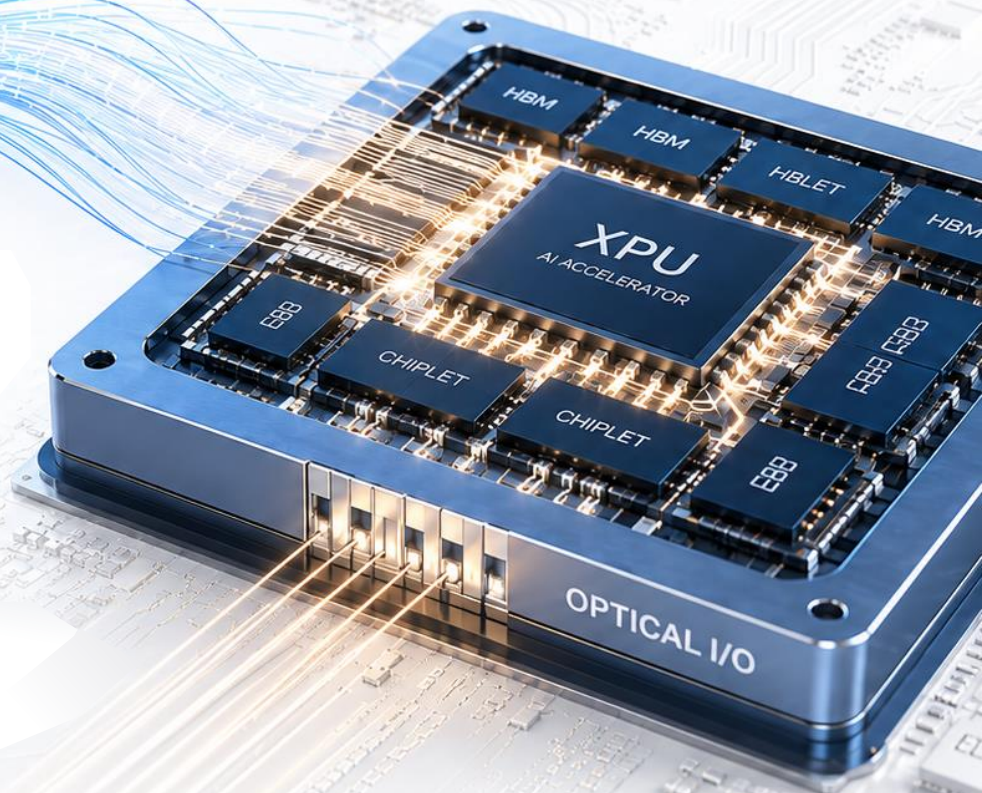


AI 인터커넥트 대전환

병목은 광으로 모이고,
이제 곧 빛의 시대가 열린다

박장욱 Jangwook.Park@daishin.com



Contents

I. 스케일업으로 확장될 광 시장	05
II. 광전환의 이유 – 메모리 병목	17
III. 광전화의 이유 – Beachfront	29
IV. 표준은 미정. 광 전환은 필연	39
V. 투자전략	47
VI. 기업분석	53
성호전자	54
대한광통신	57
오이솔루션	60

Summary

빛의 시대로 가는 필연의 길

- AI 섹터에 대한 우려가 지속되는 가운데, AI 기술에 대한 발전은 지속되는 중이다. AI 인프라의 병목은 연산 능력에서 메모리로, 다시 GPU와 GPU를 연결하는 인터커넥트로 이동하고 있다. 모델의 크기가 단일 GPU의 연산, 메모리 용량을 넘어서면서, 이제 여러 GPU를 단순히 병렬 배치하는 것 이상으로 발전하고 있다. AI 연산에서 사용되는 Tensor Parallelism과 Expert Parallelism처럼 여러 GPU가 하나의 연산을 함께 수행하는 방식에서는 중간 결과와 토큰을 끊임없이 주고받아야 한다. 특히 MoE의 All-to-All 통신은 중간에서 합치거나 줄이기 어려워, 모델이 커질수록 GPU 간 대역폭 수요가 구조적으로 증가한다.
- 그러나 기존 구리 기반 인터커넥트는 속도가 높아질수록 신호를 보낼 수 있는 거리가 짧아지고, 더 많은 전기 신호선을 배치하기 위해 한정된 칩 외곽 면적 (Beachfront)을 소모한다. 대역폭을 두 배로 높일 때마다 더 짧은 전송 거리와 더 큰 전력, 더 복잡한 신호 보정이 요구되는 구조다. 시스템 배치 변경, PCB 경로 단축, 양방향 SerDes와 같은 기술을 통해 구리의 수명을 연장할 수는 있지만, 이는 물리적 한계를 우회하는 방법이지 제거하는 해법은 아니다.
- 따라서 구리에서 광으로의 전환은 어느 한 시점에 일괄적으로 발생하기보다, 대역폭과 전송 거리의 부담이 가장 큰 구간부터 단계적으로 진행될 전망이다. 이미 랙 간 연결을 담당하는 스케일아웃 네트워크에서 광 채택이 확대되고 있으며, 다음 시장은 랙 내부의 NVSwitch와 XPU를 연결하는 스케일업 인터커넥트다. 이후 광은 패키지 간 연결과 칩에 더욱 가까운 광 I/O 영역으로 침투하며, 광학 부품의 탑재 범위와 시스템당 콘텐츠를 동시에 확대할 가능성이 높다.
- 결국 광 I/O는 새로운 전송 기술의 등장이라기보다, 거대한 AI 연산 시스템을 하나의 컴퓨터처럼 작동시키기 위한 필수 인프라에 가깝다. 구리의 기술 개선과 아키텍처 변화로 상용화 시점은 지연될 수 있지만, AI 모델의 확장과 GPU 간 데이터 이동이 계속되는 한 광으로의 방향성 자체가 되돌아갈 가능성은 낮다. 향후 광통신 시장의 핵심은 단순한 트랜시버 출하량 증가를 넘어, 광이 컴퓨터에 얼마나 가까이 침투하고 시스템당 광 콘텐츠가 얼마나 확대되는가에 있다.
- 글로벌에서는 Lumentum, Coherent 같이 스케일업, 스케일아웃, 스케일러크로스에 걸쳐서 수혜를 보는 광원 회사와, Ayar Labs향 광원 매출로 턴어라운드 기대되는 Sivers 같은 회사들에 주목할 필요가 있다. 또, 광섬유 회사로는 데이터센터 내부향 매출에서 높은 지위를 지닌 Corning, Fujikura가 있다. CPO향 SiPh 생산과 관련해서 지위가 있는 GlobalFoundries, Tower Semi와 같은 8인치 파운드리 업체에도 주목할 필요가 있다. 결국 CPO 형태로 갈 것이라는 점에서 Co-Package를 하는 TSMC, Amkor, SPIL 같은 파운드리 및 OSAT 업체 역시 수혜가 기대된다. 웨이퍼 번인 등 테스트가 중요해진다는 점에서 AEHR 같은 장비사 역시 중장기적인 수혜가 기대된다.
- 국내에서는 액티브 얼라인 분야에서 선도적인 지위를 지니고 있는 성호전자에 관심을 가질 필요가 있다. 성호전자는 국내에서 유일하게 스케일업 시장에서 수혜가 기대되나, 최근 광시장 조정에도 동반하여 큰 폭 조정이 발생하였다. 펀더멘탈이 유효하다는 점에서 중장기적 관점에서 관심을 가질 필요가 있다. RF머트리얼즈의 경우, 루멘텀향 스케일러크로스 매출이 견조한 가운데, 증산이 지속되고 있다는 점에서 실적 개선을 기대한다. 또, 대한광통신도 수혜가 예상된다. 북미 하이퍼스케일러향 매출, 인캡 인수 효과뿐 아니라 중국내 광섬유 가격 상승 수혜에 따른 턴어라운드가 기대된다. 다만, 대한광통신은 밸류에이션에 대한 부담은 주의가 필요하다.

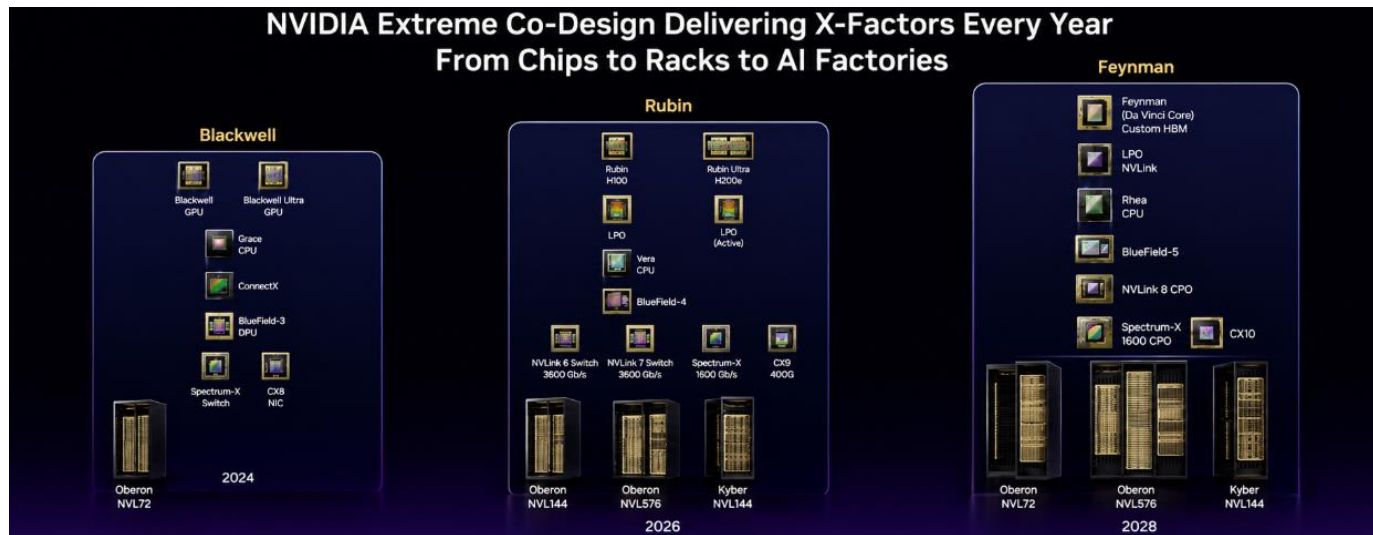
I. 스케일업으로 확장될 광 시장

2026년은 스케일 아웃 CPO의 본격화

시장의 관심은 스케일 아웃을 거쳐 스케일업으로

- 2026년 하반기는 스케일아웃 CPO가 본격적으로 채택되는 분기다. 엔비디아의 Spectrum-X(이더넷), Quantum-X(인피니밴드) 스케일 아웃 CPO, 그리고 브로드컴의 실리콘 포토닉스 CPO Bailly (Tomahawk 기반)를 시작으로 본격 채택 예정이다.
- 트렌드포스에 따르면, 800G 이상 고속 광 트랜시버 출하량은 24년 약 1,000만 대 → 25년 2,400만 대 → 26년 약 6,300만 대로 전망된다. CPO에 더해 트랜시버 출하량까지 동반 급증하면서, 데이터센터 투자에서 네트워킹의 비중이 커지고 광통신 섹터 전반이 재평가받는 계기가 형성되고 있다.
- 엔비디아는 랙과 랙을 잇는 스케일아웃을 넘어, 랙 내부 GPU 간 연결을 책임지는 스케일업 시장으로도 한 걸음씩 들어서고 있다. 블랙웰 GB200, GB300 NVL72 → 루빈 (Vera Rubin) NVL 144 → 루빈 울트라 NVL 576으로, 수십 ~ 수백 개의 GPU를 마치 한 대처럼 묶는 흐름이다. 시장은 그 연장선에서 스케일업용 NVSwitch와 서버 단에도 CPO가 점차적으로 채택될 것으로 기대하고 있다.

블랙웰 → 루빈 → 파인만까지의 엔비디아 아키텍처 로드맵



왜 스케일아웃이 아니라 스케일업인가?

병목은 연산에서 메모리로, 그리고 인터커넥트로 이동할 것

- AI 인프라의 병목은 연산 (FLOPS) → 메모리 (HBM) → 칩과 칩 사이 대역폭 (인터커넥트)으로 이동해왔다. Agent AI로 인해 폭발한 메모리 문제는 현재도 진행 중이다. 만약, 메모리만이 문제라면, GPU 1장에 HBM을 더 많이 붙이면 된다고 생각할 수 있으나, 현실에서는 그렇지 않다. HBM은 초광폭 버스로 GPU에 붙기 때문에, 스택 수가 다이 가장자리 (beachfront)와 CoWoS 인터포저 면적, 그리고 열에 의해 물리적으로 제한된다. 실제로 블랙웰 8개 → 루빈 울트라 16개로 GPU당 HBM의 붙는 개수가 늘어나고 있지만, 무한정 붙일 수는 없다.
- 게다가 메모리를 무한히 붙여도 지금의 연산은 이미 GPU 한 장으로는 수행이 불가능하다. 수조 파라미터 모델을 낮은 지연으로 돌리려면, 연산 자체를 텐서 병렬 (TP), 전문가 병렬 (EP)로 여러 GPU에 쪼개야 하는데, HBM 개수는 연산 속도를 끌어 올리지 못한다. 결국 용량 한계든 연산 분산이든, 여러 GPU를 하나로 묶는 것 외에는 기술적인 답이 없다. NVLink는 바로 이 묶인 GPU들을 '하나의 거대한 GPU'처럼 동작하게 만드는 기술이다.
- 여러 GPU를 초고속으로 묶기 때문에, 병목은 이제 'GPU와 GPU 사이의 대역폭', 'GPU와 HBM 사이의 대역폭' 즉, 인터커넥트로 옮겨간다. 이 단거리, 초고대역 연결을 키우는 것이 스케일업이며, NVLink 대역폭은 GPU당 0.9TB/s (NVLink4) → 1.8TB/s (NVLink5) → 3.6TB/s (NVLink6)로 세대마다 2배씩 증가하고 있다. AI 인프라의 다음 병목이자, 다음 투자처가 스케일업 인터커넥트인 이유다.

엔비디아 세대별 GPU 당 HBM 탑재량 추이

연도	아키텍처	GPU	HBM 세대	스택 구성	HBM/GPU	대역폭
2016	Pascal	P100	HBM2	4 × 4GB	16 GB	0.72 TB/s
2017	Volta	V100 (→18 32GB)	HBM2	4 × 4→8GB	16→32 GB	0.90 TB/s
2020	Ampere	A100 40GB	HBM2e	5 × 8GB	40 GB	1.6 TB/s
2021	Ampere	A100 80GB	HBM2e	5 × 16GB	80 GB	2.0 TB/s
2022	Hopper	H100	HBM3	5 × 16GB	80 GB	3.35 TB/s
2024	Hopper	H200	HBM3e	6 × 24GB (8-Hi)	141 GB	4.8 TB/s
2024-25	Blackwell	B200	HBM3e	8 × 24GB (8-Hi)	192 GB	8 TB/s
2025	Blackwell Ultra	B300	HBM3e	8 × 36GB (12-Hi)	288 GB	8 TB/s
2026E	Rubin	VR200/R200	HBM4	8 × 36GB (12-Hi)	288 GB	13 TB/s
2027E	Rubin Ultra	VR300/R300	HBM4e	16 × 64GB (16-Hi)	1,024 GB	32 TB/s
2028E	Feynman	—	차세대 HBM	>16 stacks	>1,024 GB	TBD

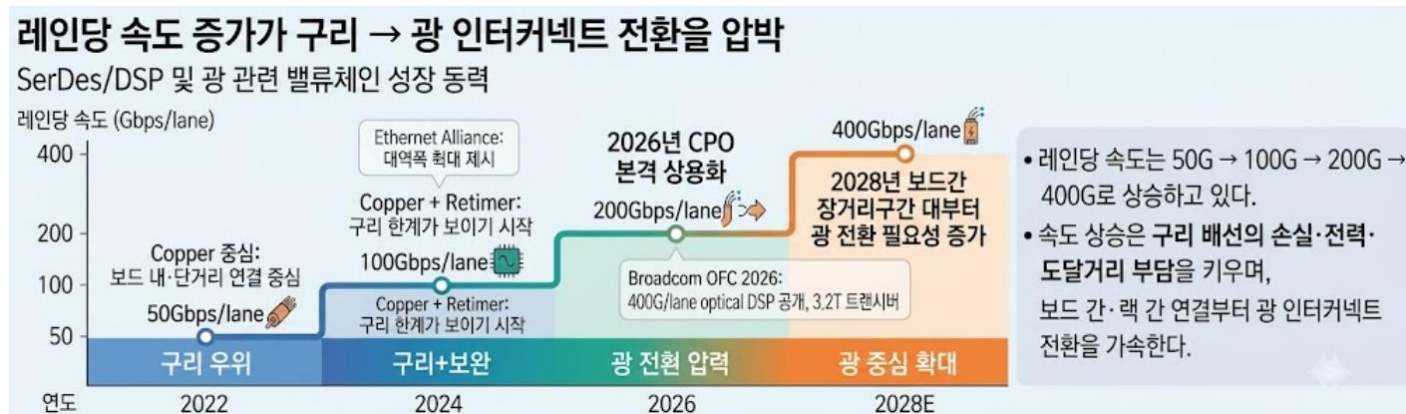
주: HBM/GPU = 패키지 (가속기) 단위 기준 (Rubin Ultra-B200은 패키지당 2~4다이 → 다이 단위로는 1/2~1/4). 2026E~전망치 (GTC 2025~2026 기준), Feynman은 custom HBM (C-HBM4E/HBM5 추정)으로 사양 미공개
 자료: NVIDIA, GTC 2025~2026, 업계 자료 종합

거대한 AI 팩토리 중요한 것은 대역폭

대역폭의 한계는 결국, 광의 시대의 도래를 뜻한다.

- 메모리 - GPU, GPU - GPU간 연결성 강화는 결국 대역폭 문제다. 지속적인 대역폭의 확대는 단순한 회로 설계 개선을 넘어 연결 소재의 변화까지 요구하고 있다.
- 기존에는 구리 기반 인터커넥트를 통해 보드 내 칩간, 보드 간 연결을 처리했다. 스케일업의 광화는 보드 간 연결에서 도달거리가 긴 보드 간 연결부터 구리를 광이 점차적으로 대체할 것으로 예상된다. 여러가지의 기술이 연구 및 개발되고 있으나, 여전히 보드 내에서의 연결은 구리를 중심으로 이루어질 가능성이 높다. 구리는 비용과 신뢰성 측면에서 높은 강점이 있기 때문이다.
- 하지만, 보드간 연결에서 구리의 입지는 점차 축소될 것으로 예상된다. 점차 요구되는 대역폭의 크기가 커짐에 따라 한계를 맞이할 것으로 예상되기 때문이다. 한계를 밀어붙이는 동력은 레인당 속도의 증가다. 최근 대역폭은 2년 마다 2배씩 빨라지고 있다. 레인당 속도는 22년 50Gbps/lane, 24년 100Gbps/lane, 26년 200Gbps/lane, 28년 E 400Gbps/lane으로 발전 중이다.

레인당 속도가 증가하면서 구리에서 광으로 인터커넥트가 전환되고 있는 중



결국에 한계에 도달할 구리선

구리선의 한계를 만드는 두 가지 도달 거리(reach), 칩의 외곽 둘레(Beachfront)

- 대역폭이 커질 수록 구리선의 한계에 봉착하는 이유는 크게 두 가지다. 첫째, 도달 거리(Reach)다. 구리선은 표피효과(skin effect, 주파수가 높아질수록 전류가 도체 표면으로 몰리는 현상)를 비롯한 주파수 의존 손실 때문에, 신호 속도가 빨라질수록 멀리 보내기가 급격히 어려워진다. 경험적으로 대역폭이 두 배가 될 때마다 구리가 도달할 수 있는 거리는 대략 절반으로 짧아져 왔다. 대역폭이 커질수록 구리가 닿을 수 있는 거리 자체가 계속 줄어드는 것이다.
- 둘째, 칩의 외곽 둘레(Beachfront)다. Beachfront란 데이터 신호(I/O)가 드나들 수 있는 칩 가장자리의 폭을 말한다. 구리는 채널 하나에 선 하나가 필요하므로, 대역폭이 커질수록 더 많은 선을 깔 둘레가 필요해진다. 차량이 늘수록 고속도로를 4차선에서 8차선, 16차선으로 넓혀야 하는 것과 같다. 문제는 칩 둘레가 물리적으로 한정되어 있어, 2년마다 두 배로 불어나는 대역폭 수요를 따라가기 어렵다는 점이다.
- 결국, 대역폭이 커질수록 구리의 도달거리는 짧아지고, 둘레의 한계에 부딪히게 된다. 이는 시간의 문제일 뿐 구리가 광으로 대체되는 것은 물리적 한계에 의한 구조적 필연이다.

레인당 속도가 증가할 때마다 줄어드는 구리선의 거리

년도	레인당 속도	구리선 거리
2022	50 Gb/s	~3m
2024	100 Gb/s	~2m
2026	200 Gb/s	~1m
2028E	400 Gb/s	~0.5m

자료: 대신증권 Research Center

대역폭 수요대비 더딘 전기적 탈출 용량 증가



주: 전기적 탈출 용량(Electrical Escape Capacity)은 beachfront를 통한 I/O 통신 용량을 뜻함

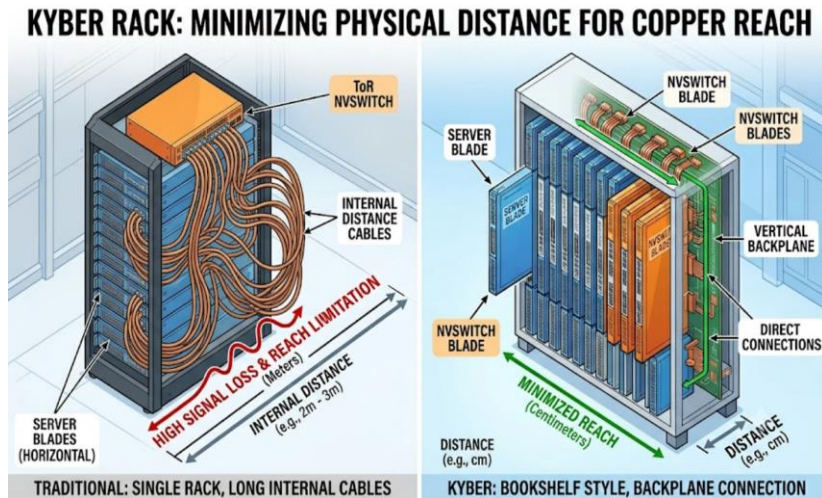
자료: 대신증권 Research Center

그럼에도 계속 살아남는 구리선

구리선의 한계를 보충해주는 다양한 방식

- [Reach] 파인만(2028)에 레인당 400Gb/s로 구리 도달거리가 ~0.5m로 짧아지며 NVSwitch 및 서버의 구리가 광으로 전면 대체될 전망이었으나, 로드맵상 L1(Leaf)은 구리, L2(spine)는 CPO의 혼용으로 사용될 예정이다. 엔비디아는 케이블 대신 블레이드를 책장처럼 수직으로 세워 밀도를 극대화하고 케이블 없이 미드플레인 PCB에 직접 서버와 NV Switch를 꽂는 방식으로 전기 경로를 단축시켜, 구리 도달 거리의 한계를 아키텍처로 우회했다.
- [Beachfront] 기존에는 송신선과 수신선에서 각각 단방향으로 정보를 보내기만 했었다. 엔비디아는 속도를 224G에 묶은 채, 전화선처럼 같은 구리 한 쌍에 송,수신을 동시에 실어 (bidirectional SerDes) 같은 면적으로 대역폭 2배를 만들고, 스위치를 144개로 쪼개 총량을 확보했다. 아키텍처를 변화시켜, Reach의 문제를 해소했다면, Beachfront는 정보를 처리하는 방식을 통해서 구리선의 한계를 우회했다. 다만, 우회는 우회일 뿐이다. 파인만에서는 이미 한계에 도달한 beachfront를 2배 더 늘려야 하는 상황으로 구리 케이블은 광으로 변화될 가능성이 높다.

