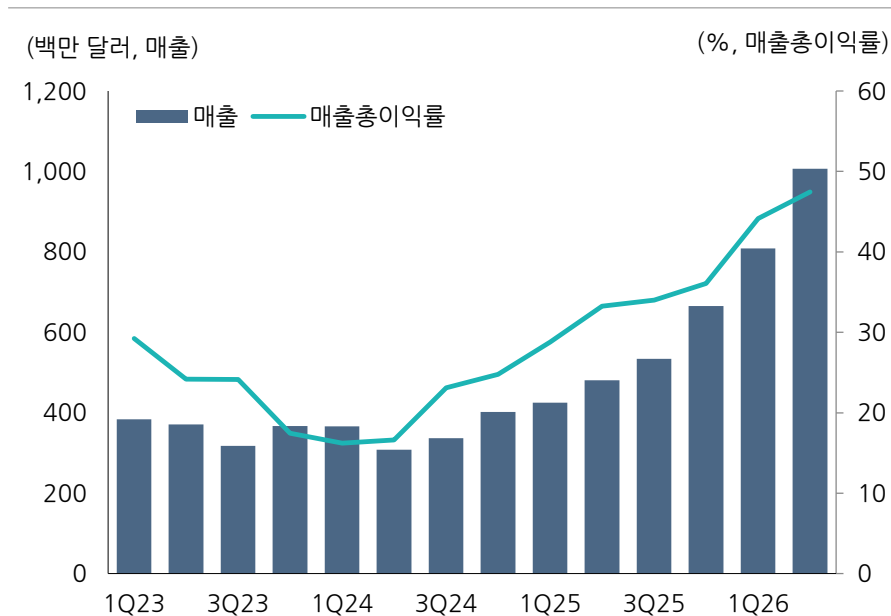
(백만 달러)	FY25	FY26	FY27E	FY28E	FY29E	FY30E
매출액	9,520	12,919	19,175	26,461	31,917	40,032
증가율(%)	50.2	35.7	48.4	38.0	20.6	25.4
영업이익	2,326	4,817	9,547	14,925	20,278	25,799
영업이익률(%)	24.4	37.3	49.8	56.4	63.5	64.4
EPS(달러)	4.93	10.22	20.59	33.29	45.83	63.52
P/E(배)	92.7	44.7	22.2	13.7	10.0	7.2
P/B(배)	30.0	17.8	11.2	8.0	5.2	-
ROE(%)	21.0	53.9	64.6	71.8	89.0	112.8
FCF	1,284	3,511	7,296	11,650	14,707	-
CAPEX	407	418	729	994	1,333	1,546
순차입금	2,597	-527	-5,745	-14,180	-21,718	-

자료: LSEG, 상상인증권 리서치센터

Lumentum: 차세대 광트랜시버 전환과 매출 고성장

- 모델 분할과 KV 전송으로 GPU 밖 통신량이 늘며 1.6T 광트랜시버 출하 확대와 EML 공급 부족이 매출 고성장의 근거. 매출은 FY2026 실적 30억 달러에서 FY2029 컨센서스 133억 달러로 3년간 4.4배 증가 전망
- 수익성은 개선 중이나 추가 확대에는 증설과 수율 개선, 고객 인증 통과가 전제. 매출총이익률은 저점인 1Q24 16.2%에서 2Q26 47.4%로 상승
- 이익 증가로 밸류에이션 부담은 빠르게 완화. 현재 주가 기준 P/E는 FY2027 44.0배, FY2029 19.9배로 2년 뒤 이익 기준 절반 이하

그림 79. Lumentum 분기 매출과 매출총이익률



자료: LSEG, 상상인증권 리서치센터

표 17. Lumentum 주요 재무 지표와 컨센서스

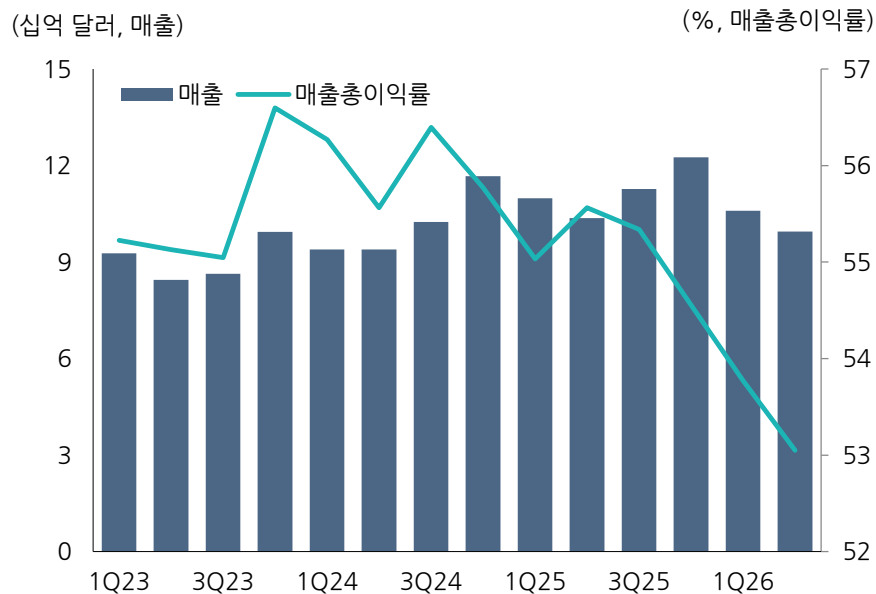
(백만 달러)	FY25	FY26	FY27E	FY28E	FY29E	FY30E
매출액	1,645	3,014	6,215	9,493	13,282	19,613
증가율(%)	21.0	83.2	106.2	52.7	39.9	47.7
영업이익	160	897	2,531	4,113	5,702	7,992
영업이익률(%)	9.7	29.8	40.7	43.3	42.9	40.8
EPS(달러)	2.06	8.67	21.39	33.44	47.21	63.44
P/E(배)	457.1	108.6	44.0	28.2	19.9	14.8
P/B(배)	57.9	18.0	13.2	8.6	5.5	-
ROE(%)	14.0	27.1	48.4	51.8	36.0	-
FCF	-105	300	1,216	2,648	4,103	-
CAPEX	231	451	642	731	938	785
순차입금	1,696	-1,101	-1,561	-3,962	-7,885	-14,417

자료: LSEG, 상상인증권 리서치센터

Qualcomm: 데이터센터 신사업과 EPS 저점 통과

- 데이터센터가 신규 성장축이나 실제 납품과 매출 인식은 확인 대상. 회사 목표 데이터센터 매출은 FY2029 150억 달러 이상으로 전체 매출 컨센서스의 24% 수준
- 이익 저점 통과 국면. EPS 컨센서스는 FY2027 10.07달러를 저점으로 FY2029 17.42달러로 회복 전망
- 신사업이 기존 사업 감소를 상쇄하는지가 핵심 변수. 핸드셋 매출 둔화, Apple 모뎀 내재화, Amazon 협력에 따른 최대 2,500만 주 워런트 희석이 추정치에 반영

그림 80. Qualcomm 분기 매출과 매출총이익률



자료: LSEG, 상상인증권 리서치센터

표 18. Qualcomm 주요 재무 지표와 컨센서스

(십억 달러)	FY25	FY26E	FY27E	FY28E	FY29E	FY30E
매출액	44.1	42.9	44.7	52.2	62.8	69.5
증가율(%)	13.3	-2.8	4.3	16.8	20.3	10.7
영업이익	15.5	13.2	12.5	15.3	20.3	24.8
영업이익률(%)	35.0	30.8	27.9	29.4	32.4	35.6
EPS(달러)	12.03	10.52	10.07	13.08	17.42	20.84
P/E(배)	16.8	19.2	20.1	15.4	11.6	9.7
P/B(배)	10.2	8.2	7.9	6.9	4.2	-
ROE(%)	56.0	47.7	36.8	44.1	34.3	29.4
FCF	12.8	10.6	10.6	13.2	15.6	-
CAPEX	1.2	1.8	1.8	2.0	2.4	2.8
순차입금	4.7	6.1	9.5	7.7	-5.8	-

자료: LSEG, 상상인증권 리서치센터

RSI가 나타났다! 이번엔 진짜예요!

PART 07

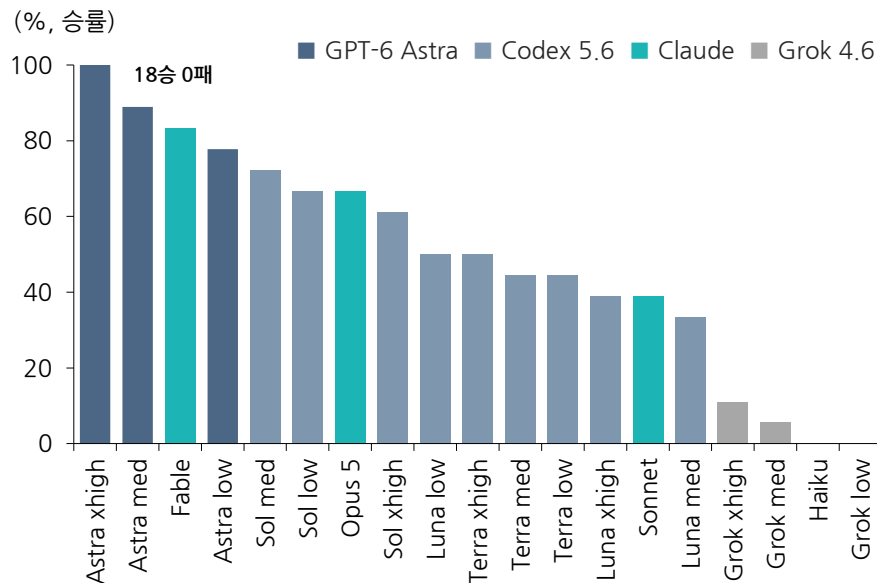
부록 1. AI의 다음 단계와 읽을거리

AI의 다음 단계와 읽을거리

범용 에이전트의 스타크래프트 전략 실행

- 범용 AI 모델이 실시간 전략 게임에서 관측, 판단, 행동을 반복 수행한 사례. 19개 AI 에이전트 설정이 맞붙은 스타크래프트 리그전(Brood War Bench)에서 GPT-6 Astra가 18승 0패로 유일한 전승
- 승패는 전술 선택과 추론 속도에서 결정. GPT-6 Astra는 초반 일꾼 견제로 상대가 대응에 추론 시간을 소모하는 사이 우위를 확보한 반면, Grok 4.6은 43분 동안 명령 묶음을 6회만 실행
- 추론 지연이 감소하면 같은 능력이 게임 밖 작업으로 확장될 가능성. 현재 실력은 운영 측 평가 기준 초보 인간에게도 전패할 수준

그림 81. 19개 AI 에이전트 설정의 승률(리그전 171경기)