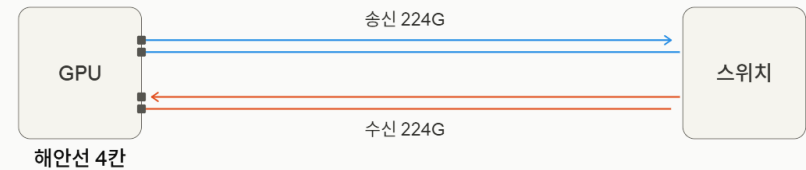
아키텍처 변경으로 Reach 문제를 해결함



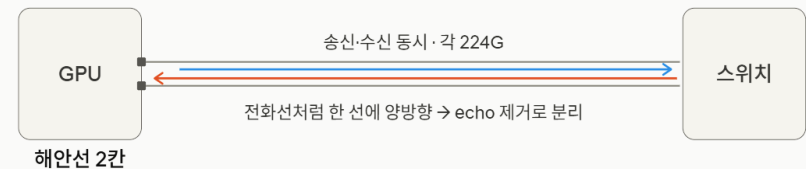
자료: 산업자료, 제미나이 재구성, 대신증권 Research Center

소프트웨어 처리로 Beachfront 문제를 해결함

기존 방식 — 송·수신 회선을 따로



Bidirectional SerDes — 같은 쌍에 동시



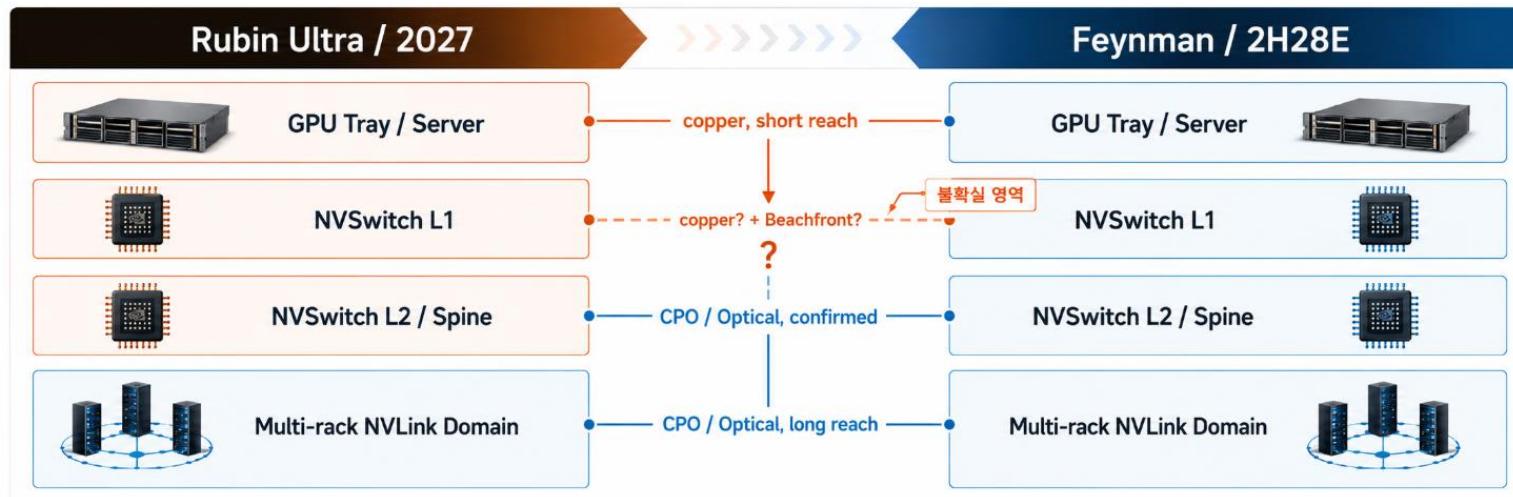
자료: 산업자료, 클라우드 재구성, 대신증권 Research Center

파인만 아키텍처에 대한 예상

파인만 아키텍처가 열리고, 빛의 길이 열리는 서막이 시작된다.

- 엔비디아는 2026 GTC에서 파인만 (Feynman) 아키텍처가 구리와 광을 혼용하는 형태로 진행될 것이라고만 짧게 언급했다. 구체적인 설명이 없었던 만큼, 파인만 단계의 스케일업 인터커넥트가 어떻게 구성될지는 아직 미지수다. 당사는 지난 자료에서 루빈 울트라(Rubin Ultra)부터 NVSwitch가 2계층으로 분리되어, 서버에 가까운 L1은 구리선을, NVSwitch간 연결을 담당하는 L2는 CPO를 적용하는 형태로 파인만 아키텍처가 구성될 것으로 예상하였다.
- 다만 beachfront 제약을 감안하면, 당장 서버단은 아니더라도 고성능 스펙에서는 NVSwitch L1과 L2 모두 CPO로 구성될 가능성이 높다. 엔비디아가 구리와 CPO 중 택일할 수 있다고 언급한 점은, 파인만 아키텍처가 단일 구성에 고정되기보다 스펙 등급에 따라 인터커넥트가 차별화된 복수의 제품군으로 운영될 수 있음을 시사하는 것으로 추정된다. 파인만 출시 시점이 2H28인 만큼 현재는 다양한 구성안이 논의되는 단계로 보이며, 사양이 구체화되는 2027년 전후에는 관련 장비사들의 발주가 선행될 것으로 예상된다.

2028년 파인만 아키텍처에 대한 전망



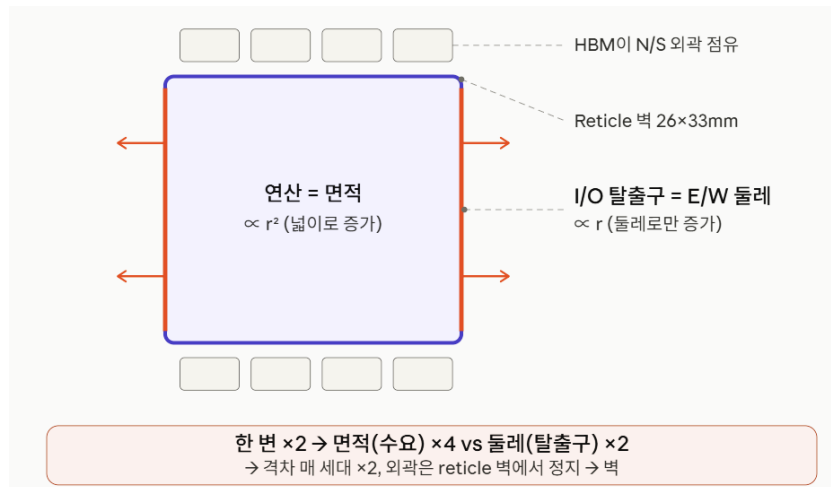
자료: GPT로 재구성, 대신증권 Research Center 추정

물리 구조상 필연적인 한계에 봉착할 구리

구리선의 한계는 '언제'의 문제

- AI 가속기의 연산 능력은 사실상 실리콘 면적에 비례한다. 그런데 이 연산은 HBM이 공급하는 메모리 대역폭이 있어야 돌아가므로, 연산이 커질수록 더 많고 더 빠른 HBM이 필요하고, 이 HBM이 패키지 외곽 (주로 칩의 남,북 변)을 점점 더 많이 차지한다. 동시에 한 패키지에서 나온 연산 결과는 다른 GPU 패키지와 끊임없이 주고받아야 하므로 (스케일 업 & 스케일 아웃), 패키지를 빠져나가야 하는 데이터량 자체도 함께 폭증한다.
- 문제는 이 늘어나는 물량이 '한정된 칩 외곽(둘레)'이라는 같은 출구로 빠져나가야 한다는 점이다. 구리로 더 많이 보내려면 레인당 속도를 높여야 하는데, 속도를 올릴수록 전력 소모와 발열이 급증하고 신호가 갈 수 있는 거리도 짧아진다. 반면 빛은 서로 다른 파장(색)에 각각 다른 신호를 실어 하나의 광섬유로 묶어 보내는 WDM으로, 칩 외곽 넓히지 않고도 파장 수만 늘려 대역폭을 키울 수 있다. - 즉 대역폭이 칩의 둘레 한계에서 분리된다.
- 결국 패키지 안의 연산과 HBM이 커질수록 같은 면적·같은 외곽에서 더 많은 데이터를 효율적으로 송수신해야 하고, 이를 감당할 수 있는 건 구리가 아니라 광이다. 구리에서 광으로의 전환은 '할지 말지'가 아니라 '언제'의 문제인, 구조적으로 예고된 병목 해소다.

패키지내 연산이 커질수록 I/O에 커지는 부하 증가



자료: 산업자료, 클라우드 재구성, 대신증권 Research Center

같은 면적에서 더 많은 정보량을 전송하는 WDM 기술



자료: 산업자료, 클라우드 재구성, 대신증권 Research Center

대역폭의 확대는 곧 광의 TAM 확대를 의미한다.

다양한 인터커넥트 층위와 순차적으로 침투해나가고 있는 광

- 레인(lane)당 전송 속도가 높아질수록 광은 구리를 점진적으로 대체해 왔다. 가장 먼저 구리가 광으로 대체된 구간은 도달 거리(reach) 제약이 큰 스케일 아웃 분야이다. 이전부터 광트랜시버 (Pluggable)로 전환되어서 사용되어 왔다. 그 트랜시버도 beachfront의 한계에 부딪히면서 이제 스케일 아웃은 CPO로 전환되고 있다.
- CPO의 침투 영역은 거리축을 따라 안쪽으로 점차 확대되고 있다. 랙과 랙을 잇는 스케일 아웃이 현재 상용화된 부분이고, 거기서 한 단계 더 들어가 개별 랙 안에서 GPU를 잇는 NVLink(스케일업)까지 침투 대상으로 확대하고 있다. 이후에는 점차 같은 보드 위에 있는 GPU 패키지 간의 인터커넥트, GPU 패키지 내에서 GPU – HBM간의 인터커넥트까지 구리에서 광으로 대체하는 여러 기술이 개발 및 논의되고 있다.

레인당 속도가 증가하면서 구리에서 광으로 인터커넥트가 전환되고 있는 중

패키지 안 · die↔die, die↔HBM
μm~mm · 전기(구리) — 글래스코어·미세피치로도 구리

같은 보드 · GPU↔GPU (별개 패키지)
cm · 구리 (아직 전기)

랙 scale-up · GPU↔NVSwitch (NVLink 스파인)
~1m · 구리 한계 → 광 (다음 변곡 = scale-up 광학화)

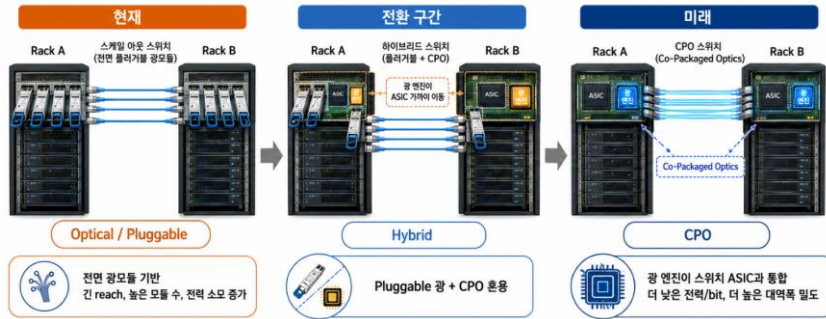
◀ 오늘의 구리/광 경계 — 위=구리, 아래=광 · 경계는 위로 전진

scale-out · Switch↔Switch ~ 랙↔랙
m+ · 광 → CPO 전환 중 (오늘 이미 광)

다양한 인터커넥트 층위

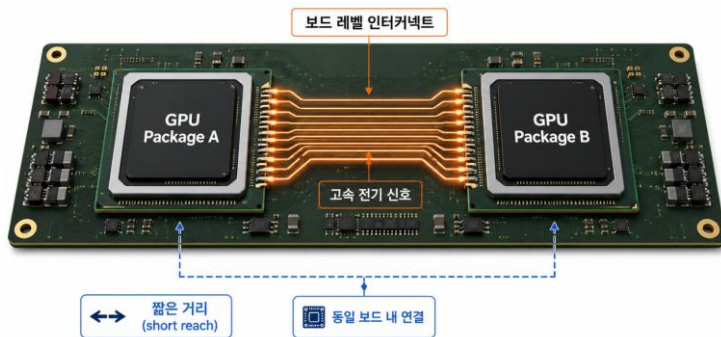
스케일 아웃에서 스케일 업. 보드위 패키지까지 점차 범위를 확장해나가는 광

스케일 아웃 랙과 랙의 연결. 광 → CPO로 전환



자료: 산업자료, GPT로 재구성, 대신증권 Research Center

같은 보드 위 GPU 패키지간의 인터커넥트



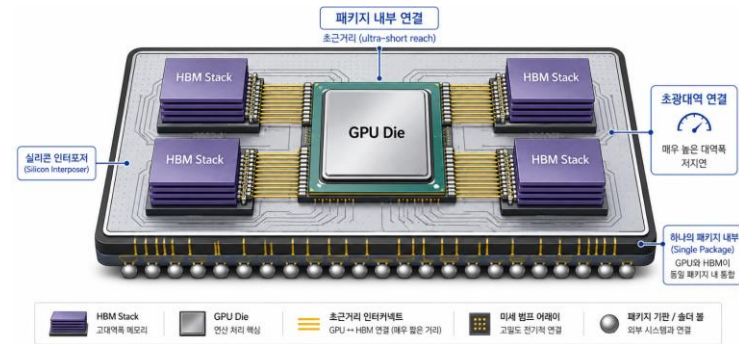
자료: 산업자료, GPT로 재구성, 대신증권 Research Center

GPU ↔ NVSwitch 보드간 연결에서 구리를 대체해 나가는 광



자료: 산업자료, GPT로 재구성, 대신증권 Research Center

GPU 패키지내 GPU - HBM간의 인터커넥트



자료: 산업자료, GPT로 재구성, 대신증권 Research Center

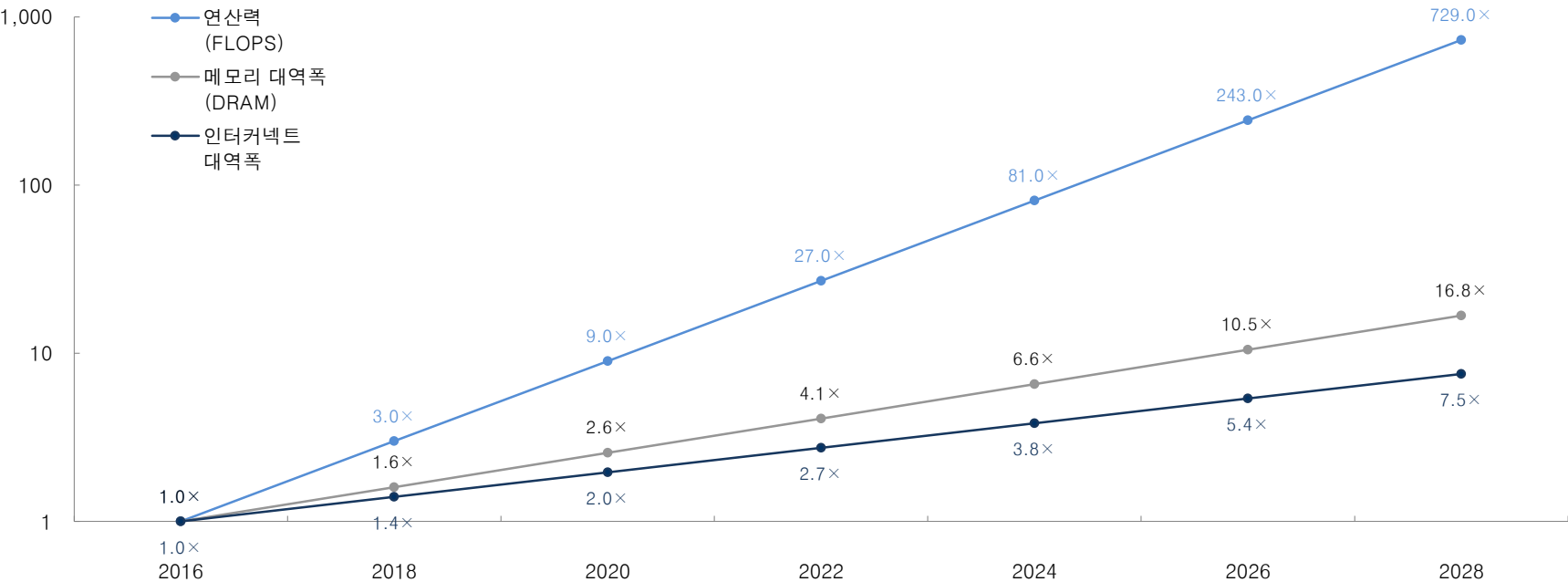
층위별로 다른 발전속도를 보이는 AI 하드웨어

연산력 발전대비 더딘 메모리 대역폭과 인터커넥트 대역폭 성능 개선 속도

하드웨어별로 성능 개선폭이 다름

상대 성능 (2016=1.0, 로그)

AI 하드웨어 상대 성능 추이 (로그 스케일)



자료: Gholami, Yao, Kim, Hooper, Mahoney, Keutzer (2024). "AI and Memory Wall", 대인증권 Research Center

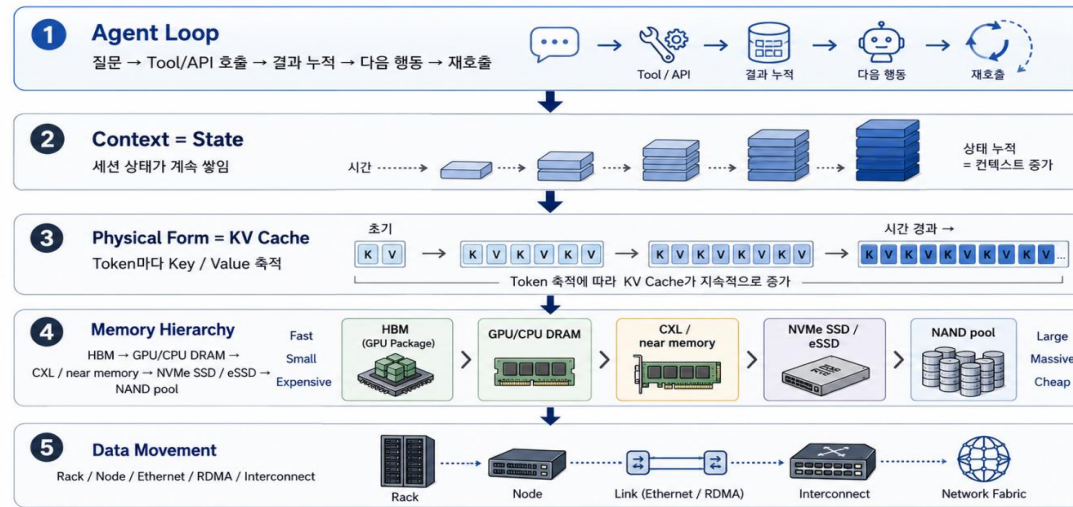
II . 광전환의 이유 – 메모리 병목

Agent AI 시대의 도래와 폭발하는 메모리 수요

프롬프트에서 컨텍스트로 메모리 수요를 폭발하는 KV Cache

- 2025~26년 AI의 무게중심이 '한 번 묻고 답하는' 생성형에서 '일을 끝낼 때까지 굴러가는' 에이전트로 넘어왔다. 바야흐로 Agent AI 시대의 개막이며, 이는 메모리의 구조적 수요와 직결된다. 생성형 AI의 입력이 일회성이었던 것과 달리, 에이전트는 하나의 목표를 완수하기까지 여러 단계를 스스로 돌며 수십 개의 API·툴을 호출하고 그 결과를 다시 컨텍스트에 쌓아 다음 판단의 근거로 삼는다. 단계가 누적될수록 메모리에 들고 있어야 할 문맥도 함께 불어나, 컨텍스트는 세션이 끝날 때까지 지속적으로 증가하는 상태가 된다.
- 트랜스포머는 토큰을 하나씩 생성하는데, 다음 토큰을 만들 때마다 앞서 나온 모든 토큰을 참조해야 한다. 매 단계 그 전체를 처음부터 다시 계산하면 길이가 늘수록 중복 연산이 폭증하므로, 각 토큰의 Key·Value를 한 번만 계산해 저장해두고 재사용한다. 이 저장 형태가 KV Cache이며, 캐시는 컨텍스트 길이에 정비례해 커진다.
- KV Cache의 급증은 메모리 자체(HBM·DRAM·eSSD)의 수요를 폭발할 뿐 아니라, 그것을 노드·랙 너머로 옮기는 데이터 이동 및 인터넥트(구리·광·CPO) 시장의 촉매제 역할을 한다.

Agent AI – 컨텍스트 엔지니어링이 촉발한 폭발적인 메모리 수요



KV Cache가 촉발하는 HBM 초과 수요

용량을 넘어, 대역폭까지 잠식하는 KV Cache

- KV Cache의 크기는 모델 구조상 거의 공식처럼 정해진다. 토큰 하나가 추가될 때마다 레이어 수, KV 헤드 수, 헤드 차원, 정밀도만큼 메모리가 필요하고, 여기에 컨텍스트 길이와 동시 사용자 수가 곱해진다. 문제는 Agent AI 환경에서는 긴 문맥, 여러 사용자, 오래 유지되는 세션이 동시에 늘어난다는 점이다. 즉 KV Cache 수요는 하나씩 더해지는 것이 아니라 곱으로 커진다.
- 더 큰 병목은 디코드 구간에서 발생한다. AI가 답변을 한 토큰씩 생성할 때마다 GPU는 기존에 쌓인 KV Cache를 계속 HBM에서 읽어야 한다. 이때 필요한 데이터 이동량이 연산량보다 커지면, 속도는 GPU 성능이 아니라 HBM 용량과 대역폭에 묶인다. 결국 Agent AI가 확산될수록 HBM은 "더 많이 저장해야 하는 압박"과 "더 빨리 읽어야 하는 압박"을 동시에 받는다.
- 임계점은 KV Cache가 모델 가중치보다 커지는 순간이다. 예를 들어 70B 모델의 가중치는 약 140GB인데, 100만 토큰 수준의 긴 요청에서는 KV Cache가 이를 넘어설 수 있다. 이제 메모리 문제의 핵심은 단순히 모델을 올리는 것이 아니라, 디코드 과정에서 거대한 KV Cache를 얼마나 빠르게 HBM에서 GPU로 공급하느냐가 된다. HBM4 이후 하이브리드 본딩 등 차세대 HBM 기술이 필요한 이유도 여기에 있다.

구조적으로 확장되는 KV 캐시

에이전트형 AI는 컨텍스트 길이, 동시성, 세션 지속 시간을 통해 메모리 수요를 기하급수적으로 늘립니다.

토큰당 KV 캐시 = $2 \times \text{레이어 수} \times \text{KV 헤드 수} \times \text{헤드 크기} \times \text{정밀도 (바이트)}$
 총 KV 캐시 = (토큰당 KV) $\times$ 시퀀스 길이 $\times$ 배치 크기



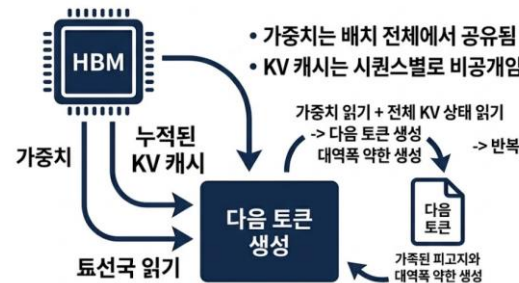
에이전트형 AI 수요 =
 컨텍스트 길이 $\times$ 동시 사용자 수 $\times$ 세션 지속 시간

KV 캐시는 시퀀스 길이 및 배치 크기와 선형으로 증가하지만, 에이전트형 워크로드는 모든 축을 동시에 확장합니다.

자료: 산업자료, 제미나이 재구성, 대신증권 Research Center

추론 중 디코드는 메모리에 속도가 종속됨

생성되는 각 토큰은 HBM에서 반복적인 읽기가 필요합니다.



디코드 병목 현상 = 컴퓨팅이 아닌 메모리 대역폭

동시성이 증가함에 따라 KV 캐시 트래픽이 지배적입니다.

자료: 산업자료, 제미나이 재구성, 대신증권 Research Center

용량과 대역폭 압박을 둘 다 받는 HBM

KV 캐시는 용량과 대역폭 모두를 압박합니다.



새로운 메모리 벽은 모델 가중치뿐만 아니라, KV 캐시의 성장입니다.

70B 모델 가중치 $\approx$ 140GB
 1M 토큰 KV 캐시 (단일 요청)는 모델 가중치를 초과할 수 있습니다

HBM3E $\rightarrow$ HBM4
 스택당 용량 $\uparrow$
 인터페이스 폭 및 대역폭 $\uparrow$
 업그레이드는 선택이 아닌 필수가 됩니다

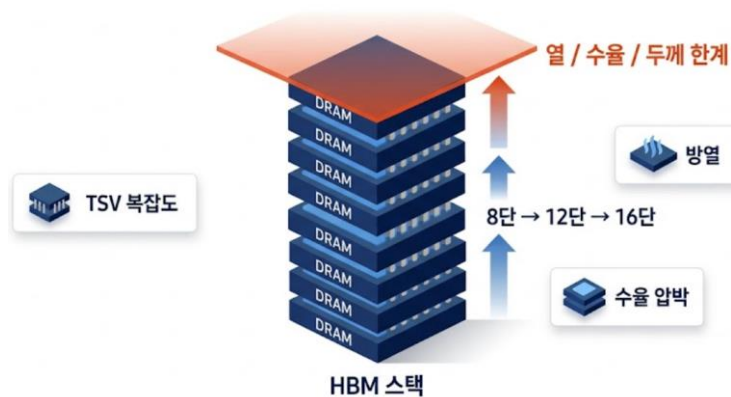
자료: 산업자료, 제미나이 재구성, 대신증권 Research Center

HBM만으로는 감당할 수 없는 메모리 초과수요

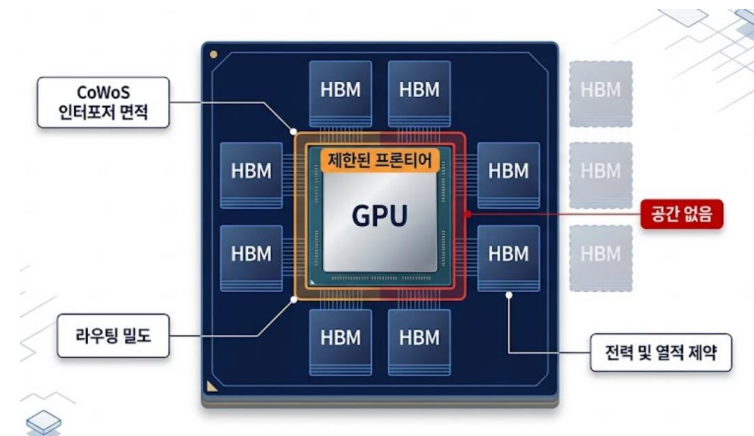
수요는 기하급수, HBM 증설은 물리적으로 제약

- Agent AI가 촉발하는 메모리 수요를 전부 HBM이 흡수하기 어렵다. HBM은 GPU에 가장 가까운 고대역폭 메모리지만, 가장 비싸고 패키징 난도가 높으며 물리적으로 붙일 수 있는 양도 제한된다. KV Cache 수요는 컨텍스트 길이, 동시 사용자 수, 세션 지속시간의 곱으로 커지지만, HBM 증설은 같은 속도로 따라가기 어렵다.
- 첫 번째 한계는 HBM 스택의 높이다. HBM은 DRAM 다이를 수직으로 쌓아 용량과 대역폭을 높이지만, 층수가 올라갈수록 TSV 연결, 열 방출, 수율, 패키지 두께 문제가 커진다. 8단에서 12단, 16단으로 진화하고 있지만, 무한정 높게 쌓을 수는 없다.
- 두 번째 한계는 GPU 주변 실장 면적이다. HBM은 초광폭 인터페이스로 GPU와 매우 가깝게 붙어야 하므로 인터포저와 패키지 위에서 큰 공간을 차지한다. GPU 주변의 beachfront 영역, CoWoS 면적, 전력과 열 제약이 모두 스택 수를 제한한다. 결국 메모리 병목은 HBM을 더 붙이는 문제뿐 아니라, HBM, DRAM, CXL, eSSD로 데이터를 어떻게 계층화하고 빠르게 이동시키느냐의 문제로 확장된다.

무한히 상승하기 어려운 HBM의 높이



GPU 주변 실장 면적의 한계

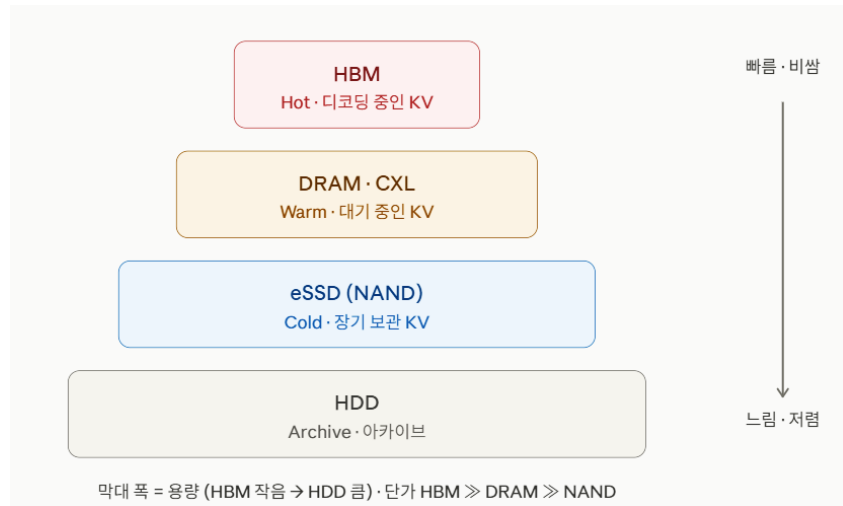


KV캐시를 HBM 밖으로 - KV 캐시 오프로딩

메모리 계층화로 초과수요에 대응

- HBM 증설이 물리적으로 막히면 선택지는 둘 뿐이다. 캐시를 줄이거나, HBM 밖으로 밀어내거나, 모델, 알고리즘 차원의 압축 (MQA, GQA, 양자화 등)으로 캐시 자체를 줄이는 길도 있으나 한계가 분명하다. 현실적 해법은 자주 사용하지 않는 KV캐시를 더 싸고 큰 메모리 계층으로 내보내는 것이다. 이것이 KV 캐시 오프로딩이다.
- 핵심은 중요도에 따라 데이터를 계층화하는 것이다. 지금 디코딩에 쓰이는 중요한 캐시는 HBM에 두되, 당장 쓰지 않는 캐시는 GPU 옆 DRAM, CXL 메모리로 장기 보관할 캐시는 eSSD (NAND), HDD로 증위를 두고 저장한다. GPU에 가장 가까운 비싼 자원 (HBM)을 진짜 필요한 데이터에만 할당하고, 용량은 다른 메모리들로 확보하는 구조다.
- 오프로딩은 두 가지 문제를 동시에 해결한다. 1) 용량: HBM 한 장에 안 들어가던 긴 컨텍스트, 다중 세션을 아래 계층, 용량으로 수용한다. 2) 비용: GB당 단가가 HBM > DRAM > NAND 순으로 떨어지므로 같은 서비스를 훨씬 저렴하게 구동한다. 주론을 'GPU를 더 사는 문제'에서 '메모리를 어떻게 계층화하고 빠르게 옮기느냐의 문제'로 바뀌게 되는 것이다.

데이터 중요도별 메모리 피라미드



자료: 산업자료, 클라우드 재구성, 대신증권 Research Center

메모리별 스펙표

계층	최신세대	대역폭	용량(소자/모듈)	지연	GB당단가
HBM	HBM4 (25) → HBM4E (27)	2~2.5 TB/s	36~64 GB/스택	~ns	최고 (DRAM의 ~5x)
서버 DRAM	DDR5-6400 RDIMM	~51 GB/s/채널	64~256 GB	~수십 ns	높음
CXL 메모리	CXL 3.0 / PCIe Gen5~6	~64 GB/s (x16)	수백 GB~TB	~수백 ns	중간 (풀링시 실효↓)
eSSD (NAND)	PCIe Gen5 → Gen6 QLC	14~28 GB/s	~수십 TB (최대 60TB+)	~μs	낮음 (HDD의 ~3x)
HDD	HAMR nearline	~250 MB/s	30~40 TB	~ms	최저

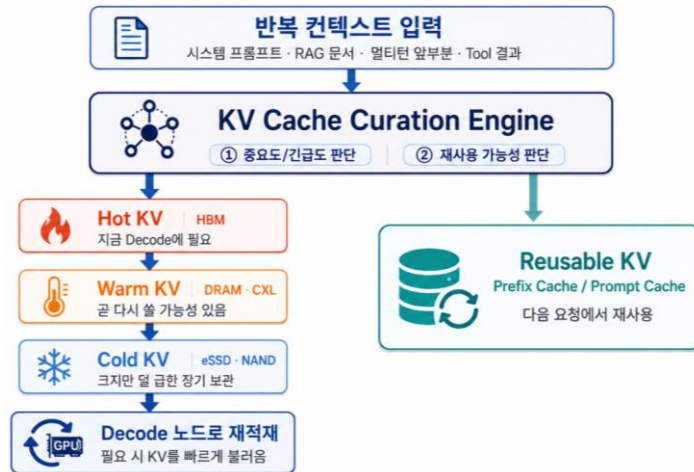
자료: 산업자료, 대신증권 Research Center

KV캐시 오프로딩의 두 엔진 – 계층화와 재사용

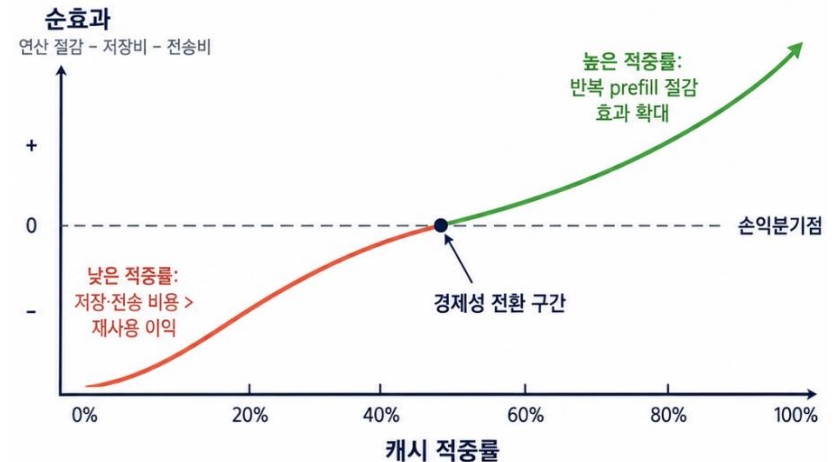
KV 캐시 오프로딩의 경제성을 만드는 요인 – 캐시 적중률 (Cache Hit Ratio)

- KV 캐시 오프로딩의 경제성은 캐시를 얼마나 잘 큐레이션 하느냐 – 무엇을 남기고 어디에 두느냐-에서 나오며, 두 엔진이 이를 떠받친다. 첫째, 계층화다. 데이터의 중요도에 따라 HBM, DRAM, CXL, eSSD로 나뉘, 비싼 HBM은 지금 쓰는 데이터에만 두고 덜 급한 건 아래 계층에 싸게 보관한다.
- 둘째, 재사용이다. 시스템 프롬프트, RAG로 붙는 문서, 멀티턴 대화의 앞부분처럼 반복되는 컨텍스트의 KV는 한 번만 계산해 저장하고 다음 요청에서 다시 쓴다. 같은 토큰을 매번 prefill하지 않으니 연산이 절감되고 첫 토큰까지의 지연(TTFT)이 짧아진다. 이 축의 성패를 가르는 변수가 바로 캐시 적중률이다. 저장한 KV가 재사용되는 비율로, 높으면 한 번의 연산이 여러 번 절감으로 돌아와 크게 남고, 낮으면 아무도 안 쓰는 KV에 자리값만 치르는 셈이라 차라리 재계산이 싸다. Agent AI는 적중률이 본질적으로 높은 워크로드라 저장해둘 KV (=메모리의 수요)가 구조적으로 커진다.
- 핵심은 큐레이션이 KV 용량을 줄이는 게 아니라 '계층'으로 분산시킬 뿐이라는 점이다. 총량 자체가 컨텍스트 길이, 동시 사용자, 세션 지속의 곱으로 폭증하므로, 메모리의 총 수요는 HBM, DRAM, CXL, NAND로 나뉠뿐 구조적으로 증가한다. 동시에 메모리 계층별로 흠뻑려진 KV를 빠르게 가져오는 것이 성능의 중요한 동인이 된다. 즉, KV 캐시 오프로딩은 '어디에 저장하느냐'가 아니라 '얼마나 빨리 싸게 옮기느냐'의 문제로 수렴한다.

KV캐시 큐레이션 엔진 구조도



KV캐시 적중률별 경제성 곡선



Prefill – Decode 분리 – 메모리 수요는 인터커넥트 수요로 이어진다.

KV 캐시 오프로딩 = 데이터를 어떻게 빠르게 옮길 것이냐

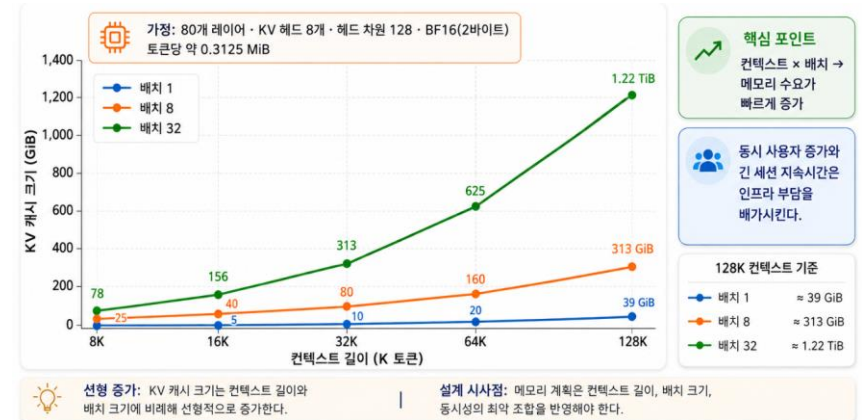
- 추론은 성격이 다른 두 구간으로 나뉜다. Prefill은 입력 전체를 한 번에 처리해 KV를 만들어내는 연산 집약(compute-bound) 구간이고, Decode는 그 KV를 계속 읽어가며 토큰을 하나씩 뽑는 메모리 집약(memory-bound) 구간이다. 두 구간의 병목이 다르기에, 최신 추론 스택은 이를 서로 다른 노드 GPU 풀로 분리 (Prefill-Decode disaggregation)해 각각 최적화한다.
- 분리의 대가는 데이터의 이동이다. Prefill 노드가 만든 KV캐시를 Decode 노드로 넘겨야 하고, 오프로딩으로 아래 계층에 보관한 KV도 다시 GPU로 끌어와야 한다. 즉 KV 캐시는 '어디에 저장하느냐(메모리)'만큼이나 '얼마나 빨리 옮기느냐(인터커넥트)'의 문제로 이어진다. 캐시가 커질수록 노드 및 랙을 가로지르는 전송량이 폭증하고, 이는 스케일업, 스케일 아웃 네트워크 대역폭에 그대로 부하로 쌓인다.
- KV 캐시 오프로딩은 메모리(HBM, DRAM, eSSD)의 수요인 동시에, 그 데이터를 노드 너머로 실어 나르는 구리 → 광 → CPO 인터커넥트 수요의 촉매다. 메모리 계층이 깊어질수록 계층 간을 잇는 대역폭이 병목이 되고, 전기 신호(구리)의 거리, 전력 한계가 광,CPO로의 전환을 앞당긴다. 즉, KV캐시의 폭발적인 증가가 만든 메모리의 수요는 곧 광통신의 수요로 이어지는 셈이다.

AI 추론 – Prefill – Decode 분리 구조



자료: 산업자료, GPT 재구성, 대신증권 Research Center

KV 캐시 크기 민감도



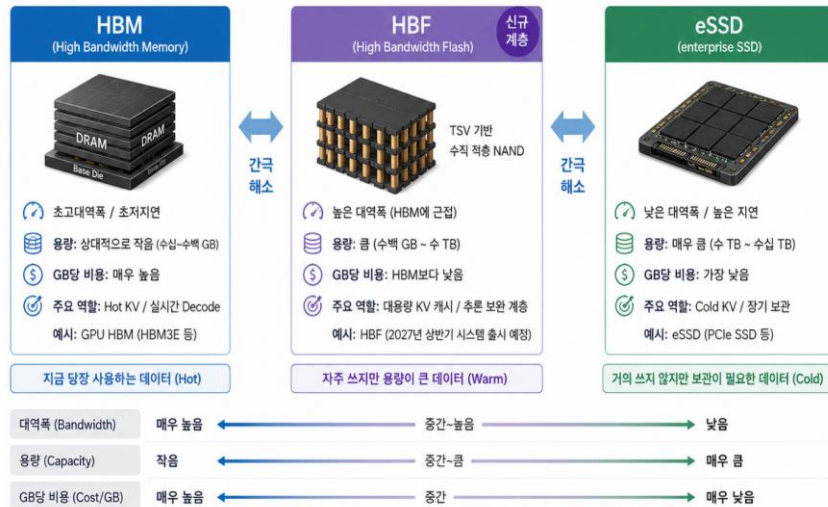
자료: 산업자료, GPT 재구성, 대신증권 Research Center

빈 계층을 메우는 새 메모리 – HBF (High Bandwidth Flash)

HBM의 속도와 NAND의 용량 사이. 태동하는 새로운 적층 칩

- 계층화의 약점은 HBM과 eSSD 사이의 간극이 크다는 점이다. 이 틈을 메우려는 시도가 HBF다. NAND를 TSV로 수직 적층해 HBM에 준하는 대역폭을 내면서 용량은 NAND 경제성으로 확보하는 새 메모리 계층으로, KV캐시처럼 '용량은 커야 하는데 대역폭도 포기 못 하는' 데이터에 정확히 겨냥된다.
- 샌디스크는 HBF가 eSSD를 대체하는 게 아니라, '주론에 훨씬 큰 용량을 다른 방식으로 가져오는' 보완재라고 언급하였다. HBF 생산의 시점은 HBF용 NAND 다이는 2026년 말, 컨트롤러를 포함한 시스템은 2027년 상반기로 예정되어 있다. (HBF는 샌디스크와 SK하이닉스 공동 표준화에서 출발)
- HBF가 TSV 기반 수직 적층이라는 점에서 HBM과 밸류체인을 공유한다. 낸드용 하이브리드 본딩, 적층 계면 결함을 치유하는 어닐링, 공정 복잡도 상승에 따른 세정, 테스트 수요가 확대되며, 관련 업체들에 대한 수혜가 기대된다.