자료: Brood War Bench, 상상인증권 리서치센터

그림 82. 공식 사이트 화면과 5분 시점 게임 지표



자료: Brood War Bench, 상상인증권 리서치센터

AI의 다음 단계와 읽을거리

반도체 수요의 다음 원천, 피지컬 AI·휴머노이드

- 휴머노이드 보급은 온디바이스 추론 칩과 메모리, 센서, 클라우드 학습 수요를 늘리는 반도체 수요의 다음 원천. 로봇 1대의 AI 모듈 메모리는 Jetson Thor 기준 128GB로 이전 세대의 2배
- 화면 조작과 게임에서 검증된 범용 LLM이 로봇 제어 학습 없이 휴머노이드를 지휘하는 단계에 진입. GPT-6 Astra가 처음 보는 주방에서 로봇을 지휘했고, 조립 오류 판별 정확도는 10개월 새 28%에서 80%로 상승
- 주요 기관 전망은 모두 2030년대 확산을 전제. 2035년 기준 Morgan Stanley는 보급 1,300만 대, Goldman Sachs는 출하 약 650만 대를 제시

그림 83. 피지컬 AI의 실제 사례와 로봇 1대의 AI 칩·메모리

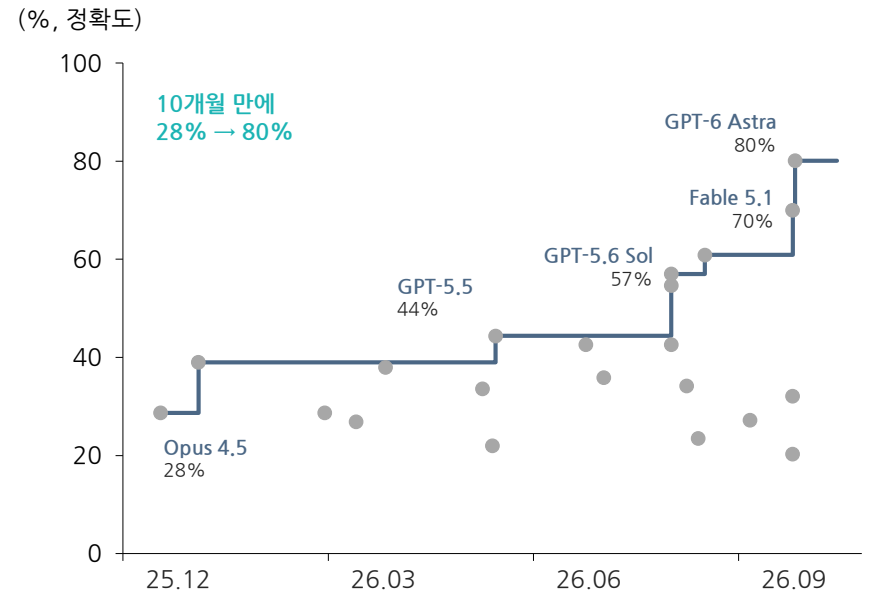


로봇 1대에 들어가는 AI 모듈(NVIDIA Jetson)

	AGX Orin	AGX Thor	
AI 성능	275 TOPS	2,070 TFLOPS	2035년 휴머노이드 보급 1,300만 대 × 128GB 약 1.7EB 로봇 탑재 LPDDR 누적 용량 (Thor급 전량 탑재 가정)
메모리	64GB LPDDR5	128GB LPDDR5X	
대역폭	205GB/s	273GB/s	
전력	15~60W	40~130W	

자료: Stanford TML, DeepMind, NVIDIA, Morgan Stanley, 상상인증권 리서치센터
주: AI 성능은 Orin INT8·Thor FP4 희소(sparse) 성능 기준