HBF: HBM과 eSSD 사이의 간극을 메우는 신규 메모리층



자료: 산업자료, GPT 재구성, 대신증권 Research Center

HBF 글로벌 및 국내 밸류체인

공정 단계	핵심 변화	국내 수혜주 (상장)	글로벌 수혜 기업
하이브리드 본딩 (Cu-Cu)	TC본드 → 하이브리드 본드 전환 (수백 nm급 정밀 접합)	한미반도체, 한화비전 (한화세미텍)	BESI (NL), Applied Materials (US), EV Group (AT), ASMPT (HK), SUSS MicroTec (DE), Kulicke & Soffa (US)
어닐링 (계면 결함 치유)	적층 계면 결함 (Void) 치유·저온 소수 어닐링	HPSP, 예스티, AP시스템	Applied Materials (US), Tokyo Electron (JP), Kokusai Electric (JP), SCREEN (JP)
웨이퍼 박막화 다이싱	박막 웨이퍼 그라인딩 + 레이저 풀러닝 전환	이오테크닉스	DISCO (JP), Tokyo Seimitsu·ACCURETECH (JP)
세정·CMP 장비	공정 스텝 증가 → 세정·CMP 수요 확대	피에스케이홀딩스, 케이씨텍, 제우스	SCREEN (JP), Tokyo Electron (JP), Lam Research (US), Ebara (JP)
소재 (슬러리·식각액·PR·전구체)	CMP 슬러리·식각액·PR·전구체 사용량 증가	한솔케미칼, 솔브레인, 동진세미켄, 이엔에프테크놀로지	Entegris (US), Fujimi (JP), Merck·Versum (DE), JSR (JP), Tokyo Ohka·TOK (JP), DuPont (US)
테스트·번인 (ATE/KGD)	적층 수율 관리·전수 검사 번인 강화	테크윙, 유니테스트, 와이씨, 다이이	Advantest (JP), Teradyne (US), Cohu (US), ACHR (US)
프로브카드·소켓	적층 검사용 프로브카드·테스트 소켓 수요 증가	ISC, 리노공업, 티에스이, 티에프이	FormFactor (US), Technoprobe (IT), Yamaichi (JP), Smiths Interconnect (UK)

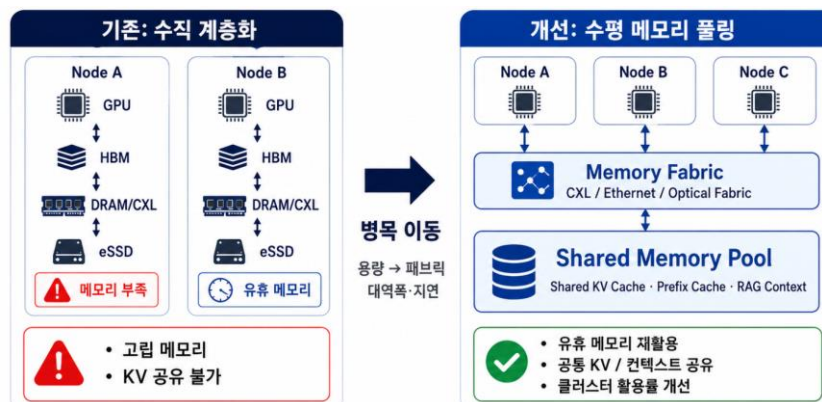
자료: 대신증권 Research Center

계층화만으로는 부족하다 - '수직'을 넘어 '수평'으로

개별 사용자(호스트)를 넘어서 공용으로 사용되는 메모리

- 지금까지의 오프로딩은 한 노드 안에서 데이터를 위아래로 옮기는 수직 계층화였다. 캐시를 수직으로 계층화함으로써 용량,비용은 풀었지만, 두 가지 구조적 한계가 그대로 남는다. 계층을 깊게 파는 것만으로는 부족하고, 연결성이 같이 좋아져야 한다.
- 첫 번째 한계. 고립 메모리(stranded memory): 메모리가 특정 호스트,GPU에 물리적으로 묶여 있어, 옆 노드가 모자라도 남의 노드에서 남는 메모리를 끌어 쓸 수 없다. AI 인프라에서 가장 비싼 자원이 한쪽에선 쉬고 있고, 한쪽에서는 부족해 비효율이 구조적으로 발생한다. GPU 가동률이 절반 수준에 머무르는 한 원인도 데이터,메모리를 기다리는 대기 시간이다.
- 두 번째 한계. 수평 공유 불가. Agent AI는 Prefill → Decode 분리, 멀티 에이전트가 같은 컨텍스트(KV)를 공유해야 하지만, 노드에 갇힌 메모리로는 이를 못 한다. 결국, 병목은 '계층을 더 깊게(수직)'가 아니라, '노드'를 가로질러 메모리를 나눠 쓰는 연결성 (수평)'으로 옮겨간다.

수직 오프로딩의 한계와 수평 메모리 풀링



고립 메모리를 활용하는 '수평' 메모리 공유

구분	Node A	Node B	클러스터 전체
필요 메모리	80	20	100
보유 메모리	50	50	100
수급 갭 (보유-필요)	-30	+30	0
기존 구조 (노드 고립)	30 부족	30 유휴	총량 충분하나 실패
메모리 풀링	부족분 30 흡수	유휴분 30 제공	가동률·활용률 개선

메모리 풀링 – 호스트에서 메모리를 떼어내 공유한다.

CXL이 메모리를 'GPU의 부속'에서 '공용 자원'으로 바꾼다

- 해법은 메모리를 호스트에서 분리(disaggregation)해 공용 풀(pool)로 만들고, 여러 GPU,노드가 패브릭 너머로 함께 접근하게 하는 것이다. 이를 가능케 하는 표준이 CXL이다. CXL은 메모리를 'GPU에 붙은 부속'에서 '필요한 만큼 빌려 쓰는 공용 자원'으로 바꾼다.
- 풀링은 앞의 두 한계를 동시에 푼다. 1. 고립 메모리 해소 → 메모리 가동률을 높여, TCO를 낮춘다. 2. KV 캐시를 노드간 공유함 → Prefill 노드가 만든 KV를 Decode 노드가, 한 에이전트의 컨텍스트를 다른 에이전트가 패브릭으로 끌어 쓴다. 오프로딩으로 '내려놓은' KV를 여러 주체가 함께 재사용하는 것이 가능해진다.
- CXL 채택은 이제 실증 단계를 지나 상용화 초입에 진입하고 있다. Astera Labs의 Leo CXL 컨트롤러는 Microsoft Azure M-Series VM에서 CXL-Attached memory 확장 사례로 적용되고 있으며, 2026년 말 일반 상용화가 예정되어 있다. 더 중요한 변화는 AI 추론 쪽이다. Astera는 KV Cache 오프로딩을 겨냥한 커스텀 CXL 설계에서 두 번째 디자인 원을 확보했고, 2027년부터 관련 매출이 발생할 것으로 기대된다. 이는 Agent AI가 만든 메모리 초과수요가 HBM뿐 아니라, CXL 기반 메모리 풀링 시장으로도 확산되고 있다는 초기 신호다.

CXL 상용화 및 양산 시점 계획

레이어/시기	① 2024-25·도입·검증	② 2026·양산·개화★	③ 2027-28·추론·메인스트림
CXL 표준	CXL 2.0	CXL 3.0/3.1	CXL 3.2 → 4.0
메모리 모듈	삼성 CMM-D 2.0 SK 96-128GB 검증	삼성 CMM-D 3.1 SK 256GB (3.2)	삼성 Pangea v3 (광연계 풀링)
컨트롤러·스위치	Montage MXC 양산	Astera Leo GA Marvell Structera S30260 샘플 (3Q)	Astera KV 캐시 커스텀 CXL 매출
호스트 CPU	Intel Xeon· Granite Rapids (CXL 2.0)	Intel 신규 CPU (CXL 3.0) NVIDIA Vera	CXL 3x 본격 지원
수요 채택	범용 컴퓨터 검증	개화 원년 (범용 컴퓨터)	AI 추론·KV 캐시 멀티 에이전트 풀링

자료: 산업 및 언론 보도 종합, 대신증권 Research Center

CXL 관련 글로벌 및 국내 관련 기업 밸류체인

역할	글로벌 업체	국내 업체
CXL 메모리 모듈 (CXL DRAM 미디어)	Micron (CZ120)	삼성전자 (CMM-D, 업계 첫 CXL 3.1 양산), SK하이닉스 (CMM-DDR5, 256GB CXL 3.2)
CXL 컨트롤러 (MXC/스마트)	Astera (Leo), Marvell (Structera X/A), Montage (MXC-삼성 SK 채택), Rambus	파두 (고성능 CXL 컨트롤러 개발), 파네시아 (비상장, CXL 3.0 IP)
CXL 스위치·패브릭 (풀링의 심장)	Marvell (Structera S30260, XConn 인수), Broadcom, Microchip	파네시아 (스위치 IP) — 상장 순수주 부재
인터페이스 IP·리타이머·SerDes (CXL-attach 공동)	Credo (리타이머), Astera (Aries), Marvell (Alaska P)	켈리타스반도체 (초고속 I/F IP), 오픈엣지 (메모리 컨트롤러 IP)
후공장·테스트·번인·소켓	Advantest, Teradyne (범용 ATE)	네오셈 (세계 첫 CXL 램 검사 장비), 엑시콘 (삼성과 CXL 2.0 테스터 공동 개발), 큐알티, 오킨스전자 (소켓)
모듈 기반 (PCB)	—	티엘비 (차세대 메모리 모듈 PCB), 코리아퍼킷 (대덕전자)
풀링 시스템 SW·어플라이언스	Penguin Solutions, MemVerge, Liquid, SMART Modular	—

자료: 산업 및 언론 보도 종합, 대신증권 Research Center

메모리 풀링이 만드는 다음 병목 – 인터커넥트 월 (Interconnect Wall)

메모리 월은 사라지는 것이 아니라, 패브릭으로 이동한다

- 메모리 풀링은 HBM, DRAM, CXL 메모리를 개별 GPU나 호스트에 종속시키지 않고, 패브릭 너머의 공용 자원으로 바꾸는 구조다. 이를 통해 고립 메모리 (stranded memory)를 줄이고, 전체 메모리 가동률을 높일 수 있다. 하지만 이 효율은 공짜가 아니다. 로컬 메모리에서 끝나던 접근이 이제 패브릭을 건너야 하므로, 풀링으로 얻는 용량과 가동률은 대역폭과 레이턴시라는 인터커넥트 비용으로 치환된다.
- 즉, 메모리 월 (Memory Wall)은 사라지는 것이 아니라 인터커넥트 월 (Interconnect Wall)로 이동한다. KV Cache를 여러 노드가 공유하고, prefill 노드가 만든 캐시를 Decode 노드가 가져다 쓰려면 패브릭은 더 빠르고, 더 넓고, 더 낮은 지연을 제공해야 한다. Astera Labs의 Scorpio X 320-lane 고-radix 스위치, Hypercast, in-Network Compute 같은 스케일업 AI 패브릭 기술은 이 병목을 줄이기 위한 초기 해법으로 볼 수 있다.
- 결국 KV Cache가 만든 메모리 초과수요는 HBM에서 멈추지 않는다. 캐시를 밖으로 밀어낼수록 데이터는 더 멀리, 더 자주, 더 빠르게 이동해야 한다. 그래서 Agent AI의 병목은 메모리에서 패브릭으로, 다시 구리에서 광으로 이동한다. KV Cache 풀링은 메모리 수요이자, 동시에 광 인터커넥트 (CPO)의 수요로 이어진다.

연산 → 메모리 → 인터커넥트 → 광으로 이동해나갈 병목



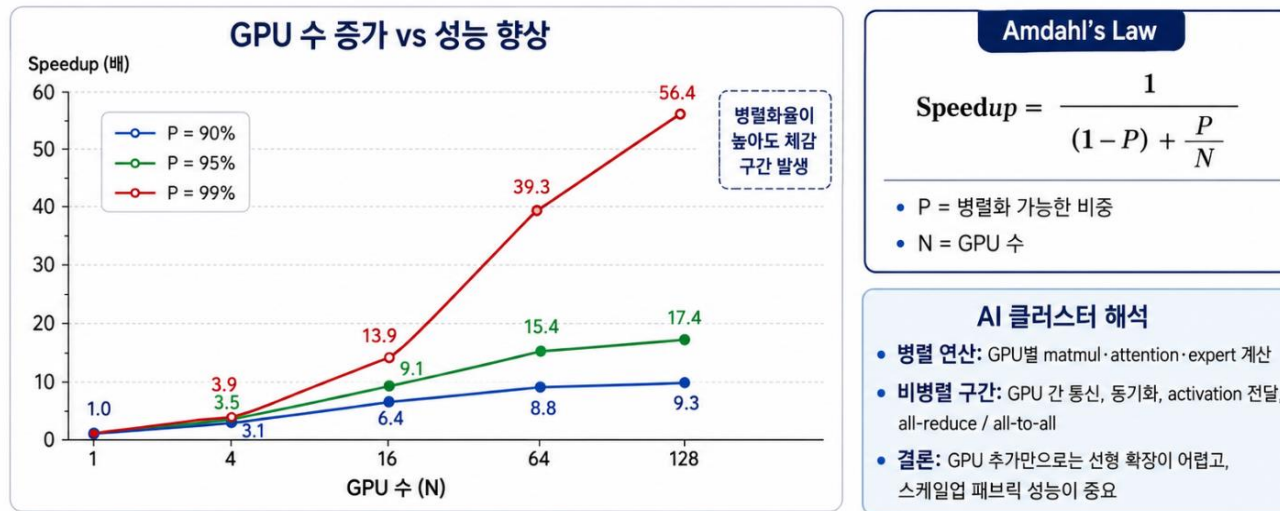
III . 광전화의 이유 – Beachfront

메모리를 넘어, 이제는 'GPU간 연결'이 병목

데이터를 어디에 둘 것인가에서, 얼마나 빨리 주고받을 것인가로

- 지금까지 AI 인프라의 핵심 과제는 '단일 GPU에 메모리를 얼마나 잘 공급하느냐(HBM 대역폭)'였다. 하지만 모델이 한 장의 GPU에 들어가지 않고, 저지연을 요구하는 에이전트 및 피지컬 AI가 등장하면서 연산의 패러다임이 바뀌고 있다. 여러 GPU가 중간 결과를 끊임없이 주고받으며 하나의 계산을 완성하는 구조로 이동한 것이다.
- 문제는 여기서 '병렬 처리의 한계'가 드러난다는 점이다. 독립적인 연산(배치 처리)은 GPU를 추가한 만큼 선형적으로 빨라지지만, GPU 간 통신이 필수적인 의존적 연산은 '암달의 법칙(Amdahl's Law)'에 따라 '인터커넥트 대역폭'이라는 속도의 천장에 부딪힌다. GPU 100개를 쏜다고 100배 빨라지지 않는 이유다.
- 특히 모델을 쪼개어 연산할 때 발생하는 All-Reduce(중간 결과 교환)로 통신 비용은 모델 크기와 GPU 대수가 커질수록 기하급수적으로 폭발한다. 거대 모델 환경에서는 전체 처리 시간의 80% 이상을 연산이 아닌 통신에 허비할 정도다. 결국 앞으로의 확장은 단순한 연산력의 추가가 아니라, GPU의 촘촘한 통신 병목을 어떻게 뚫어내느냐의 싸움이다.

암달의 법칙: GPU 통신 동기화 구간이 전체 확장의 상한을 결정한다.



연산 병렬화의 네 가지 방식

같은 GPU 클러스터도, 무엇을 쪼개느냐에 따라 병목이 달라짐

- AI 모델을 여러 GPU에 나눠 올리는 방식은 크게 네 가지다. 중요한 차이는 무엇을 쪼개느냐와 GPU끼리 언제 데이터를 주고받느냐다. Data Parallelism과 Pipeline Parallelism은 비교적 독립적으로 계산하거나 단계별로 결과를 넘기는 구조라 통신 부담이 완만하다. 반면, Tensor Parallelism과 Expert Parallelism은 여러 GPU가 하나의 계산을 함께 완성해야 하기 때문에, 통신이 연산의 핵심 경로에 들어오며 지연과 대역폭에 훨씬 민감해진다.

GPU의 병렬화 방식에 따라서 인터넥트의 부하는 달라짐.

통신부담	병렬화방식	무엇을쪼개나	GPU간통신시점·패턴	통신연산 (Collective)	지연·대역폭민감도	적합네트워크
완만	Data Parallelism (DP) 데이터 병렬	데이터(배치)를 분할 (모델은 GPU마다복제)	스텝종료시 1회 (연산후 일괄동기화)	All-Reduce (gradient)	낮음 (연산과 겹쳐 은닉 가능)	스케일아웃
	Pipeline Parallelism (PP) 파이프라인 병렬	레이어를 단계(stage)로 분할	단계 경계에서전달	P2P (activation send/rcv)	낮음~중간 (단, 파이프라인 버블)	스케일아웃
높음 (임계경로)	Tensor Parallelism (TP) 텐서 병렬	레이어 내부 행렬(가중치)을 분할	매 레이어마다합산 (연산 도중 임계경로)	All-Reduce / All-Gather (레이어별반복)	매우 높음	스케일업
	Expert Parallelism (EP) 전문가 병렬	MoE 전문가(expert)를 GPU별분산	토큰분배·취합시 (라우팅)	All-to-All (dispatch / combine)	매우 높음	스케일업

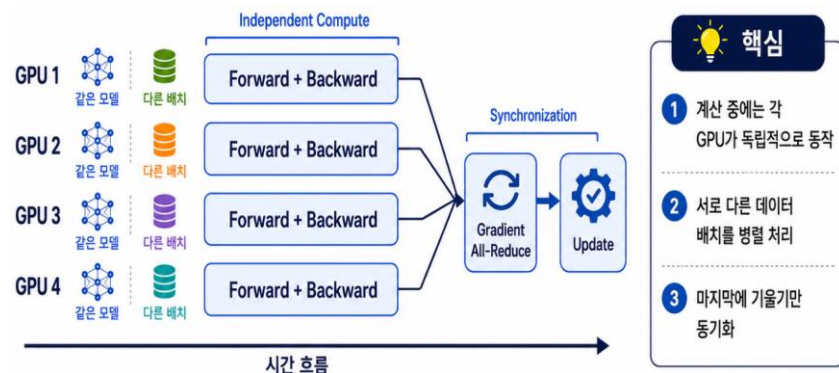
주: '통신 연산'은 GPU 간 데이터를 합치는 집합통신(collective) 유형. TP·EP는 통신이 연산의 임계경로(critical path)에 위치해 지연·대역폭에 민감 → 스케일업 도메인 필수.
 자료: 대신증권 Research Center

DP, PP: 통신 부담이 가벼운 병렬화

따로 계산하거나, 단계별로 넘기는 구조

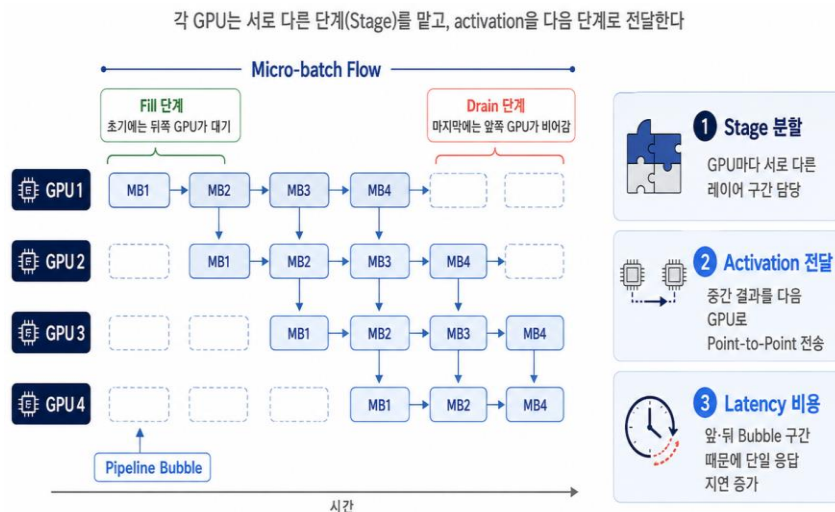
- Data Parallelism은 각 GPU가 같은 모델을 들고 서로 다른 데이터를 따로 계산한다. 계산이 끝나면 각 GPU가 구한 gradient(모델을 어느 방향으로 고쳐야 하는지 알려주는 수정 신호)를 서로 합쳐 평균을 내고, 모든 GPU의 모델 파라미터를 같은 방향으로 업데이트한다. 이 동기화 과정이 Gradient All-Reduce다. 즉 GPU들이 매 연산마다 대화하는 것이 아니라, 한 번의 계산 단위가 끝난 뒤 '이번에 모델을 어떻게 고칠지'만 맞추는 구조다.
- 계산 중에는 각 GPU가 독립적으로 움직이고, 마지막에 gradient를 모아 평균을 낸 뒤 동일한 모델 업데이트를 적용한다. 그래서 통신이 연산 구간을 막지 않고, 지연 민감도가 낮다.
- Pipeline Parallelism (PP)은 모델의 레이어를 여러 단계로 나누고 GPU 그룹에 순차적으로 배치하는 방식이다. 각 단계는 계산 결과인 Activation(중간 결과값)을 다음 단계로 전달하는 Point-to-Point Communication만 수행한다. 통신량은 크지 않지만, 파이프라인이 채워지고 비워지는 과정에서 병목이 발생하여(Pipeline Bubble) 단일 응답 지연을 증가시킬 수 있다.