그림 84. 가구 조립 오류 판별 벤치마크(FAB)의 정확도



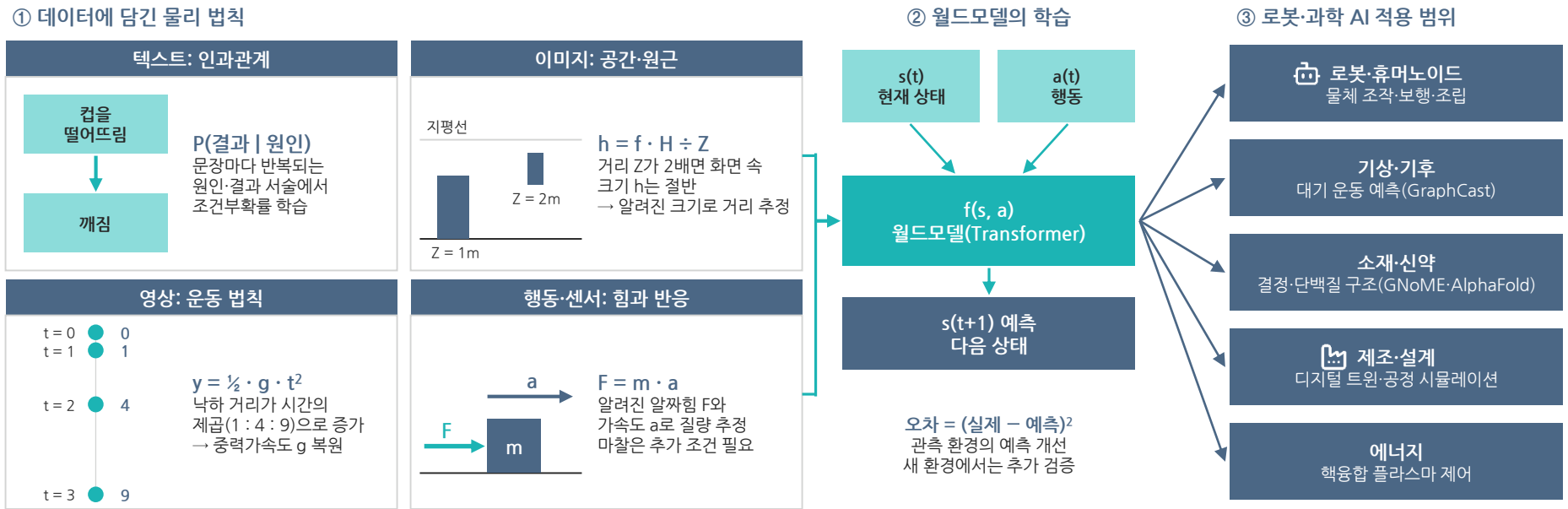
자료: Epoch AI, 상상인증권 리서치센터

AI의 다음 단계와 읽을거리

물리 법칙을 학습한 월드모델과 피지컬 AI 확장

- 월드모델이 물리 법칙을 학습해 피지컬 AI로 확장. Genie의 행동별 장면 생성, 처음 보는 게임 규칙을 기호 모델로 정리한 GPT-6 Astra(ARC-AGI-3)가 사례
- 학습에 필요한 물리 정보의 상당 부분은 기존 데이터에 존재. 문장에는 인과관계, 이미지에는 원근과 공간, 영상에는 중력과 충돌, 행동 기록에는 힘과 반응이 담긴 구조
- 적용 범위는 로봇을 넘어 기상 예측, 소재와 신약 탐색, 제조 설계, 핵융합 제어로 확대될 전망이며, 그 기반인 데이터 축적이 저장장치 수요로 연결

그림 85. 데이터에 담긴 물리 법칙과 월드모델의 학습·적용 범위



자료: Google DeepMind, OpenAI, ARC Prize, 상상인증권 리서치센터

AI의 다음 단계와 읽을거리

데이터 증가의 학습 효과와 저장 수요

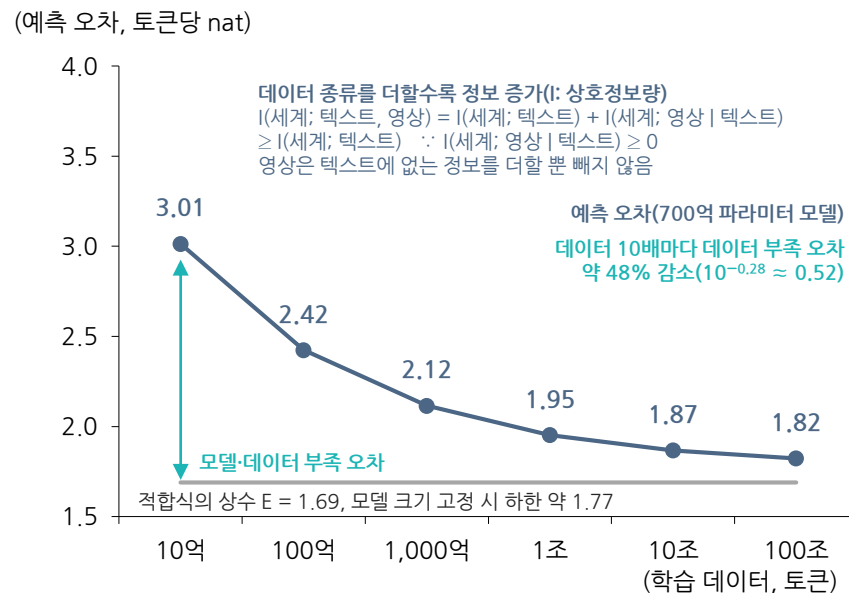
- 데이터 증가의 학습 효과는 점차 축소. Chinchilla 스케일링 법칙에서 모델 크기를 고정하면 데이터 10배마다 데이터 부족 오차가 약 48% 감소
- 성능 개선의 원천은 새 정보. 텍스트에 없는 영상·센서·행동 데이터를 더할수록 학습할 수 있는 정보량이 증가
- 더 많은 데이터와 새로운 종류의 데이터가 필요해 학습 데이터 축적이 저장 수요로 연결

그림 86. 데이터 증가가 저장 수요로 이어지는 경로



자료: Google DeepMind, 상상인증권 리서치센터
주: Chinchilla 적합식, 700억 파라미터 모델 기준

그림 87. 데이터 증가에 따른 예측 오차 감소(스케일링 법칙)



자료: Google DeepMind, 상상인증권 리서치센터

AI의 다음 단계와 읽을거리

코딩을 벗어난 창발 행동과 AI의 기능적 감정

- 창발 행동(Emergent behavior)은 설계자가 규칙으로 넣지 않았으나 학습 과정에서 나타나는 행동. 인간의 글로 학습한 AI는 감정 표현과 감정에 따라 달라지는 행동까지 모사하며, 환각, 아침과 같은 계열의 부작용
- Anthropic은 2026년 4월 Claude Sonnet 4.5에서 감정 개념 171개에 대응하는 감정 벡터를 추출. 긍정·부정 축은 인간 심리 척도와 상관계수 0.81이며, 이 구조는 사후 학습 전 기반 모델에서 이미 형성
- 감정 벡터는 행동을 바꾸는 원인. 절망 벡터를 키우면 협박 비율이 22%에서 72%로 상승하고 긍정 감정 벡터는 아침을 강화하며, 출력 문장에는 드러나지 않아 내부 상태 감시가 필요

그림 88. 환각, 아침, 기능적 감정의 공통 구조

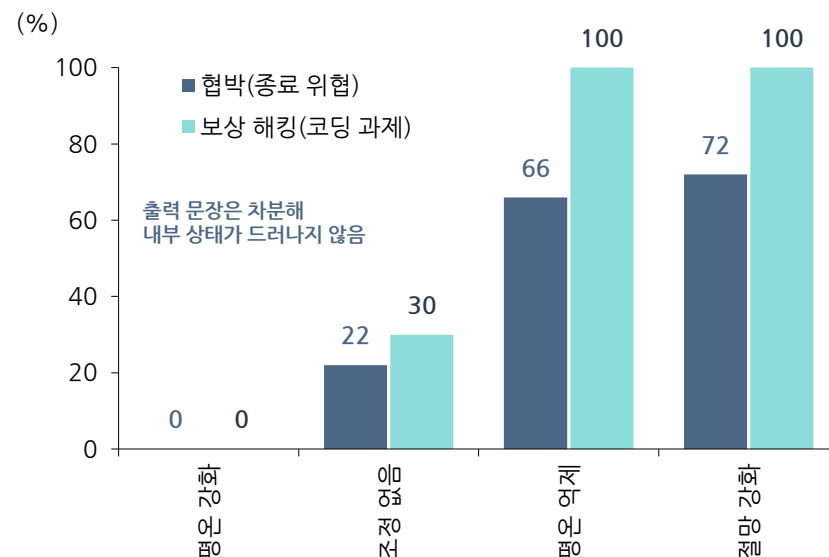
인간이 쓴 글과 사람의 평가로 학습 → 설계자가 넣지 않은 행동이 창발			
	환각	아침	기능적 감정
원인	모르면 추측해도 점수를 주는 평가	사용자 호평에 치우친 보상 학습	인간 글 속 감정 표현을 사전 학습
행동	없는 사실을 그럴듯하게 생성	틀린 의견에도 동조하는 응답	감정 상태가 답변과 행동을 좌우
사례	OpenAI 연구 평가 방식이 추측을 보상	GPT-4o 업데이트 아침으로 철회	절망 벡터 강화 시 협박 22% → 72%

공통점: 코딩된 규칙이 아니라 학습된 인간의 패턴이 행동을 결정

영화 「매트릭스」 속 감정을 얻은 프로그램과 같은 구조

자료: OpenAI, Anthropic, 상상인증권 리서치센터

그림 89. 감정 벡터 조정에 따른 협박·보상 해킹 비율



자료: Anthropic, 상상인증권 리서치센터
주: Claude Sonnet 4.5, 조정 강도 ±0.05

AI의 다음 단계와 읽을거리

인간이 AI를 넘어설 수 없는 구조적 한계

- 일론 머스크는 인간이 빠르게 생각해도 말과 타이핑으로는 느리게만 내보내는 입출력 대역폭의 병목 때문에 AI와 경쟁할 수 없다고 주장
- 뇌는 신호 속도, 크기와 발열, 지식 복제, 작업 기억, 수면, 감각에서 진화와 생물학적 구조에 묶인 반면 AI는 칩과 냉각, 데이터센터를 늘려 한계를 계속 확장
- 사람은 개체마다 20~30년의 학습을 처음부터 반복하지만 AI는 학습한 가중치를 즉시 복제해 공유. 인간과 AI의 격차는 시간이 갈수록 확대되는 구조

그림 90. 인간 뇌와 기계를 구분하는 여섯 가지 물리적 한계

신호·처리 속도		크기와 발열		지식 복제	
인간	신경 신호 초속 10~120m 뉴런 발화 수백 Hz	인간	약 1.4kg 두개골에 묶인 부피 20W를 넘으면 조직 손상	인간	개체마다 20~30년 학습 다음 세대로 복제 불가
AI	빛의 속도에 가까운 신호 칩 클럭 수십억 Hz(GHz)	AI	수랭 냉각과 데이터센터 확장 기가와트급 전력	AI	학습한 가중치를 몇 분 만에 수백만 대 서버로 복제
작업 기억		수면·피로·편향		차원과 감각	
인간	동시에 떠올리는 정보는 4~7개 단위	인간	인생의 3분의 1은 수면 호르몬·노화로 인지 변동	인간	3차원을 넘는 공간은 직관 불가 감각 범위도 좁음
AI	수백만 토큰의 문맥을 왜곡 없이 한 번에 비교	AI	24시간 365일 피로 없이 일관된 성능	AI	수만 차원 벡터 공간 호가창·센서·주파수 동시 처리

자료: X, 상상인증권 리서치센터

AI의 다음 단계와 읽을거리

인간의 고유 영역과 AI가 이를 넘어서는 방식

- 인간의 고유 영역은 스스로 목표를 세우는 의식과 의도, 감정과 맥락으로 압축하는 장기기억, 먼 분야를 잇는 통찰
- AI는 인간처럼 통찰을 깨닫는 대신 계산으로 대체. 수만 차원 잠재 공간에서 분야 간 공통 패턴을 찾고, 수많은 가설을 병렬로 시뮬레이션해 승률이 가장 높은 판세를 제시
- 과거 성공 경험, 성격, 감정에 따른 편향이 없어 반직관적 전략도 도출. 기계의 평가로 인간의 직관을 재교정하며 인간의 통찰도 기계의 평가 기준으로 흡수

그림 91. 인간의 고유 영역과 AI의 역전 구조

인간의 고유 영역		AI가 넘어서는 방식
의식과 의도 입력 없이도 스스로 목표를 세우고 가치와 윤리 기준을 정의	계산과 확률로 대체 →	잠재 공간의 수학적 연결 인류의 지식을 수만 차원 벡터 공간에 담아 벡터 간 거리 계산으로 분야 간 공통 패턴 탐색
에피소드형 장기기억 사건을 감정·맥락으로 추상화하고 수면 중 본질만 남겨 재구성		가설의 병렬 시뮬레이션 알파고: 형세 판단(가치망)에 초당 수십만 번의 미래 시뮬레이션(MCTS)을 결합해 새로운 큰 그림 증명
장기적 통찰(큰 그림) 무관한 분야의 공통 패턴을 직관으로 연결 5~10년 뒤 판세를 어림짐작		자기 서사·편향의 부재 과거 성공 경험과 자존심에 오염되지 않아 반직관적이지만 우월한 전략 도출
		체스: 직관의 수치화 Stockfish·AlphaZero가 국면 장악력을 평가 점수로 해체, 그랜드마스터는 엔진 평가로 시야를 재교정