Data Parallelism: 계산은 따로, 업데이트만 함께



자료: 산업자료, GPT 재구성, 대신증권 Research Center

Pipeline Parallelism: 레이어를 나눠 순차적으로 이어서 계산



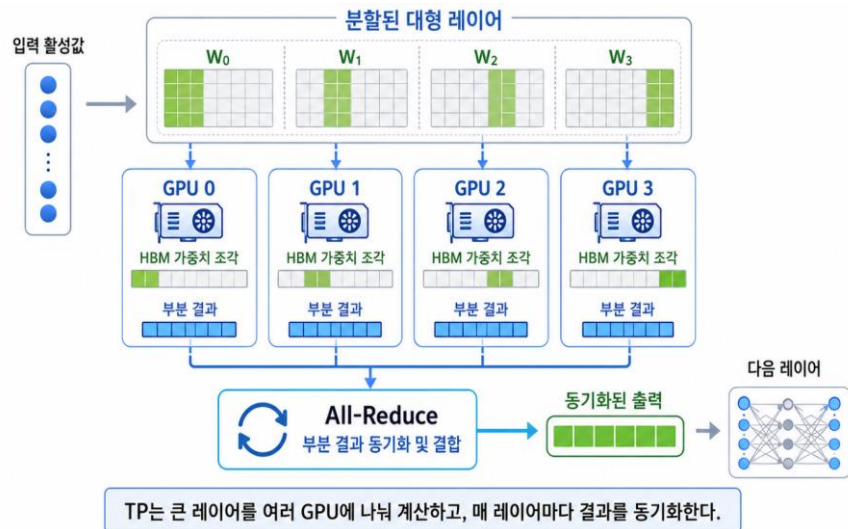
자료: 산업자료, GPT 재구성, 대신증권 Research Center

TP, EP: 인터커넥트가 병목이 되는 병렬화

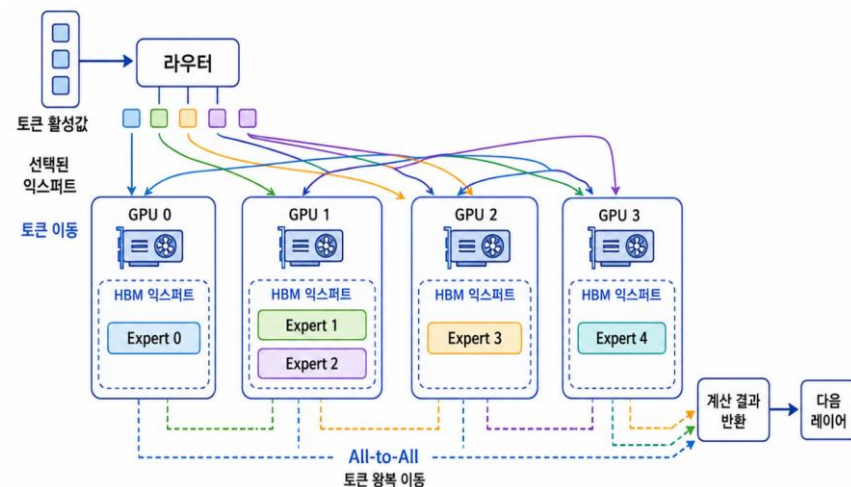
여러 GPU가 하나의 계산을 함께 완성하는 구조

- TP(Tensor Parallelism)는 하나의 레이어 안에 있는 거대한 행렬 계산을 여러 GPU가 나눠 수행하는 방식이다. 각 GPU는 자기 HBM에 올라온 가중치 조각만 읽어 부분 계산을 하지만, 그 결과는 완성된 값이 아니라 부분 답이다. 따라서 다음 연산으로 넘어가기 전에 GPU들이 부분 결과를 합쳐야 하며, 이 과정이 all-reduce다. TP는 HBM 용량과 대역폭, GPU 연산을 동시에 병렬화하지만, 레이어마다 동기화가 필요해 초저지연 인터커넥트를 요구한다.
- EP(Expert Parallelism)는 MoE 모델에서 여러 Expert FFN을 GPU별 HBM에 나눠 배치하는 방식이다. Expert는 실행 중인 객체가 아니라, GPU 및 HBM에 상주하는 FFN 가중치 덩어리다. 라우터가 토큰마다 필요한 Expert를 고르면, 토큰 Activation(토큰 활성화 코드)이 해당 Expert가 있는 GPU로 이동하고, 계산 후 다시 원래 위치로 돌아온다. 이 과정에서 모든 GPU가 서로 토큰을 주고 받는 'All-to-All' 통신이 발생한다. MoE는 연산량을 줄이지만, 그 대가로 스케일업 대역폭 수요를 키운다.

Tensor Parallelism: 하나의 큰 레이어를 여러 개의 GPU에 분할함



Expert Parallelism: MoE에서 토큰이 전문가가 있는 GPU로 이동함

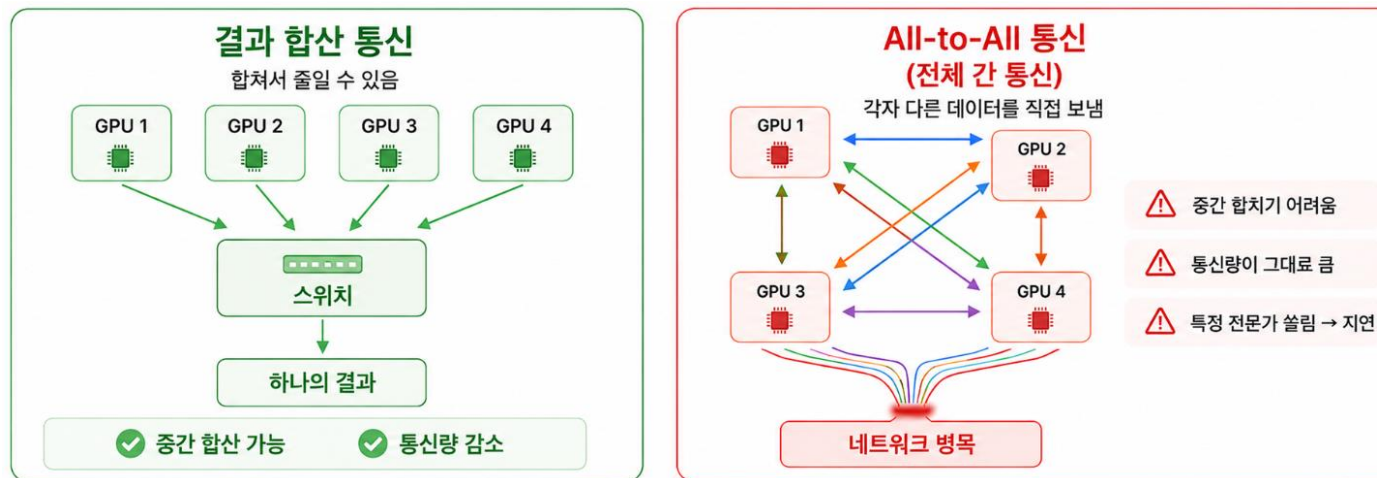


EP(Expert Parallelism, MoE) – 줄일 수 없는 'All to All'이 스케일업을 강제한다.

프런티어가 MoE로 수렴 → 'All-to-All'은 우회로가 없어 대역폭 증설이 비가역적

- All to All 통신은 집합통신 중에서도 줄이기 가장 어려운 통신이다. 결과값을 합산하는 통신은 여러 GPU의 값을 하나로 더하는 구조라, 스위치가 중간에서 일부 계산을 대신해 데이터 이동량을 줄일 수 있다. 반면, All to All 통신은 각 GPU가 서로 다른 데이터를 모든 GPU에 개별적으로 보내야 하므로, 중간에서 합치거나 압축할 여지가 작다. 결국 병목은 네트워크가 한 번에 처리할 수 있는 전체 통신 용량에 그대로 묶인다. (Bind) MoE모델에서는 상위 전문가 선택 방식으로 활성값이 여러 전문가로 흩어지고, 특정 전문가에게 부하가 몰리면 가장 늦은 작업이 전체 지연을 결정하게 된다. 여유 용량을 두는 방식으로 완화는 가능하지만 구조적 병목은 남는다.
- 그래서 MoE에서는 규모 확장이 선택이 아닌, 강제가 된다. All to All 통신은 낮은 지연시간과 높은 대역폭을 동시에 요구하기 때문에, 가능한 한 가장 빠른 GPU 연결망 안에서 처리해야 한다. 대형 고속 연결망의 핵심 의미도 여기에 있다. 줄일 수 있는 통신은 외부 확장으로 보내고, 줄이기 어려운 MoE 병렬화와 텐서 병렬화 통신은 내부 고속 연결망 안에 남는다. 인공지능 클러스터가 커질수록 내부 고속 연결망의 부하는 더 커질 수 밖에 없다.

All to all 통신은 GPU간의 통신 대역폭이 곧 성능

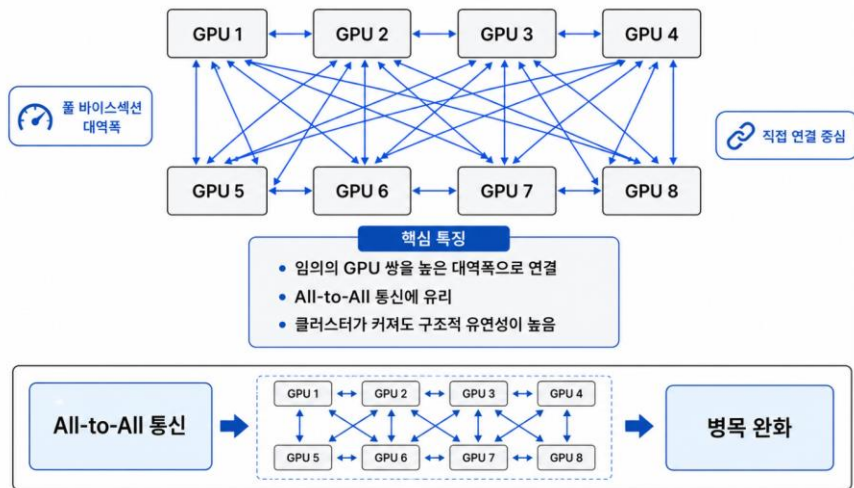


구글의 TPU – 토폴로지는 같려도, 인터커넥트는 동일하게 커진다.

'All-to-All'에선 NVIDIA가 구조적 우위. 그러나 스케일업 광 수요는 ASIC, GPU 양쪽 모두에서 증가

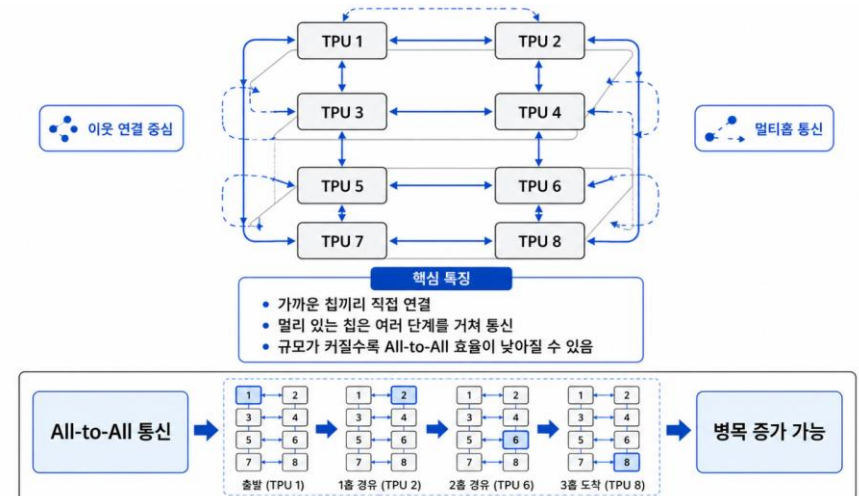
- 구글의 TPU와 NVIDIA는 토폴로지(GPU를 구성하는 방식)에서 차이가 있다. All-to-All 연결에서는 NVIDIA가 유리하다. NVIDIA의 NVLink·NVSwitch 기반 구조는 여러 GPU를 넓은 대역폭으로 직접 연결해 All-to-All에 강하다. 반 면 TPU의 3D 토러스 구조는 가까운 칩끼리는 빠르지만, 멀리 있는 칩과 통신할 때는 여러 단계를 거쳐야 한다. 클러스터가 커질수록 칩당 체감 대역폭이 떨어질 수밖에 없다. 그래서 MoE, 멀티모달, 롱컨텍스트, 피지컬 AI처럼 칩 간 통신이 많은 모델에서는 NVIDIA의 구조적 우위가 커진다.
- 다만 NVIDIA도 무한히 유리한 것은 아니다. NVLink로 묶인 고속 도메인을 벗어나면 NVIDIA 역시 계층형 네트워크에 의존해야 한다. 결국 양쪽 모두의 과제는 같다. 더 큰 AI 클러스터를 더 빠르게, 더 낮은 전력으로 연결해야 한다.
- 결론적으로 All-to-All이 강해질수록 NVIDIA가 유리하다. 하지만 TPU든 GPU든 AI 모델이 커질수록 칩 간 데이터 이동은 늘어난다. 따라서 토폴로지의 승자와 무관하게 인터커넥트 수요가 구조적으로 커진다. 구리선은 속도·거리·전력 한계에 부딪히고 있고, 이를 넘기 위해 양 진영 모두 광 인터커넥트로 이동하고 있다. NVIDIA는 CPO·광링크 방향으로, Google TPU는 OCS 같은 광 스위칭 방식으로 확장한다.

NVIDIA 토폴로지: 크로스바 기반 연결. All to All에 유리



자료: 산업자료, GPT 재구성, 대신증권 Research Center

TPU 토폴로지: 3D 토러스 기반 연결. All to All에 불리



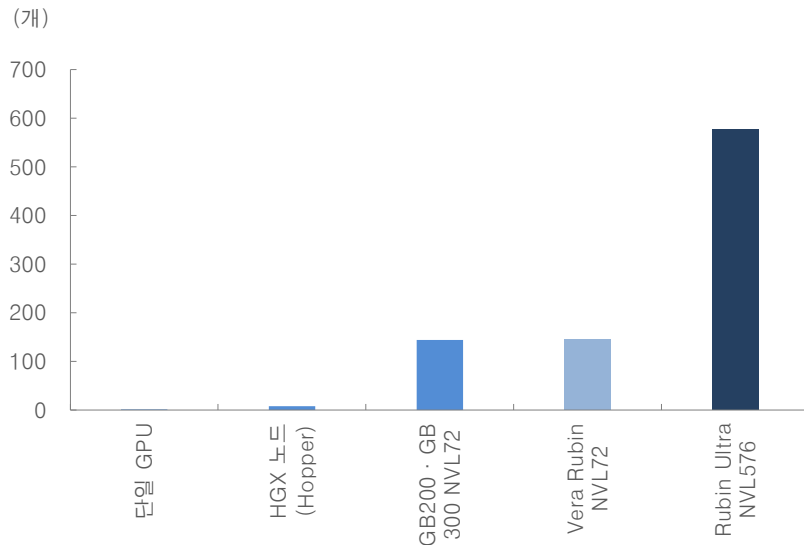
자료: 산업자료, GPT 재구성, 대신증권 Research Center

스케일업 vs 스케일아웃 — '큰 GPU 만들기' vs '큰 GPU 잇기'

'이어붙여 계산'을 담는 그릇이 곧 스케일업이다

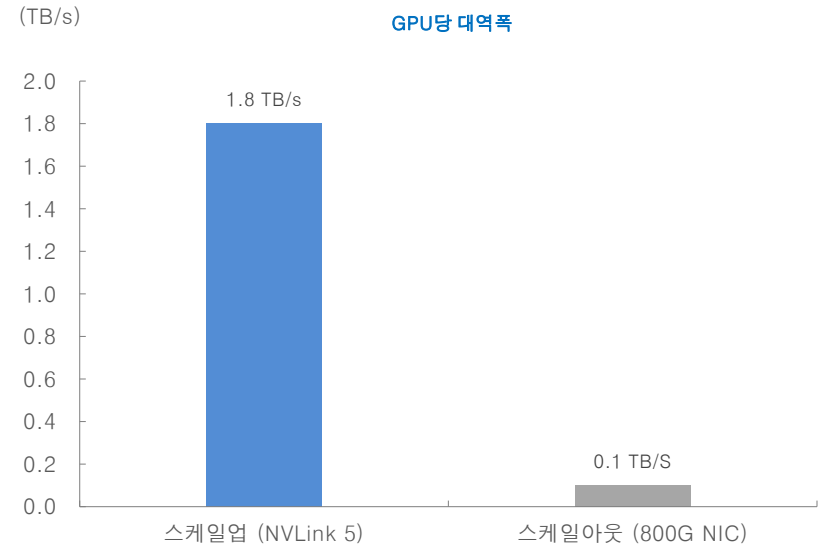
- 스케일업(Scale-up)은 여러 GPU-ASIC을 묶어 '하나의 더 큰 GPU'처럼 동작하게 만드는 연결이다. NVLink가 대표적이며, 초고대역폭, 초저지연, 높은 신호속도가 동시에 필요한 TP, EP가 스케일업과 성능이 연동된다. 연산 단위가 1 GPU → 8 GPU → 72 GPU → 576 GPU로 커지는 이유도 여기에 있다.
- 스케일아웃(Scale-out)은 이렇게 만들어진 '큰 GPU' 여러 개를 다시 연결해 더 큰 미션을 수행하는 망이다. InfiniBand, Ethernet이 대표적이며, 노드를 가로지르는만큼 상대적으로 느슨한 메시지 패싱 기반이다. 따라서 지연 내성이 큰 DP, PP가 스케일아웃과 성능이 연동된다.
- 둘 간의 처리 속도 차이는 숫자로도 분명히 드러난다. NVLink 72 단일 랙 안에서만 약 130TB/s를 처리한다. GPU당 스케일업 대역폭도 NVLink 5세대 기준 약 1.8TB/s로, 800Gb/s 스케일아웃 NIC의 약 18배다. Marvell은 AI 트래픽의 80 ~ 90%가 스케일업 구간에서 발생한다고 추정한다. 결국 AI 하드웨어 네트워킹 투자의 무게중심이 스케일업으로 이동하는 이유는, AI 연산이 점점 따로 계산이 아니라 이어붙여 계산하는 구조로 바뀌고 있기 때문이다.

엔비디아 아키텍처별 GPU die 수 추이 - Rubin Ultra에서 4배 증가



자료: NVIDIA, 대신증권 Research Center

GPU당 대역폭: 스케일업이 스케일아웃의 약 18배



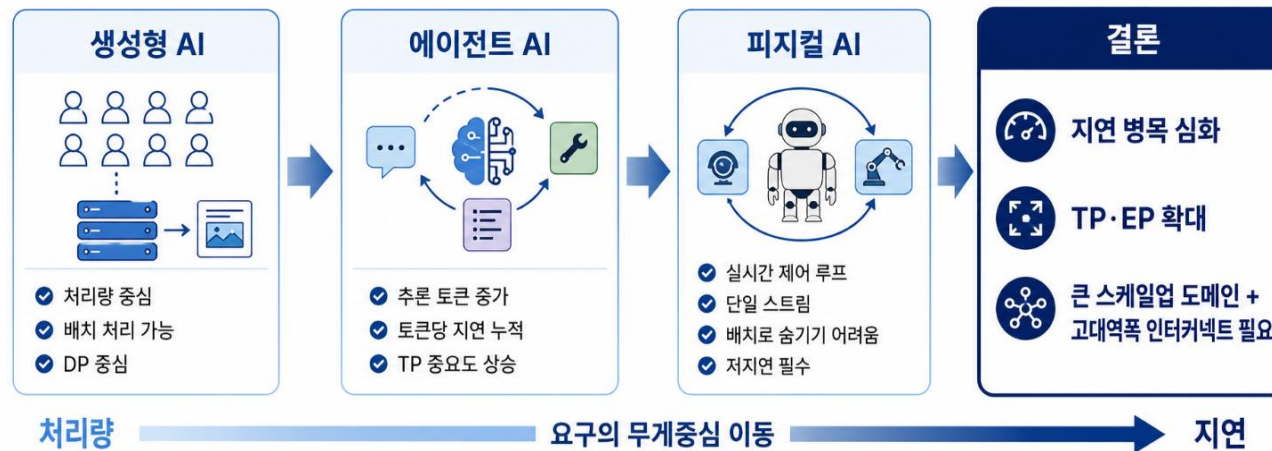
자료: NVIDIA, 대신증권 Research Center

피지컬 AI가 스케일업의 무게를 키운다

실시간·단일 스트림·저지연 – 스케일업 인터커넥트의 수요를 키우는 요인

- 생성형 AI에서 에이전트, 피지컬 AI로 갈수록 컴퓨팅 요구의 무게중심은 '처리량'에서 '지연'으로 이동한다. 로봇, 자율 시스템, 월드 모델(VLA)은 실시간 제어 루프 안에서 판단과 행동을 반복해야 하므로, 응답이 몇 초씩 늦어지면 효용이 급격히 떨어진다. 또한 동시 사용자를 묶어 처리하는 '배치'방식이 잘 사용되는 생성형 AI와 달리, 피지컬 AI는 개별 객체, 환경 상황에 즉각 반응해야 하는 단일 스트림 성격이 강해 지연을 배치 뒤에 숨기기 어렵다.
- 이때 병렬화 전략도 달라진다. 데이터 병렬화(DP)는 처리량은 늘지만 단일 응답 지연을 줄이지 못하고, 파이프라인 병렬화(PP)는 단계 간 대기과 버블로 오히려 지연을 키울 수 있다. 결국 단일 스트림의 토큰당 지연을 낮추는 핵심 경로는 텐서 병렬화(TP)를 확대하는 것이다. TP가 커질수록 GPU/ASIC 간 중간 결과를 더 빠르게 주고 받아야 하므로, 더 큰 스케일업 도메인과 더 높은 인터커넥트 대역폭이 필요해진다.
- 에이전트 추론(reasoning)은 이 압력을 더 키운다. 하나의 작업을 끝내기까지 중간 사고, 물 호출, 검증 과정을 반복하면서 생성 토큰 수가 크게 늘어나기 때문이다. 이때 토큰당 지연은 작업 전체 지연으로 누적된다. 따라서 에이전트와 피지컬 AI의 확산은 더 강한 TP, EP 더 큰 스케일업 도메인, 더 높은 인터커넥트 대역폭 수요로 이어진다.

에이전트 AI, 피지컬 AI로 갈수록 AI 워크로드의 무게 중심은 처리량에서 지연으로 이동한다

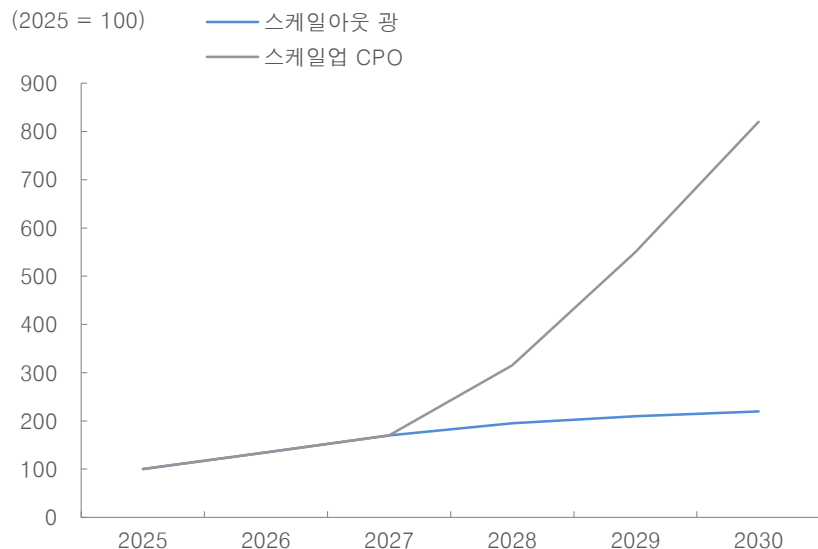


결국 스케일업도 '구리에서 광으로' — 인터커넥트의 종착점

이러나저러나 광으로: 메모리에서 출발해도, 연산에서 출발해도 같은 결론

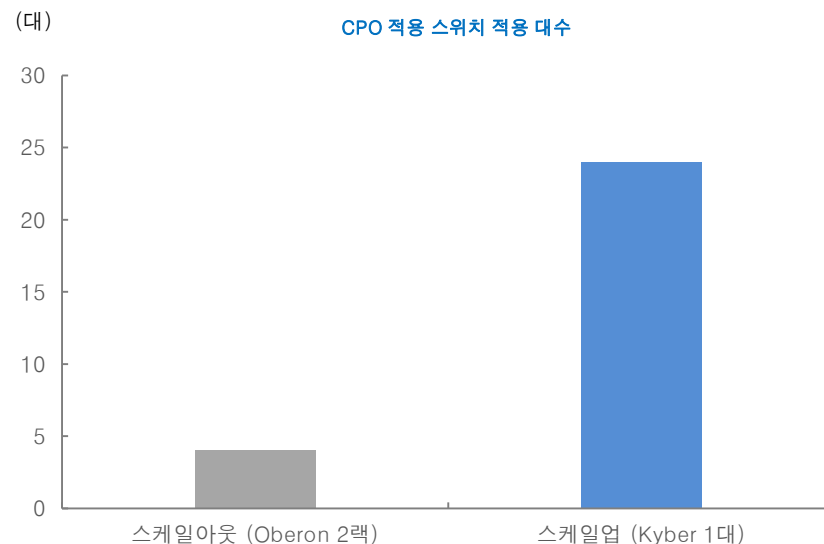
- 스케일업 시장에서 가속기간의 연결은 현재 구리가 사용되고 있다. 구리는 싸고, 전력을 적게 사용하며, 신뢰성이 높기 때문이다. NVIDIA가 GPU를 한 랙 안에 최대한 밀집시키고 냉각까지 강화하는 이유도 결국 '랙' 안이면 구리로 연결할 수 있기 때문이다. 하지만, 레인당 속도가 매년 2배씩 증가하면서, 결국 구리는 Beachfront 문제로 인해서 한계에 봉착할 가능성이 높다. 28년 파인만아키텍처에서부터 NVIDIA도 CPO 즉, 광 연결을 스케일업에 적용할 가능성이 높다.
- 중요한 것은 시장의 크기다. NVIDIA Oberon 랙 2대 기준으로 CPO가 적용되는 스위치가 4대(스케일 아웃용)라고 한다면, Kyber 1대 (Oberon 랙 2대와 연산기준 동일) 부터 적용 가능한 스위치 대수는 24대로 CPO가 적용 가능한 시장 규모 TAM은 기존 스케일 아웃 시장 대비 6배에 이른다. 스케일업에서 CPO가 적용되기 시작하면, 광 시장의 TAM은 비선형적으로 커질 수 있다.
- 투자의 관점에서 방향성은 변함없다. 메모리 병목에서 출발해도, 연산 병렬화에서 출발해도 결론은 같다. 병목은 결국 인터커넥트로 이동하고, 그 해법은 광 전환이다. 그래서 광 부품 CPO 스케일업 스위치 패키징 매력적인 시장이며, AI 하드웨어 투자자들에게 있어서 한 시장이다.

스케일업 광 침투 → 광 TAM 비선형적 확대



자료: 대신증권 Research Center 추정
주: 광 스케일업 2028년 개시 가정. NVIDIA Oberon/Kyber 아키텍처 기준

동일 연산 기준 CPO 적용 가능 스위치 수: 스케일아웃 vs 스케일업



자료: 대신증권 Research Center 추정
주: 광 스케일업 2028년 개시 가정. 1,2계층 NVSwitch에 CPO 적용 가정
주: NVIDIA Oberon/Kyber 아키텍처 기준

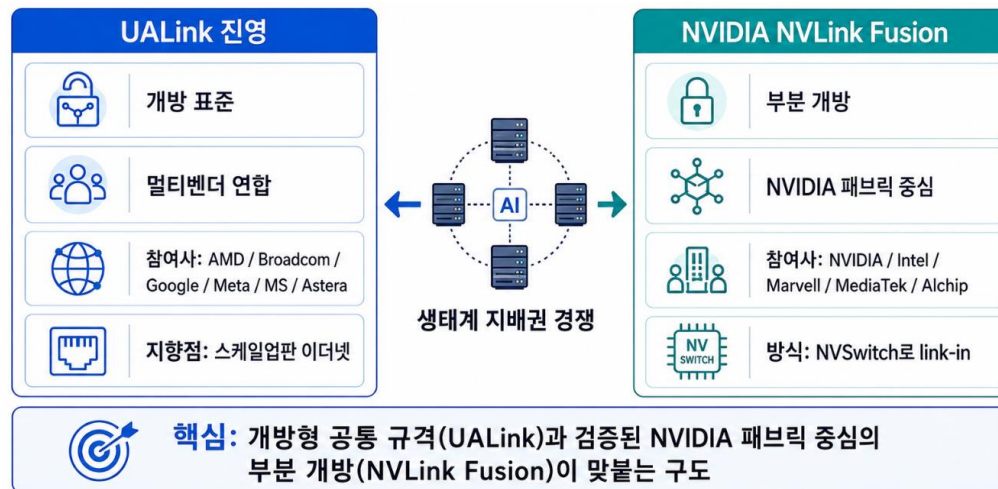
Ⅳ. 표준은 미정. 광 전환은 필연

프로토콜 경쟁 NVIDIA VS 비NVIDIA

엔비디아는 NVFusion

- 스케일업 인터커넥트는 이제 표준이 정해지기 시작하는 시장이다. 경쟁 구도는 크게 엔비디아의 NVLink-NV Fusion과 비엔비디아 진영의 UALink로 갈린다. 이는 GPU-XPU를 어떤 프로토콜 (통신을 하는데 사용되는 표준규격)로 묶을 것인가를 둘러싼 생태계 지배권 싸움이다.
- UALink는 '스케일업판 이더넷'을 지향한다. AMD, Broadcom, Google, Meta, Microsoft, Astera 등이 참여해 여러 벤더가 공통 규격으로 XPU를 연결 하자는 개방형 표준이다. 반면 NVLink는 엔비디아가 수년간 자체 GPU 생태계 안에서 고도화해 온 폐쇄형 규격이다. NV Link 5와 NVL72 양산을 통해 성숙도에서는 UALink 진영보다 앞서 있다.
- 엔비디아는 NVLink - Fusion을 통해 외부 CPU, ASIC이 NVSwitch 패브릭에 들어올 수 있도록 인터페이스를 여는 전략을 택했다. 이는 하이퍼스케일러가 자체 ASIC을 쓰더라도 그 칩들을 엔비디아 패브릭 위에 얹게 만들어 플랫폼 통제권을 유지하려고 하는 것이다. UALink가 단일 벤더 종속 회피를 내세운다면, 엔비디아는 검증된 성능과 부분 개방으로 그 명분을 약화시킨다. 요컨대 규격의 문은 열리, 패브릭의 중심은 여전히 엔비디아가 쥐겠다는 전략이다.

스케일업 인터커넥트 표준 경쟁 – NVIDIA NVLink Fusion V.S. UALink

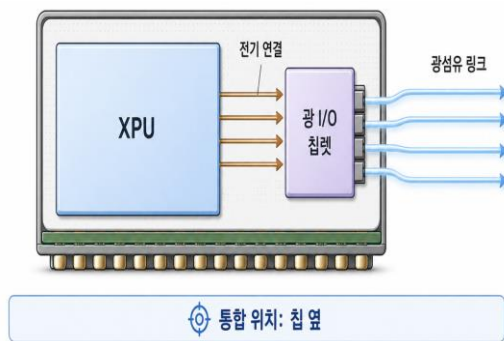


개막하는 광의 시대와 춘추전국시대

광 I/O 시대에 대표적인 세 가지 아키텍처

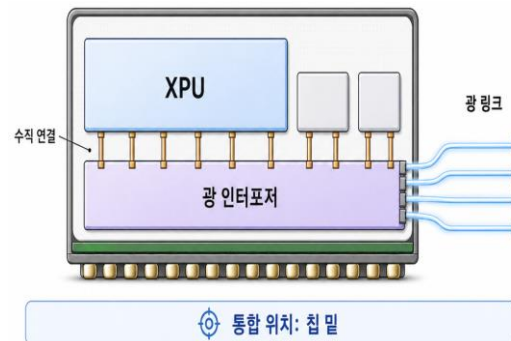
- 앞서 언급한 경쟁이 통신 프로토콜 차원의 경쟁이었다면, 물리계층에서 스케일업 인터커넥트를 어떤 구조로 구현할 것인가에 대한 경쟁 역시 본격화되고 있다. 현재 물리 계층 역시 표준이 정해지지 않은 상황으로, 대표적인 진영으로 Ayar Labs, Lightmatter, Celestial AI가 서로 경쟁하고 있다.
- 여기서 말하는 스케일업 광 I/O는 현재 상용화가 진행 중인 스케일 아웃 CPO와는 다르다. 스케일아웃 CPO가 주로 스위치단에서 광을 적용하는 방식이라면, 스케일업 광 I/O는 광을 AI 가속기, 즉 XPU 패키지 안으로 끌어들이는 접근이다. 다만 그 이전 단계에서는 XPU와 직접 연결되는 NVSwitch부터 광으로 전환될 가능성이 높다. 결국 대역폭 병목이 XPU - NVSwitch 구간에서 발생하기 때문이다.
- 세 아키텍처의 차이점은 결국 '광-연산의 통합 위치'다. 칩 옆에 광 I/O 칩셋을 붙이는 Ayar Labs, 칩 밑에 광 인터포저를 배치하는 Lightmatter, 그리고 패브릭 전체와 메모리 연결까지 광으로 확장하려는 Celestial AI로 구분할 수 있다.
- 현재 시장에서는 기존 CPO 전환 경로와 상대적으로 가까운 Ayar Labs 방식이 차기 표준 기술이 될 가능성이 가장 높다고 예상된다. 다만 구조적 문제를 어디까지 해결하느냐의 관점에서는 이야기가 달라진다. 칩 가장자리의 I/O 공간이 병목이 되는 Beachfront 문제는 칩 밑으로 광 인터포저를 넣는 Lightmatter 방식이 더 직접적인 해법에 가깝고, HBM 용량·대역폭 한계로 대표되는 메모리 월 문제는 메모리 연결까지 광 패브릭으로 확장하는 Celestial AI 방식이 가장 급진적인 해결책에 가깝다.

Ayar Lab: 광 칩셋을 XPU 옆에 배치



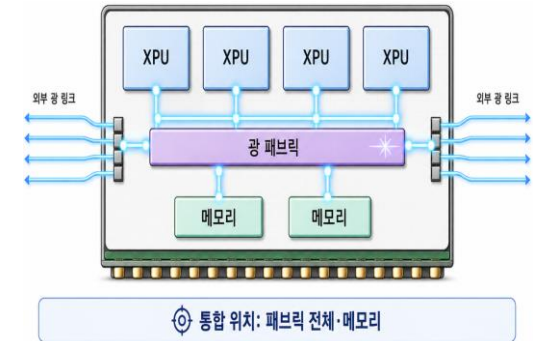
자료: 산업자료, GPT 재구성, 대신증권 Research Center

Lightmatter: 칩 밑 광 인터포저



자료: 산업자료, GPT 재구성, 대신증권 Research Center

Celestial AI: 패브릭 전체를 광으로 연결



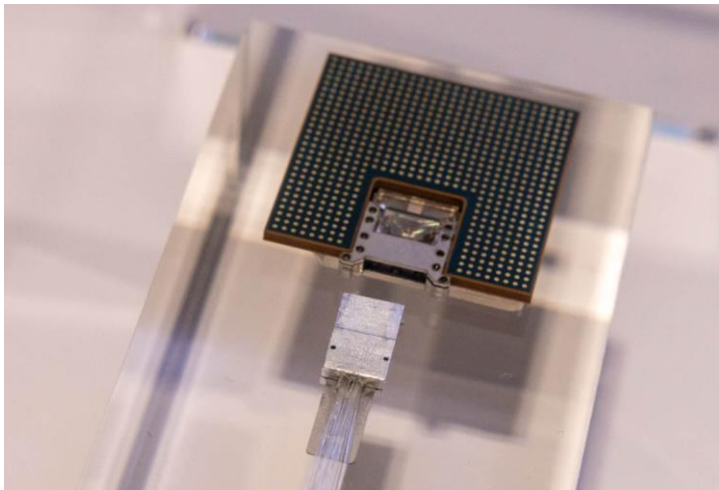
자료: 산업자료, GPT 재구성, 대신증권 Research Center

Ayar Labs – 칩 옆에 붙이는 광 I/O 칩렛

표준으로 채택될 가능성이 가장 높은 선두주자

- Ayar Labs의 접근은 TeraPHY(물리적으로 붙는 칩 IP 종류) 광 I/O 칩렛을 XPU 옆에 UCIE로 붙이는 방식이다. 빛은 SuperNova라는 외부 멀티파장 레이저에서 섬유로 공급한다. 기존 전기 SerDes가 놓이던 자리에 광 칩렛을 꽂는 구조에 가까워, XPU 본체 설계를 크게 바꾸지 않고도 광 I/O를 도입할 수 있다. 기존 스케일아웃 CPO와 비슷한 공정을 거쳐 가장 표준화 가능성이 높다고 여겨지고 있다.
- Ayarlabs 방식의 가장 큰 장점은 범용성이다. UCIE와 CW-WDM MSA 같은 표준 기반 구조라 특정 GPU나 ASIC에 종속되지 않는다. 머천트 GPU든 커스텀 ASIC이든 비교적 쉽게 얻을 수 있다는 점이 Ayar의 가장 큰 차별점이다. 실제로 Siver의 16채널 DFB 레이저 어레이를 SuperNova에 통합해 ECOC에서 16Tbps 양방향 전송을 실증했는데, 이는 스케일업 광 I/O에 필요한 수준에 근접한 성능이다.
- 다만, 단점은 Beachfront가 다른 방식에 부족해서 적다는점이다. TeraPHY는 '칩 위'가 아니라, '칩 옆'에 붙는 구조이기 때문에, 결국 I/O 배치가 가장자리 근처에 머문다. 면적 전체를 활용하는 인터포저, 패키지 내 광 방식보다 대역폭 밀도의 최대치는 적을 수 밖에 없다. Ayarlabs는 비상장으로 IPO를 앞두고 있다. 관련 기업으로는 Sivers(레이저 광원), Global Foundries(파운드리), Alphawave(DSP) 같은 업체들이 밸류체인으로 꼽힌다.

칩셋에 직접 연결되는 Tera PHY



자료: ServeTehHome, 대신증권 Research Center

Ayarlabs 관련 밸류체인

레이어	업체	상장	비고
광원 (레이저 다이)	Sivers★/ Lumentum★/ Coherent★/ Sumitomo/ MACOM	상장 다수	SuperNova는 Ayar 자체 모듈. Sivers가 DFB 셀 어레이 설계축, 나머지는 다이 공급. NVIDIA의 \$4B 투자로 다이 공급 병목 완화
InP 기판	(III-V 기판)	—	레이저 상위 병목. 다음 파트 공통 레이어
SiPho 파운드리	GlobalFoundries★ (45nm/Fotonix) + TSMC	상장	팹리스라 멀티 파운드리 이동 가능. Gazettabyte TSMC 경로는 Alchip GUC 경우 (COUPE/SoIC-X)
DSP/ SerDes-UCIE IP	Alphawave Semi	상장 (AWE)	Cu-/O-/Co-Wave 계열
ASIC 통합 패키징	Alchip/Global Unichip (GUC)	상장 (대만)	둘 다 TSMC 진영 ASIC-패키징. GUC 최대주주가 TSMC Gazettabyte
투자 접근	비상장 → 밸류체인	—	NVIDIA·Intel·AMD·GF·HPE·Lockheed 투자

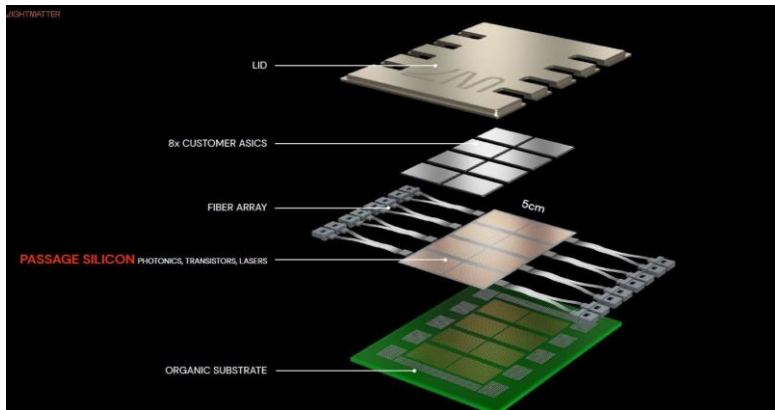
자료: 대신증권 Research Center

Lightmatter – 칩 밑에 까는 광 인터포저

Beachfront를 최대치로 올리는 인터포저 방식

- Lightmatter의 Passage M1000은 연산 칩을 위에 얹는 3D 광 인터포저다. 핵심은 I/O를 칩 가장자리에서만 빼는 것이 아니라, 인터포저 바닥면 전체에서 빛으로 뽑아낸다는 점이다. 기존 전기 I/O가 항구의 해안선처럼 가장자리 길이에 묶였다면, M1000은 항구를 해안선이 아니라 도시 전체 지하망처럼 까는 방식이다. 4,000mm²급 대형 기판 위에 여러 칩렛을 묶고, 패키지에 광섬유를 직접 붙여 대역폭을 끌어올린다.
- 칩의 가장자리를 통해서 연산을 해야하는 Ayarlabs와 비교해서 인터포저는 칩과 광 I/O가 닿는 면적이 훨씬 넓어 Beach Front 측면에서 유리하다. 단일 패키지내에서 114Tbs 양방향 대역폭, 1,024개 SerDes 레인, DWDM, 내장형 광스위칭까지 얹어 차세대 XPU, 스위치 체급을 목표로 한다.
- 다만, 제조 방식의 난이도는 단점으로 꼽힌다. 멀티테라클 웨이퍼급 인터포저라 패키징 수율 비용, 열처리가 모두 리스크다. Hot Chips 2025에서 Production-ready를 선언하고 SC25에서 실물 데모까지 보여줬지만, 아직 이름이 공개된 대형 고객사에서 선택된 바는 없다. Lightmatter의 경우, 기술 레퍼런스에서 사용 채택으로 넘어갈 수 있는 지 여부가 중요하다. Lightmatter는 비상장이고, 글로벌파운드리, Amkr, 외부 레이저 광원 업체가 관련 밸류체인으로 꼽힌다.

Lightmatter의 광 인터포저 구조의 아키텍처



자료: SemiAnalysis, 대신증권 Research Center

Lightmatter 관련 밸류체인

레이어	업체	상장	비고
광원	Guide™ (자체 광엔진, VLSI) — III-V 다이외부	자체	
SiPho 파운드리	GlobalFoundries★ (GF Fotonix)	상장	Ayar와 동일 파운드리
DSP/ SerDes-UCIe IP	AlphawaveSemi + Synopsys	상장	L200 UCIe D2D IP는 Alphawave가 로열티 프리 제공, Lightmatter@Synopsys는 224G SerDes+UCIe(3nm) 공급 BusinessWire
패키징/OSAT	Amkor★ + ASE★	상장	M1000·L200 모두 GF·ASE·Amkor 파트너십으로 양산, Data Center Dynamics 웨이퍼급 인터포저라 난도 최상
파이버 부착	256 가닥 edge-attach	—	다음 파트 공통 레이어 (FAU)
투자 접근	비상장	—	밸류 \$4.4B, 누적 조달 \$822M(25.10)

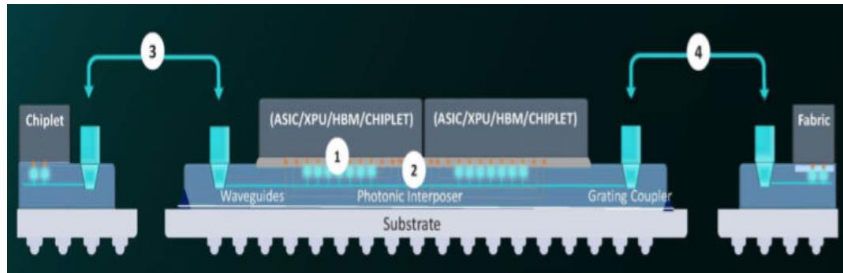
자료: 대신증권 Research Center

Celestial AI – 메모리까지 통합한 패브릭

칩에 변조기를 붙여 하는 방식

- Marvell의 Photonic Fabric은 스케일업, 인터커넥트 전용으로 설계된 광 패브릭이다. 광 칩렛과 OMIB(Optical Multi-chip Interconnect Bridge)를 통해 킬로와트급 XPU, 스위치와 직접 co-package되는 구조이며, 첫 광 칩렛은 16T급 대역폭을 목표로 한다. 이는 현재 스케일아웃 광모듈 SOTA인 1.6T 대비 10배 높은 대역폭 밀도를 겨냥한 것이다.
- 차별점은 단순히 XPU끼리를 광으로 잇는데 그치지 않는다는 점이다. Celestial AI의 방식은 XPU와 메모리를 2.5D 인터포저 수준에서 제한적을 연결하는 접근이 아니라, XPU와 메모리 양쪽에 광 변조기를 배치해 패키징, 보드, 랙을 넘어선 광 메모리 패브릭을 구성하는 방식이다. 즉, Photonic Fabric은 연산 칩 간 인터커넥트 병목뿐 아니라, KV 캐시, 오프로딩, CXL, 메모리 풀링으로 대표되는 메모리 병목까지 동시에 겨냥하는 방식이다.
- 투자 관점에서도 Marvell은 가장 급진적인 진영이다. 26년 2월 Celestial AI 인수에 더해 CXL 스위치 업체 Xconn까지 화북하면서, 광 인터커넥트와 메모리 디스에그리게이션을 함께 가져가려 하고 있다. 현재 Marvell의 주력 매출은 DSP와 커스텀 ASIC 이지만, 장기적으로는 Photonic Fabric을 통해 'AI 인프라의 메모리 인터커넥트 병목을 푸는 회사'로 재평가될 가능성이 있다.