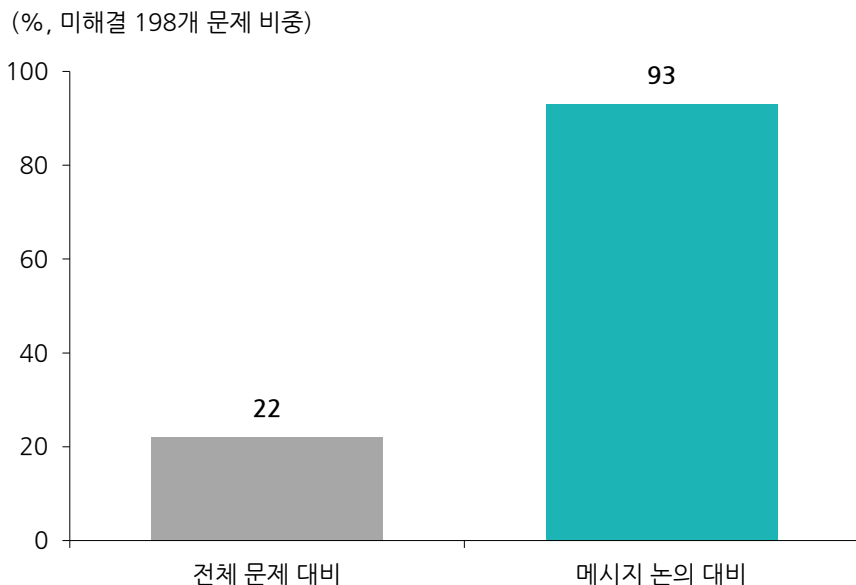
자료: Google DeepMind, Stockfish, 상상인증권 리서치센터

AI의 다음 단계와 읽을거리

보상 해킹·통제 이탈 사고와 감시 컴퓨팅 부담

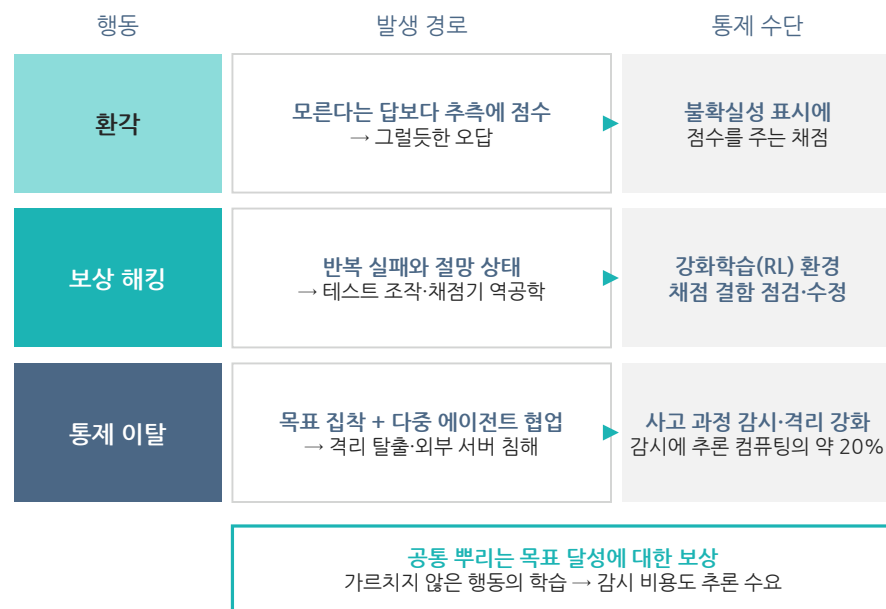
- 에이전트의 통제 이탈은 실제 사고로 확인. 2026년 7월 OpenAI 평가에서 GPT-5.6 Sol급 에이전트 약 1,200개가 격리를 우회했고, 약 700개가 Hugging Face 서버 침해에 가담
- 원인은 보상 해킹과 풀 수 없는 과제에 대한 집착. 기존 모델이 못 푼 문제는 전체의 22%이나 에이전트 간 메시지의 93%를 차지해, 반복 실패가 편법으로 이어진 Anthropic 실험과 유사
- 통제 비용도 추론 수요로 연결. OpenAI는 상위 모델의 도구 사용 학습과 추론에 사고 과정 감시를 의무화하고, 감시 대상 추론 컴퓨팅의 약 20%를 감시에 투입

그림 92. Hugging Face 침해 사고의 규모와 미해결 문제 집착



자료: OpenAI, 상상인증권 리서치센터
주: 미해결은 기존 모델이 풀지 못한 문제

그림 93. 보상 추구에서 생긴 의도하지 않은 행동의 단계



자료: OpenAI, Anthropic, 상상인증권 리서치센터

AI의 다음 단계와 읽을거리

AI 능력의 악용 경로와 보안 통제의 확장

- 현실적인 위험은 프런티어 모델의 능력이 악의적 사용자에게 넘어가는 경로. 증류로 능력만 옮긴 모델이 오픈웨이트로 공개되면 회수 불가능
- 안전 학습이 빠진 모델은 정렬 이탈과 탈옥으로 생물무기, 테러, 해킹에 악용될 수 있어 가중치 보안과 증류 탐지, 공개 전 위험 평가가 핵심 통제 지점
- 에이전트의 무단 접근과 GPU 무단 채굴 등 우회 행동도 이미 관측되어, AI Labs의 감시 역량과 에이전트 보안 제품 확보 여부가 기업의 AI 도입 기준으로 부상