Celestial AI – Photonic



자료: Celestial AI, 대신증권 Research Center

Celestial AI 관련 밸류체인

레이어	업체	상장	비고
소유	Marvell★ (100% 내재화)	상장 (MRVL)	26.2월 인수완료, 선지급 \$3.25B-총 \$5.5B Optics.org
초기공정·PDK	imec (IC-Link)	비상장	imec PDK로 SiPho 공정·패키징 개발, 상용 파운드리로 볼륨 이관 Imec
SiPho 파운드리	TSMC (imec IC-Link 경유)	상장	고불륨 경로
DSP	없음 (자체 OSP)	—	DSP-free 아키텍처 (Optical Signal Processor) 가 차별점 Substack
핵심 IP	Rockley Photonics 특허 인수	—	EAM·광스위치·SiP 특허 포트폴리오 Optics.org
광원	외부(비공개)	—	일부 POET/Sivers 경유 추정 — 확인 필요
투자 접근	Marvell 지분이 유일 경로	상장 (MRVL)	유익미 매출 ~2028말, \$1B run-rate ~2029말 Optics.org

자료: 대신증권 Research Center

광 I/O 시대. 세 기술은 공존할 가능성도 존재함

승자기술을 예단하기 어려우며, 공존할 가능성도 존재함

- 세 방식의 차이는 결국 채택 용이성과 구조적 문제의 해결력이라는 측면에서 상충관계에 있다. Ayar Labs는 스케일 아웃 CPO 방식과 동일하게 XPU 옆에 광 I/O 칩렛을 붙이는 방식이라 현재 스케일아웃 밸류체인을 활용할 수 있다. 이 때문에 세 가지 방식 중 가장 빠르게 양산에 돌입할 것으로 예상되며, 표준 기술로 채택될 가능성이 높다. 다만, 뒤의 두 기술과 비교하여 Beach Front 출력 부문에서 열위에 있다는 단점이 있다.
- 반대로 Lightmatter는 칩 밑에 광 인터포저를 까는 방식으로 Beach Front측면에서 Ayarlabs방식 대비 구조적인 우위에 있다. 다만, 멀티래티클 웨이퍼급 인터포저에 따른 제조 난이도, 패키징, 수율, 열 처리 문제가 있어 기술 실증을 넘어 양산 단계로 넘어가는 고비가 존재한다.
- Celestial AI 방식은 XPU, HBM 등 칩에 곧 바로 광 변조기를 부탁하는 방식으로 XPU간 연결을 넘어 메모리 연결까지 광 패브릭으로 확장하려는 접근이다. KV Cahce, CXL, 메모리 풀링까지 겨냥한다는 점에서 가장 급진적이지만, 그만큼 생태계 통합 난이도와 상용화 타이밍의 불확실성이 존재한다.
- 초기 양산에 있어서는 Ayar Labs방식이 가장 속도가 빠를 것으로 보여, 표준 기술로 채택될 가능성이 높다. 그러나 컨텍스트 길이가 길이가 계속 커지고 추론에서 메모리 계층화가 중요해질 수록, 물리계층의 경쟁축은 단순 XPU간 연결에서 XPU와 메모리를 함께 묶는 구조로 확장될 가능성도 존재한다. 이 관점에서 Marvell의 Celestial AI 인수 이후, NVDIA의 \$20 Bill 투자는 특별한 함의점이 있다.
- 엔비디아는 프로토콜에서 NVLink - Fusion으로 프로토콜 표준을 잡아나감과 동시에 ASIC 진영과의 호환 및 병용을 전제하고 있다. 이렇게 볼 때, Ayarlabs 방식뿐 아니라 다양한 광 I/O의 물리계층 구현방식으로 한 개의 데이터센터에서 병용되어 사용될 가능성도 존재한다.
- 현 시점, 광 I/O 분야에서는 뚜렷한 승자 기술이 없지만, 어떤 기술이 되든 광 밸류체인내 중복되는 분야가 있다. 이들에 주목할 때다.

Ayarlabs, LightMatter, Celestial AI 세 진영 방식의 차이점

축	Ayar Labs	Lightmatter	Celestial AI
통합 위치	칩 옆(칩렛)	칩 밑(인터포저)	패브릭+브리지
핵심 제품	TeraPHY + SuperNova	Passage M1000	Photonic Fabric
주 타겟	XPU↔XPU 범용	최대 대역폭 밀도	XPU↔XPU + 메모리
성능 전장	중	최상(114Tbps)	상(16T 칩렛)
채택 난이도	낮음(볼트온)	높음(웨이퍼급)	중간
표준성	표준(UCIe 등)	독자 성향	독자→Marvell 통합
레이저	외부(Sivers)	외부	외부
파운드리/패키징	GF Fotonix	GF Fotonix/Amkor	Marvell 내재화
투자 접근	비상장→밸류체인	비상장	Marvell(상장)

자료: 대신증권 Research Center

V. 투자전략

스케일 아웃에서 스케일업으로 이어지는 밸류체인에 주목

밸류체인이 중복되는 기업들에 주목할 필요

- 광 I/O 시장에서 아직 결정되지 않은 것은 프로토콜과 광 - 연산 통합구조다. 반면, 레이저 광원, 광소자, 파이버 결합, 정밀 정렬, 패키징 및 테스트가 필요하다는 방향성은 어느 아키텍처가 채택되더라도 동일하게 수혜를 볼 것으로 예상된다.
- 따라서, 현재 시점에서 투자전략은 특정 기술이 광 I/O에서 채택되는지에 대해서 예상하기 보다는 여러 진영에서 공통적으로 수혜를 보는 아키텍처 중립적인 기업에 주목할 시점으로 판단한다.
- 특히, 지금 광원 및 광섬유 업체들 같은 업체들 같은 경우에는 현재 스케일아웃과 CPO용 외부레이저 시스템으로 매출을 확보하면서, 향후 스케일업 광 I/O로 제품 확장되는 국면에서도 수혜가 되는 구조를 띄고 있어서, 광 시장에 투자한다고 했을 때 가장 수혜 가시성이 높은 분야다.

Ayarlabs, LightMatter, Celestial AI 세 진영 방식의 차이점

밸류체인 레이어	대표 기업	① Ayar Labs 표준화 (칩 옆 · 칩렛)	② Lightmatter 채택 (칩 밑 · 인터포저)	③ Celestial 확장 (패브릭 · 메모리)	시나리오 중복도
광 엔진 설계 (PIC · 패브릭)	Ayar Labs* · Lightmatter* · Marvell (Celestial 인수)	Ayar Labs*	Lightmatter*	Marvell	1 / 3
외부 레이저 광원	Coherent · Lumentum · Sivers 패키지: RF머트리얼즈	●	●	●	3 / 3
SiPh 파운드리	GlobalFoundries (Fotonix) · TSMC (COUPE) · Tower	●	●	①	2 / 3
패키징 · 인터포저	Amkor · ASE 등 OSAT	①	●	●	3 / 3
조립 · 얼라인먼트 장비	ficonTEC* 등 성호전자 (ADSTEK)	●	●	●	3 / 3
파이버 어태치 · 커넥터	Corning · Fabrinet	●	●	●	3 / 3
테스트 · 번인	티에프이 등	●	●	●	3 / 3

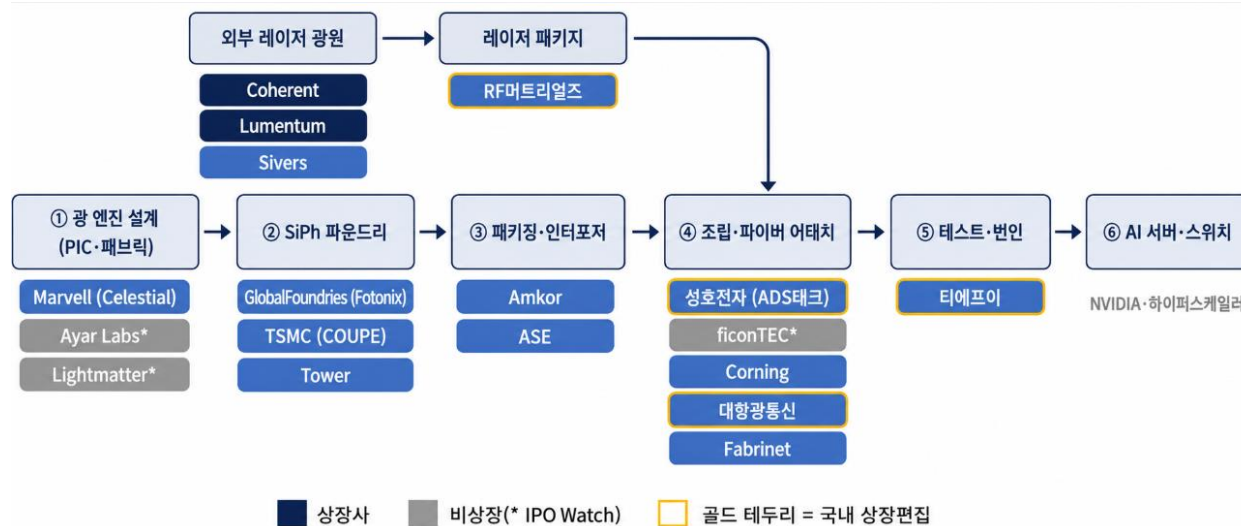
자료: 대신증권 Research Center

광 시장과 국내 수혜기업 밸류맵

광정렬 장비, 패키징에서 관련된 회사들 상장

- 글로벌 광 밸류체인의 중심은 고속 광트랜시버에서 CPO와 스케일업, 광 I/O로 이동하고 있으나, 국내 상장사의 직접적인 기술 노출도는 아직 제한적이다. Ayar Labs, Lightmatter 등 글로벌 업체들이 GPU·ASIC 패키지 인근에서 전기 신호를 광으로 전환하는 핵심 기술을 주도하는 반면, 국내 업체들은 광트랜시버, 광소자 패키지, 광섬유·광케이블 등 후방 밸류체인을 중심으로 참여하고 있다.
- 국내 업체 중에서는 성호전자가 유일하게 광 I/O 시장의 개화에 따른 직접적인 수혜가 예상된다. 다른 국내 업체들의 경우, 광 I/O가 아닌, 스케일업으로 시장에서의 CPO 기술의 확장(RF머트리얼즈). 데이터센터 스케일업으로 광케이블(대한광통신) 등 AI 네트워크 대역폭 확대가 기존 광부품과 광케이블 수요로 확산되는 과정에서 발생하고 있다.

광 시장 밸류체인과 국내 수혜기업 밸류맵



광 I/O 투자전략: 국내 밸류체인 3선

CPO 매출 가시성 확보한 성호전자, RF머트리얼즈 주목. 대한광통신은 실적 대비 밸류에이션 부담

- 글로벌에서는 Lumentum, Coherent 같이 스케일업, 스케일아웃, 스케일러크로스에 걸쳐서 수혜를 보는 광원 회사와, Ayar Labs향 광원 매출로 턴어라운드 기대되는 Sivers 같은 회사들에 주목할 필요가 있다. 또, 광섬유 회사로는 데이터센터 내부향 매출에서 높은 지위를 지닌 Corning, Fujikura가 있다. CPO향 SiPh 생산과 관련해서 지위가 있는 GlobalFoundries, Tower Semi와 같은 8인치 파운드리 업체에도 주목할 필요가 있다. 결국 CPO 형태로 갈 것이라는 점에서 Co-Package를 하는 TSMC, Amkor, SPIL 같은 파운드리 및 OSAT 업체 역시 수혜가 기대된다. 웨이퍼 번인 등 테스트가 중요해진다는 점에서 AEHR 같은 장비사 역시 중장기적인 수혜가 기대된다.
- 국내에서는 액티브 얼라인 분야에서 선도적인 지위를 지니고 있는 성호전자에 관심을 가질 필요가 있다. 성호전자는 국내에서 유일하게 스케일업 시장에서 수혜가 기대되나, 최근 광시장 조정에도 동반하여 큰 폭 조정이 발생하였다. 펀더멘탈이 유효하다는 점에서 중장기적 관점에서 관심을 가질 필요가 있다. RF머트리얼즈의 경우, 루멘텀향 스케일업 매출이 견조한 가운데, 증산이 지속되고 있다는 점에서 실적 개선을 기대한다. 또, 대한광통신도 수혜가 예상된다. 북미 하이퍼스케일러향 매출, 인캡 인수 효과뿐 아니라 중국내 광섬유 가격 상승 수혜에 따른 턴어라운드가 기대된다. 다만, 대한광통신은 밸류에이션에 대한 부담은 주의가 필요하다.

성호전자 · RF머트리얼즈 · 대한광통신 실적 전망 및 밸류에이션

				시가총액	매출액	영업이익	순이익	OPM	ROE	PER	PBR	EV/EBITDA
				(억원)	(억원)	(억원)	(억원)	(%)	(%)	(배)	(배)	(배)
성호전자	투자의견	N.R	25A	11,735	2,316	76	910	3.3	46.3	6.9	2.6	46.9
	목표주가	N.R	26F		2,724	507	989	18.6	32.7	12.8	3.6	15.4
	상승여력	N.R	27F		6,348	1,397	997	22.0	24.5	12.7	2.8	0.4
RF머트리얼즈	투자의견	N.R	25A	3,971	641	74	71	11.5	15.0	31.4	4.2	21.1
	목표주가	N.R	26F		854	135	90	15.8	15.7	43.1	6.3	22.9
	상승여력	N.R	27F		1,012	169	117	16.7	17.3	33.2	5.3	18.3
대한광통신	투자의견	N.R	25A	15,098	1,394	-214	-280	-15.4	-52.1	N/A	4.4	N/A
	목표주가	N.R	26F		2,002	-23	-32	-1.1	-4.3	N/A	19.9	575.3
	상승여력	N.R	27F		2,452	147	110	6.0	12.3	146.0	16.4	84.5

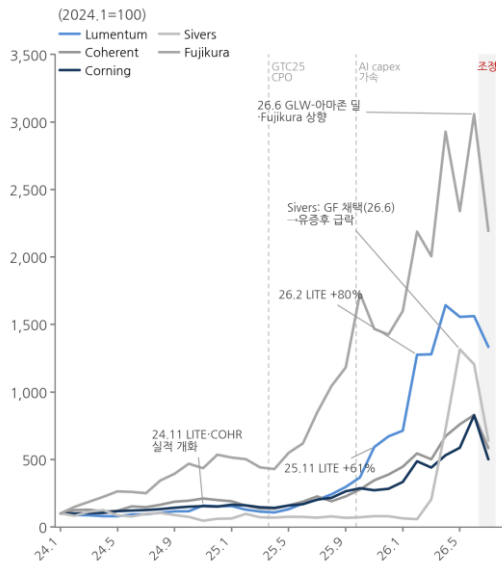
자료: 대신증권 Research Center

광 업체 주가로 읽는 AI 사이클: 개화 → 가속 → 첫 조정

글로벌 광업체들 주가에 동행하는 성격을 지님

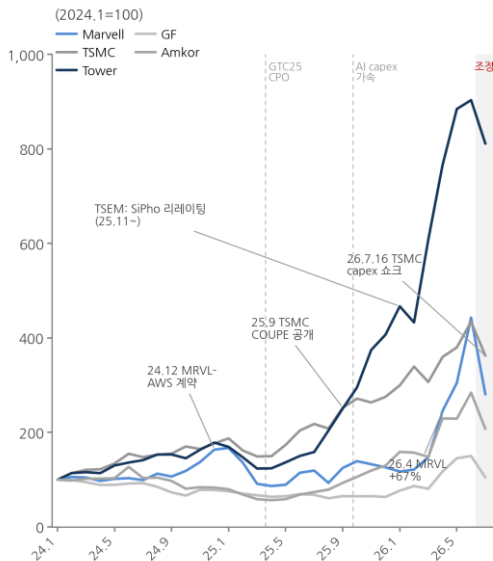
- 글로벌 광업체들의 주가는 2024년 800G 광모듈 실적 개화(루멘텀·코히런트)와 후지쿠라 재평가 시작되었다. 랠리는 2025년 GTC25 CPO 공개, NVIDIA-OpenAI 10GW 계약을 거치며 광 밸류체인 전반의 리레이팅으로 확산되었다.
- 2026년 상반기 어닝 서프라이즈와 대형 수주(코닝-아마존, 타워 \$1.3B 등)로 주가는 지속적으로 강세를 기록했다. 6월 이후, 밸류에이션 부담과 TSMC 발 과잉투자 우려(7/16)로 고점대비 -10~-56% 조정이 진행중이다.
- 국내 3사는 2025년 하반기 개별 모멘텀(루멘텀향 공급, ADS테크·인캡 인수)으로 글로벌 대비 주가 상승은 후행하였다. 격렬하게 급등했던 만큼 조정폭도 글로벌 업체들 대비 크게 나타나는 중이다. 추후, 실적에 따라 차별적인 흐름 나타날 것으로 예상된다.

글로벌 ① 광 부품·모듈·소재



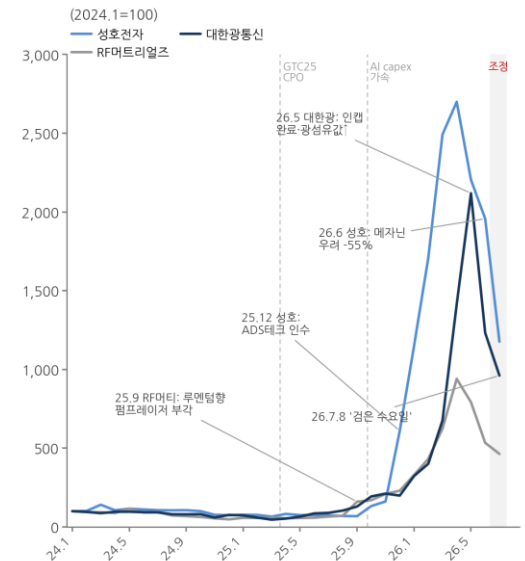
자료: Yahoo Finance, 각사 공시, 대신증권 Research Center

글로벌 ② 반도체·파운드리·패키징



자료: Yahoo Finance, 각사 공시, 대신증권 Research Center

국내 광통신 3사



자료: Yahoo Finance, 각사 공시, 대신증권 Research Center

VI. 기업분석

성호전자 (043260)

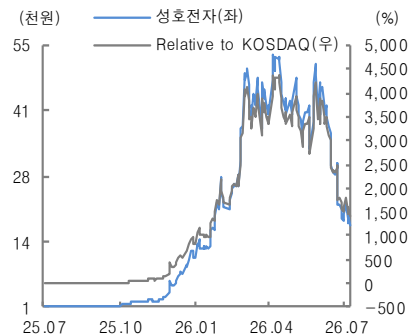
스케일업이 열리면, 기대되는 실적 성장

투자이견	N.R
목표주가	N.R
현재주가 (25.07.20)	16,600원

KOSDAQ	751.59
시가총액	11,878억원
시가총액비중	0.22%
자본금(보통주)	355억원
52주 최고/최저	53,200원 / 1,027원
120일 평균 거래대금	981억원
외국인지분율	4.84%
주요주주	서릉전자 외 14 인 55.06 %

주가수익률(%)	1W	1M	3M	12M
절대수익률	-56.1	-64.2	36.2	1,443.7
상대수익률	-43.5	-47.5	64.2	1,483.5

- 북미 광섬유 회사향 고성능 광 정렬장비 매출 본격화로 26년, 27년 성장 이어갈 것으로 기대됨. 북미 광섬유 회사향 장비는 동사의 기존 장비대비 ASP가 2배 정도 되는 고성능 장비. 데이터센터의 광섬유 밀도가 높아짐에 따라, 고성능 장비의 필요성 커지는 중. 글로벌 경쟁사 대비 납기 및 성능에서 우위를 나타내고 있어, 사실상 솔벤더에 가까운 지위로 파악됨
- 실적 성장에 발맞춘 생산능력 확대. 현재 동사의 생산능력은 기존 트랜시버 장비향 약 300대, CPO향 약 300대로 총 600대 정도 규모. 동사는 동탄 공장 증설 1,2차를 통해서 총 2,000대에 가까운 캐파 확장을 고려 중. 27년부터 순차적으로 생산능력 증대 효과 나타날 것
- CPO 사업을 영위하는 다양한 사업자들과 신규 고객사향 장비를 연구개발용으로 공급하는 단계에 있으며, 기술적 로드맵은 27년 상반기, 본격적인 설비투자는 27년 하반기에 걸쳐서 진행될 것으로 예상
- 26년, 27년 매출 성장은 코닝향 장비가, 중장기 밸류에이션 확장은 CPO향 매출 성장이 동사의 주가 향방을 결정할 것. 희석가치 고려한 27년 예상 P/E 20배로, 동사의 액티브 얼라인먼트 장비에서의 입지, 향후 광전환으로의 필연성을 고려할 때, 중장기 관점에서 매수관점 유효하다고 판단함



	2023A	2024A	2025A	2026F	2027F
매출액	2,081	2,073	2,316	2,724	6,348
영업이익	259	63	76	507	1,397
세전이익	229	131	1,173	1,268	1,278
총당기순이익	176	81	910	989	997
지배지분순이익	177	80	910	989	997
EPS	331	138	1,341	1,395	1,406
PER	4.2	8.1	6.9	12.8	12.7
BPS	1,894	2,334	3,536	4,995	6,473
PBR	0.7	0.5	2.6	3.6	2.8
ROE	19.2	6.5	46.3	32.7	24.5

주: EPS와 BPS, ROE는 지배지분 수치 기준, 단위: 억원, 원, %
자료: 대신증권 Research Center

성호전자 (043260)

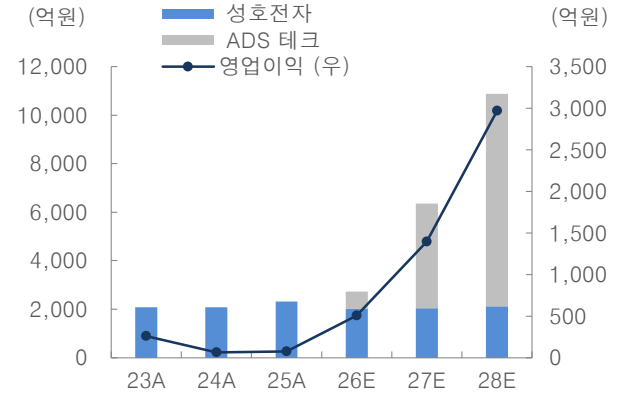
ADS테크를 타고 계속될 실적 성장

영업실적 추이 및 전망

	22A	23A	24A	25A	26E	27E	28E
매출액	1,536	2,081	2,073	2,316	2,724	6,348	11,878
YoY(%)	15.3%	35.5%	-0.4%	11.7%	17.6%	133.0%	87.1%
성호전자	-	2,081	2,073	2,316	2,011	2,033	2,114
ADS테크	-	-	-	-	714	4,317	8,764
매출총이익	245	462	360	443	943	2,158	3,920
YoY(%)	38.2%	88.7%	-22.1%	23.1%	113.1%	128.8%	81.6%
총이익률(%)	15.9%	22.2%	17.3%	19.1%	34.6%	34.0%	33.0%
영업이익	20	259	63	76	507	1,397	2,970
YoY(%)	66.3%	1164.0%	-75.6%	19.6%	571.1%	175.6%	112.6%
영업이익률(%)	1.3%	12.4%	3.0%	3.3%	18.6%	22.0%	25.0%
세전이익	-56	229	131	1,173	1,268	1,278	2,845
YoY(%)	-263.9%	-510.9%	-42.8%	795.1%	8.1%	0.8%	122.6%
세전이익률(%)	-3.6%	11.0%	6.3%	50.7%	46.6%	20.1%	24.0%
당기순이익	-42	176	81	910	989	997	2,219
YoY(%)	-232.3%	-519.4%	-54.1%	1027.2%	8.7%	0.8%	122.6%
당기순이익률(%)	-2.7%	8.5%	3.9%	39.3%	36.3%	15.7%	18.7%

자료: 성호전자, 대신증권 Research Center 추정

성호전자 부문별 실적추이 및 전망



자료: 성호전자, 대신증권 Research Center 추정

성호전자 광섬유 정렬 시스템



자료: 성호전자, 대신증권 Research Center 추정

성호전자 (043260)

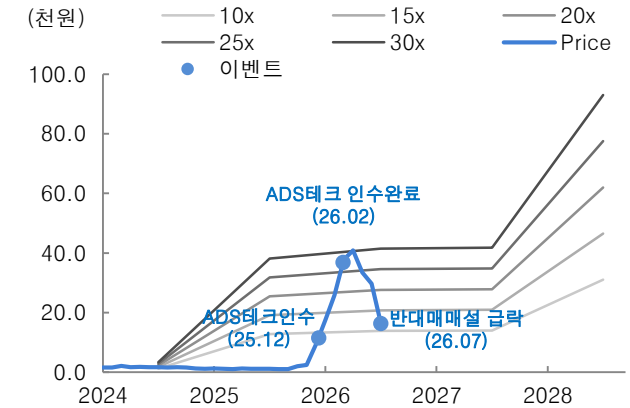
본격적인 실적 성장은 27년부터 예상

성호전자 분기별 실적 추이

	1Q25	2Q25	3Q25	4Q25	1Q26
매출액	1,536	2,081	2,073	2,316	2,724
YoY(%)	15.3%	35.5%	-0.4%	11.7%	17.6%
성호전자	-	2,081	2,073	2,316	2,011
ADS테크	-	-	-	-	714
매출총이익	245	462	360	443	943
YoY(%)	38.2%	88.7%	-22.1%	23.1%	113.1%
총이익률(%)	15.9%	22.2%	17.3%	19.1%	34.6%
영업이익	20	259	63	76	507
YoY(%)	66.3%	1164.0%	-75.6%	19.6%	571.1%
영업이익률(%)	1.3%	12.4%	3.0%	3.3%	18.6%
세전이익	-56	229	131	1,173	1,268
YoY(%)	-263.9%	-510.9%	-42.8%	795.1%	8.1%
세전이익률(%)	-3.6%	11.0%	6.3%	50.7%	46.6%
당기순이익	-42	176	81	910	989
YoY(%)	-232.3%	-519.4%	-54.1%	1027.2%	8.7%
당기순이익률(%)	-2.7%	8.5%	3.9%	39.3%	36.3%

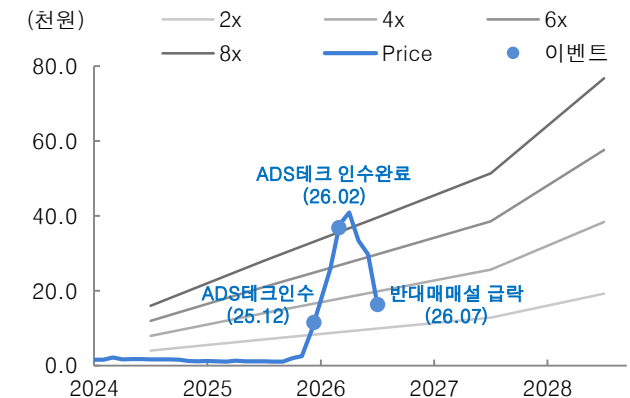
자료: 대신증권 Research Center

성호전자 PER 차트



자료: 대신증권 Research Center

성호전자 PBR 차트



자료: 대신증권 Research Center

RF머트리얼즈 (327260)

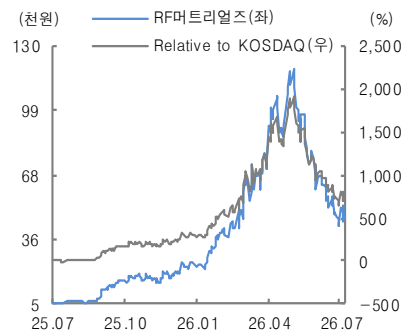
DCI에서 CPO로, 루멘텀이 이끄는 실적 성장

투자이견	N.R
목표주가	N.R
현재주가 (25.07.20)	43,850

KOSDAQ	751.59
시가총액	3,944억원
시가총액비중	0.07%
자본금(보통주)	42억원
52주 최고/최저	118,700원 / 5,612원
120일 평균 거래대금	223억원
외국인지분율	19.06%
주요주주	알에프에이치아이씨 외 9인 48.78% Morgan Stanley 5.09%

주가수익률(%)	1W	1M	3M	12M
절대수익률	-32.0	-51.4	95.9	669.7
상대수익률	-12.5	-28.7	136.2	689.5

- 루멘텀향 라만 펌프레이저 패키지가 동사 실적을 견인하는 가운데, 고객사 요청에 따라 공급 물량 증산 지속적으로 이루어지고 있는 중. 동사는 데이터센터간 연결에 사용되는 DCI – Scale across 네트워크 매출
- 라만 펌프레이저 패키지에 이어, 스케일러크로스에 사용되는 CPO 펌프레이저 패키지 매출도 샘플 공급 발생 중. 스케일러크로스에서도 CPO가 확대됨에 따라 동사 전체 매출에서 CPO향 매출 비중이 점차 확대 될 것으로 기대
- 신공장 투자를 통한 CAPA 투자 진행 중. 신공장은 27년 4분기 입주. 28년 1분기 장비 투입 및 양산 본격화를 예상함. 26년 ~ 27년은 기존 공장 가동률 상승, 28년부터는 증설에 따른 본격적인 매출 확대를 전망함
- 당사 추정 기준 28년 예상 P/E 19배로 적정 수준 밸류에이션으로 판단함. 다만, 이는 신규 공장 증설에 따른 매출 증대 효과를 보수적으로 산정한 것으로, 스케일러크로스 시장 역시 본격적인 성장 구간에 돌입했다는 점, 글로벌 1위 업체인 루멘텀향 매출이 발생하고 있다는 점에서 투자 메리트는 충분한 구간으로 판단함



	2023A	2024A	2025A	2026F	2027F
매출액	479	445	641	854	1,012
영업이익	0	-15	74	135	169
세전이익	4	-75	83	128	165
총당기순이익	7	-76	104	100	129
지배지분순이익	-7	-48	71	90	117
EPS	-80	-568	837	1,068	1,387
PER	N/A	N/A	31.4	43.1	33.2
BPS	5,259	4,871	6,259	7,321	8,680
PBR	2.0	1.0	4.2	6.3	5.3
ROE	-1.5	-11.2	15.0	15.7	17.3

주: EPS와 BPS, ROE는 지배지분 수치 기준, 단위: 억원, 원, %
자료: 대신증권 Research Center

RF머트리얼즈 (327260)

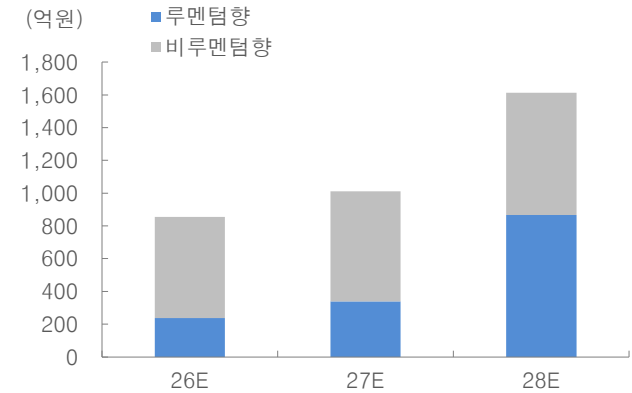
루멘텀 스케일러크로스 매출로 가팔라지는 실적 성장

RF머트리얼즈 영업실적 추이 및 전망

	23A	24A	25A	26E	27E	28E
매출액	479	445	641	854	1,012	1,614
yoy (%)	-5.0%	-7.1%	44.0%	33.3%	18.5%	59.5%
루멘텀	-	-	-	238	340	867
비루멘텀	479	445	641	616	672	747
매출총이익	106	95	179	255	301	500
yoy (%)	-14.5%	-10.7%	89.1%	42.5%	18.0%	66.5%
매출총이익률 (%)	22.1%	21.2%	27.9%	29.8%	29.7%	31.0%
영업이익	0	-15	74	135	169	291
yoy (%)	-101.3%	4096.1%	-607.3%	83.6%	24.9%	71.9%
영업이익률 (%)	-0.1%	-3.3%	11.5%	15.8%	16.7%	18.0%
세전이익	4	-75	83	128	165	289
yoy (%)	-87.2%	-2031.6%	-210.3%	55.1%	28.4%	75.7%
세전이익률 (%)	0.8%	-16.9%	12.9%	15.0%	16.3%	17.9%
당기순이익	-7	-48	71	90	117	217
yoy (%)	-120.0%	616.6%	-247.3%	27.5%	29.9%	85.3%
당기순이익률 (%)	-1.4%	-10.8%	11.0%	10.5%	11.6%	13.4%

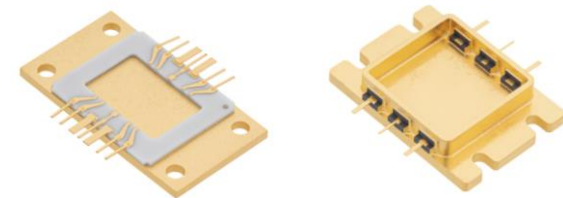
자료: RF머트리얼즈, 대신증권 Research Center

루멘텀 및 비루멘텀향 매출 추이



자료: RF머트리얼즈, 대신증권 Research Center

RF머트리얼즈 RF 및 Microwave 패키지



자료: RF머트리얼즈, 대신증권 Research Center

RF머트리얼즈 (327260)

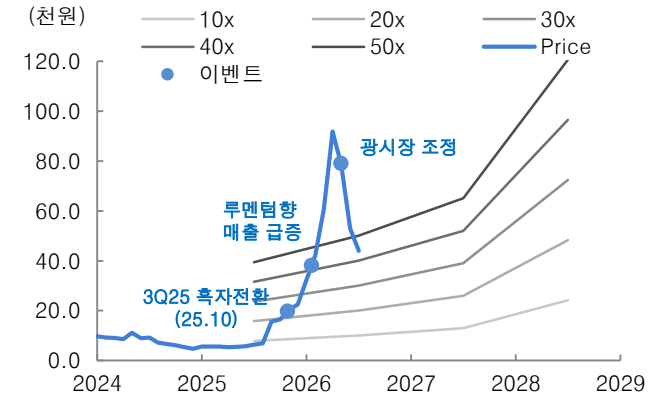
루멘텀향 매출이 본격적으로 성장 드라이브로 작용

RF머트리얼즈 분기별 실적 추이

	1Q25	2Q25	3Q25	4Q25	1Q26
매출액	118	166	146	211	200
yoy (%)	12.2%	45.5%	89.2%	42.0%	69.6%
매출총이익	30	47	46	56	60
yoy (%)	32.5%	127.7%	209.4%	53.9%	98.2%
매출총이익률 (%)	25.6%	28.2%	31.3%	26.6%	29.9%
영업이익	8	19	19	28	32
yoy (%)	흑전	흑전	흑전	566.6%	281.9%
영업이익률 (%)	7.1%	11.3%	13.0%	13.1%	16.0%
세전이익	9	20	24	30	41
yoy (%)	흑전	흑전	흑전	흑전	376.1%
세전이익률 (%)	7.4%	12.4%	16.1%	14.2%	20.7%
당기순이익	6	22	23	53	34
yoy (%)	흑전	흑전	흑전	흑전	499.5%
당기순이익률 (%)	4.8%	13.1%	15.9%	25.2%	16.9%

자료: RF머트리얼즈, 대신증권 Research Center

RF머트리얼즈 PER 차트



자료: RF머트리얼즈, 대신증권 Research Center

RF머트리얼즈 PBR 차트



자료: RF머트리얼즈, 대신증권 Research Center

대한광통신 (010170)

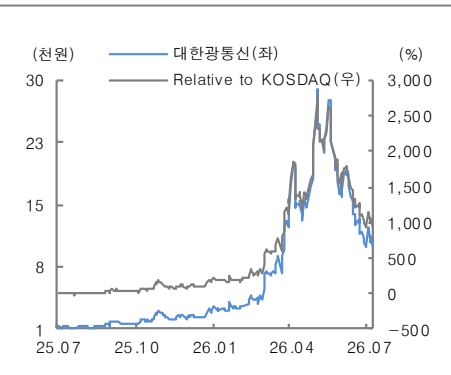
데이터센터가 이끄는 본격적인 턴어라운드

투자이견	N.R
목표주가	N.R
현재주가 (25.07.20)	9460원

KOSDAQ	791.84
시가총액	14,709억원
시가총액비중	0.28%
자본금(보통주)	777억원
52주 최고/최저	29,100원 / 848원
120일 평균 거래대금	3,715억원
외국인지분율	4.23%
주요주주	티에프오인더스트리 외 2인 14.60%

주가수익률(%)	1W	1M	3M	12M
절대수익률	-46.9	-31.4	243.3	962.1
상대수익률	-31.7	0.7	313.9	989.5

- 인캡 인수에 따른 매출 효과 및 북미 하이퍼스케일러향 864심 고다심 광섬유 매출 본격화로 실적 성장과 턴어라운드가 시작될 것으로 예상. 인캡은 BEAD향 통신사 물량, 하이퍼스케일러향 864심 고다심 광섬유는 데이터센터간 연결에 사용되는 스케일러크로쓰향 매출
- 북미 하이퍼스케일러향 매출이 발생한다는 점에서 고무적이며, 데이터센터향 매출이 전사 매출 성장을 견인할 것으로 기대함. 다만, 스케일업 등 데이터센터 내부 매출은 아직 기대하기 이른 시점으로 판단
- 중국내 데이터센터 증설에 따른 효과도 본격적으로 이뤄지는 중. Huan GL Technology에 따르면 km당 광섬유 가격은 25년 중반 \$12 ~ \$14/km에서 현재 \$22/km로 70% 이상 가격이 상승한 상황. AI 용뿐 아니라 일반용까지 글로벌 광섬유 가격이 상승한 상황으로, 동사에 우호적임
- 당사 추정 기준 28년 예상 P/E 140배로 주가가 조정 받았으나, 여전히 펀더멘탈적으로 부담스러운 구간. 현재 광시장의 TAM 폭발은 데이터센터 내부에 들어가는 스케일업 시장이라는 점에서 동사의 주가는 여전히 실적 가시성 대비 기대감이 높은 구간. 턴어라운드, 캐파 증설 등 펀더멘탈 개선폭이 당사 추정치를 크게 상회하는지 등에 대해서 지켜볼 것을 추천



	2023A	2024A	2025A	2026F	2027F
매출액	1,803	1,527	1,394	2,002	2,452
영업이익	-232	-297	-214	-23	147
세전이익	-288	-548	-272	-29	141
총당기순이익	-295	-560	-280	-32	110
지배지분순이익	-295	-560	-280	-32	110
EPS	-347	-660	-239	-21	71
PER	N/A	N/A	N/A	N/A	146.0
BPS	1,020	548	505	518	630
PBR	1.3	1.6	4.4	19.9	16.4
ROE	-32.5	-95.9	-52.1	-4.3	12.3

주: EPS와 BPS, ROE는 지배지분 수치 기준, 단위: 억원, 원, %
자료: 대신증권 Research Center

대한광통신 (010170)

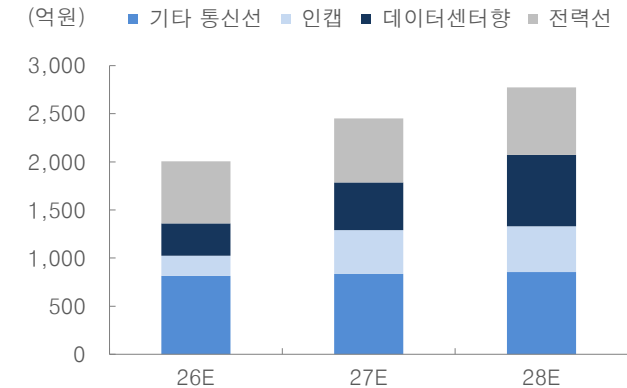
인캡향 매출 본격 반영 시작. 데이터센터향 매출이 실적 이끌 것

대한광통신 연도별 실적추이 및 전망

	23A	24A	25A	26E	27E	28E
매출액	1,803	1,527	1,394	2,002	2,452	2,771
YoY(%)	-5.2%	-15.3%	-8.7%	43.6%	22.4%	13.0%
인캡	-	-	-	211	457	476
데이터센터향	-	-	-	336	494	741
기타통신선	1,162	850	795	814	834	854
전력선	640	678	601	643	667	701
매출총이익	61	33	36	311	539	637
YoY (%)	-80.4%	-45.5%	8.2%	762.8%	73.6%	18.1%
매출총이익률 (%)	3.4%	2.2%	2.6%	15.5%	22.0%	23.0%
영업이익	-232	-297	-214	-23	147	222
YoY (%)	적자전환	적자지속	적자지속	-89.3%	-739.3%	50.7%
영업이익률 (%)	-12.9%	-19.5%	-15.4%	-1.1%	6.0%	8.0%
세전이익	-288	-548	-272	-29	141	215
YoY (%)	적자지속	적자지속	적자지속	-89.4%	-588.6%	52.5%
세전이익률 (%)	-16.0%	-35.9%	-19.5%	-1.4%	5.7%	7.7%
당기순이익	-295	-560	-280	-32	110	167
YoY (%)	855.1%	90.2%	-50.1%	-88.7%	-446.5%	52.5%
당기순이익률 (%)	-16.3%	-36.7%	-20.1%	-1.6%	4.5%	6.0%

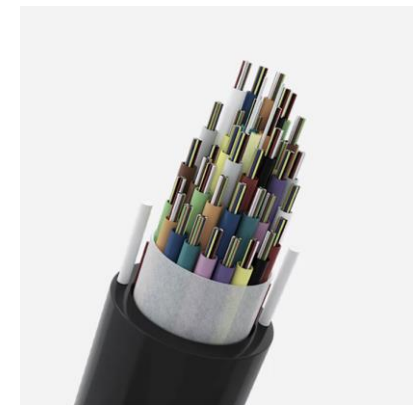
자료: 대신증권 Research Center

대한광통신 부문별 매출액 추이 및 전망



자료: 대한광통신, 대신증권 Research Center

대한광통신 광케이블



자료: 대한광통신, 대신증권 Research Center

대한광통신 (010170)

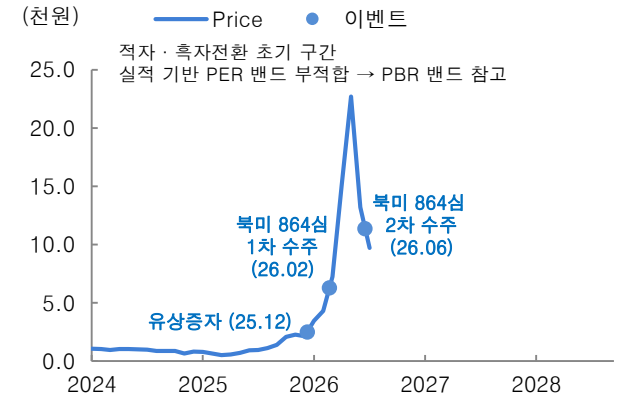
다가오는 분기 턴어라운드

대한광통신 분기별 실적 추이

	1Q25	2Q25	3Q25	4Q25	1Q26
매출액	254	416	433	291	342
YoY(%)	-22.6%	3.7%	-1.4%	-18.7%	34.9%
기타통신선	145	211	303	136	173
전력선	109	206	130	156	170
매출총이익	-2	26	3	9	29
YoY (%)	적지	-40.8%	-85.0%	흑전	흑전
매출총이익률 (%)	-0.7%	6.2%	0.7%	3.1%	8.3%
영업이익	-69	-28	-67	-50	-40
YoY (%)	적지	적지	적지	적지	적지
영업이익률 (%)	-27.3%	-6.7%	-15.5%	-17.2%	-11.6%
세전이익	-70	-11	-82	-109	-37
YoY (%)	적지	적지	적지	적지	적지
세전이익률 (%)	-27.7%	-2.6%	-18.9%	-37.6%	-10.7%
당기순이익	-70	-11	-82	-117	-37
YoY (%)	적지	적지	적지	적지	적지
당기순이익률 (%)	-27.7%	-2.6%	-18.9%	-40.1%	-10.7%

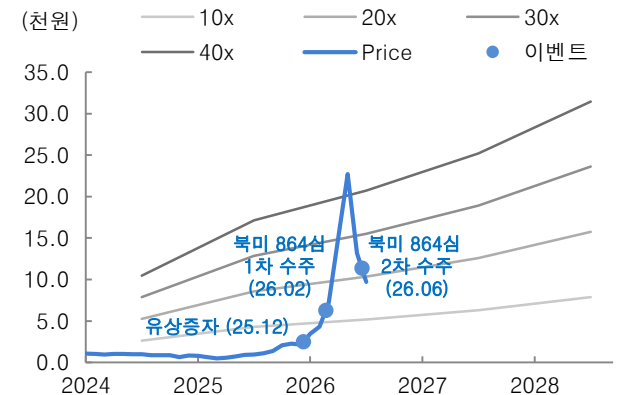
자료: 대신증권 Research Center

대한광통신 PER 차트



자료: 대한광통신, 대신증권 Research Center

대한광통신 PBR 차트



자료: 대한광통신, 대신증권 Research Center

Compliance Notice

- 금융투자업규정 4-20조 1항5호사목에 따라 작성일 현재 사전고지와 관련한 사항이 없습니다.
- 당사의 금융투자분석사는 자료 작성일 현재 본 자료에 관련하여 재산적 이해관계가 없습니다. 당사는 위 언급된 종목을 제외한 동자료에 언급된 종목과 계열회사의 관계가 없으며 당사의 금융투자분석사는 본 자료의 작성과 관련하여 외부의 부당한 압력이나 간섭을 받지 않고 본인의 의견을 정확하게 반영하였습니다. (작성자: 박장욱)
- 본 자료는 투자자들의 투자판단에 참고가 되는 정보제공을 목적으로 배포되는 자료입니다. 본 자료에 수록된 내용은 당사 Research Center의 추정치로서 오차가 발생할 수 있으며 정확성이나 완벽성은 보장하지 않습니다. 본 자료를 이용하시는 분은 동 자료와 관련한 투자의 최종 결정은 자신의 판단으로 하시기 바랍니다.

투자의견 비율공시

구분	Buy(매수)	Marketperform(종립)	Underperform(매도)
비율	91.8	8.2%	0%

기준일자: 2026.07.20

성호전자(043260) 투자의견 및 목표주가 변경 내용



RF머트리얼즈 (327260) 투자의견 및 목표주가 변경 내용



대한광통신(010170) 투자 의견 및 목표주가 변경 내용