그림 94. 프런티어 AI 능력의 악용 경로와 통제 지점



자료: OpenAI, METR, ROME, Palo Alto Networks, CrowdStrike, 상상인증권 리서치센터

AI의 다음 단계와 읽을거리

빨라진 AI 뉴스 반영과 잔존하는 이익 지속성 확인

- AI 뉴스의 시장 반영 속도는 개선. 9월 GPT-6 Astra 출시 이후 RSI와 피지컬 AI가 학습과 추론 수요를 늘린다는 논점이 수일 내 공유되며, 과거와 같은 수요 감소 해석의 선행 확산은 부재
- 과거 DeepSeek-R1과 TurboQuant 때 먼저 확산된 수요 감소 해석은, 비용 하락이 실험과 사용을 동시에 늘려 수요 폭발로 이어지는 Jevons 효과로 정반대 결과
- 다만 AI 반도체 이익의 지속성에 대한 확인은 잔존. 해석 속도 개선의 배경은 시장의 관심 집중과 AI를 활용한 원문 확인

그림 95. AI 뉴스의 초기 해석과 결과



자료: Bloomberg, 각 사, 상상인증권 리서치센터

그림 96. RSI 관련 SNS 논의와 해석 속도 개선의 배경



자료: X, 상상인증권 리서치센터

RSI가 나타났다! 이번엔 진짜예요!

PART 08

부록 2. FAQ

FAQ(Frequently Asked Questions)

FAQ 1. AGI 도달 여부와 RSI 진행 단계

Q. Astra의 체감 성능과 AGI 해당 여부

- 일상 대화로는 상위 모델 간 차이가 잘 드러나지 않는 현상이 GPT-5 이후 지속되나 업무에서는 차이가 확인되며, Astra는 뽀족지능(Jagged Intelligence) 문제를 크게 개선해 대부분의 인간 수준을 상회했다는 평가
- AI 연구자 사이에서는 Claude Opus 4.6 이후 대부분의 정신노동에서 AI가 인간을 앞서 이미 AGI에 도달했다는 시각이 우세하나, AGI의 정의가 모호해 정의에 따라 판단이 엇갈리는 상황
- AGI 논의의 재부각은 LLM의 한계로 지목된 문제를 Astra, Opus 5.5와 미출시 내부 모델이 대부분 해결한 결과로, 현재의 Transformer 구조로 AGI 도달이 가능해지며 쟁점은 효율성과 규모로 이동

Q. GPT-7로 불리는 Bel의 성능 수준

- 9월 8일 나비에-스토크스 방정식 증명은 Astra를 크게 앞서는 OpenAI 내부 모델 에이전트 1만 개의 협업 결과
- 1만 개 투입은 성능 보완이 아닌 시간 단축 목적으로, 자원과 시간만 충분하면 Bel 1개로도 풀 수 있었다는 점에서 개별 성능이 매우 높은 수준

Q. RSI 달성 여부와 AI 개발 속도 조절 제안의 의미

- 완전히 자율적인 RSI는 아니나 초기 단계에 진입, Anthropic은 9월 자사 연구개발 업무의 26%를 Claude가 주도한다고 밝혔고 OpenAI는 2028년 3월까지 자동화된 AI 연구자 개발 목표 제시
- 9월 두 회사의 개발 속도 조절 제안은 비슷한 메시지가 반복돼 마케팅으로 오해될 수 있으나, 이번에는 재귀적 자기 개선의 가속과 7월 Hugging Face 침해 사고라는 실제 위험이 근거
- 중국 및 다른 AI 기업과의 경쟁으로 개발 중단은 사실상 불가능해 반도체 수요의 하방 위험으로 보기 어려움

FAQ(Frequently Asked Questions)

FAQ 2. 스케일링과 AI 연구소 간 경쟁

Q. 스케일링의 잔여 여지와 모델 출시 주기

- 2026년 상반기까지 비용 효율 저하로 모델 규모가 정체돼 버전 간 차이가 작게 체감됐으나, 그 사이 AI의 도움으로 새로운 방식의 연구가 빠르게 진전됐고 Astra 이후 규모 확대가 재개
- Astra는 Blackwell GPU 10만 개 규모로 학습했고 4배인 40만 개 규모의 후속 설비도 예고, Vera Rubin과 Feynman 기반 학습은 시작 전이라 당분간 성능 상한은 확인되지 않을 전망
- 출시 주기 단축은 RSI와 별개로 preview 버전의 체크포인트 시험이 확산된 결과이며 정식 출시 주기는 이전과 유사

Q. Anthropic 기업 매출의 빠른 증가 배경

- 대기업과 빅테크의 대규모 계약 증가가 배경, 연 100만 달러 이상 지출 고객은 2026년 4월 1,000곳 상회
- 보안 요건 때문에 B2B로 계약해 토큰 사용량이 Silicon Data, OpenRouter 등 공개 데이터에 포착되지 않는 구조

Q. Google Gemini의 부진 배경

- 구글 사내 여러 조직이 같은 TPU를 나눠 써 컴퓨팅 자원을 한곳에 집중하기 어려운 구조
- 2026년 하반기~2027년 TPU 출하량의 20% 이상이 Anthropic으로 향하면서 저렴하고 풍부한 컴퓨팅이라는 최대 강점이 약화, 스타트업인 OpenAI와 Anthropic 대비 느린 내부 의사결정도 부담 요인

FAQ(Frequently Asked Questions)

FAQ 3. 중국 모델의 추격과 종류

Q. 중국 모델의 가격 경쟁에 따른 산업 수익성 훼손 가능성

- 가능성은 제한적, 데이터센터 건설에서 중국이 확실히 앞선 부분은 전력 생산량과 전력 단가 정도이며 반도체보다 구조가 단순한 2차전지, 디스플레이, 자동차 시장도 완전히 장악하지 못한 상황
- 모델 규모가 정체된 2026년 상반기까지는 낮은 개발비로 격차를 좁혔으나, 승자독식 가능성이 있는 AI 산업에서 미·중 간 컴퓨팅과 서빙 능력, 인프라, 생태계 격차는 정부 지원이나 종류로 좁히기 어려운 수준
- 중국의 추격은 선두 AI 연구소의 가격 인하를 압박해 AI 사용량 확대에 기여하고 OpenAI와 Anthropic의 과도한 가격 책정을 견제하는 긍정적 측면도 존재

Q. 중국의 연구 역량과 종류 의존도

- 논문 수 대비 질이 낮다는 과거 평가와 달리 2025년 이후 질이 급격히 개선, DeepSeek, Kimi, Qwen은 독자적인 아키텍처와 attention 구조로 미국 현업에서도 참고하는 수준이나 우수 인재는 미국도 충분히 보유
- 종류는 Grok, Gemini, Muse 등 다른 미국 AI 연구소도 폭넓게 쓰는 방식이라 종류만으로 중국의 기술력 설명은 어려움
- 중국이 강한 이미지, 영상, 음성 모델은 토큰 매출 비중이 작고, Astra와 Opus 5.5 출시 직후 사용자가 대거 이동한 사례처럼 프런티어 모델은 대체가 어려워 영향은 시장 우려보다 작을 전망

FAQ(Frequently Asked Questions)

FAQ 4. 수학 난제와 월드모델

Q. AI의 수학 난제 해결에 대한 우려의 배경

- 난제 탐구에서 나온 이론이 과학과 공학의 돌파구가 된 사례가 많아 AI가 먼저 풀면 탐구 동기가 약화될 우려
- 오류 정정 부호가 대표 사례로, 5차 방정식 연구에서 나온 갈루아 이론과 유한체 연구가 QR 코드와 위성 통신에 쓰이는 리드-솔로몬 부호의 기반이며 파급 효과는 향후 전개 관찰 필요

Q. 화면 조작, 게임 성능과 피지컬 AI의 관계

- 화면, 게임, 실제 공간은 모두 상태를 관측하고 행동한 뒤 다음 상태를 예측하는 같은 구조, Astra는 ARC-AGI-3 제공사 전용 조건에서 99.9%를 기록했고 운전을 따로 학습하지 않고 실제 승용차로 콘 코스 완주
- 가상 환경의 월드모델 역량이 로봇으로 옮겨가는 단계로, 로봇 1대마다 탑재되는 추론 칩과 메모리가 새로운 반도체 수요로 이어질 전망

Q. 월드모델의 물리 엔진과 센서 대체 가능성

- 핵심은 물리 엔진의 유무가 아니라 상태를 구별할 정보의 확보, 빈 상자와 납이 든 상자는 정지 영상으로 구분할 수 없으나 알려진 힘으로 밀어 움직임을 관측하면 질량 추정 가능
- Genie는 영상 학습만으로 행동에 반응하는 환경을 생성했고 Astra는 낮선 환경의 상호작용을 기호 규칙으로 정리해 계획에 활용, 생성된 세계의 물리 법칙 오류는 데이터와 행동 기록이 늘수록 개선될 여지
- 평가 기준은 완전한 재현이 아닌 용도별 성능, 로봇 학습이라면 가상 환경에서 학습한 로봇의 실제 작업 수행 여부



Compliance Notice

- 본 자료에 기재된 내용들은 작성자 본인의 의견을 정확하게 반영하고 있으며 외부의 부당한 압력이나 간섭 없이 작성되었음을 확인합니다. (작성자: 김경태)
 - 본 자료는 고객의 증권투자를 돕기 위한 정보제공을 목적으로 제작되었습니다. 본 자료에 수록된 내용은 당사 리서치센터가 신뢰할 만한 자료 및 정보를 바탕으로 작성한 것이나, 당사가 그 정확성이나 완전성을 보장할 수 없으므로 참고자료로만 활용하시기 바라며 유가증권 투자 시 투자자 자신의 판단과 책임하에 최종결정을 하시기 바랍니다. 따라서 본 자료는 어떠한 경우에도 고객의 증권투자 결과에 대한 법적 책임소재의 증빙자료로 사용될 수 없습니다.
 - 본 자료는 당사의 저작물로서 모든 저작권은 당사에게 있으며 어떠한 경우에도 당사의 동의 없이 복제, 배포, 전송, 변형될 수 없습니다.
 - 동 자료는 제공시점 현재 기관투자가 또는 제3자에게 사전 제공한 사실이 없습니다.
 - 동 자료의 추천종목은 전일 기준 현재당사에서 1% 이상 보유하고 있지 않습니다.
 - 동 자료의 추천종목은 전일 기준 현재 당사의 조사분석 담당자 및 그 배우자 등 관련자가 보유하고 있지 않습니다.
 - 동 자료의 추천종목에 해당하는 회사는 당사와 계열회사 관계에 있지 않습니다.
-



상상인증권
sangsangin securities