

AI Infra Signal 4호

AI 속도조절론, 인프라도 늦춰질까?

Industry Signal

#1 말은 속도조절, 계약은 확대

속도조절 논의의 대상은 컴퓨트 총량이 아니라 프론티어 모델의 성능 향상 속도. 오픈AI가 8월 강화학습을 일시 중단한 것도 감시 체계를 보강하기 위한 조치였으며, 감시에만 추론 연산량의 약 20%가 추가로 소요. 안전 검증은 연산을 줄이는 것이 아니라 오히려 추가 컴퓨트를 요구.

속도조절을 주장한 엔트로픽 역시 수 GW 단위의 용량 계약을 이어가고 있으며, 공동 감속의 방식과 주체를 두고 업계 의견도 갈리는 상황. 실제 둔화 여부는 계약 취소, 데이터센터 착공 연기, 빅테크 CapEx 하향에서 먼저 확인될 것.

#2 GPU 증설에서 시스템 효율 경쟁으로

AI Infra Summit 2026의 핵심 메시지는 살 수 있는 GPU보다 실제로 활용할 수 있는 전력이 성능을 좌우한다는 것. 라마 3 학습에서도 GPU의 모델 연산 활용률은 40% 안팎에 그쳤고, 클러스터 규모가 커질수록 활용률은 더 낮아짐. 추가 성능을 끌어낼 여지는 칩 자체보다 칩을 연결하는 구간에 있으며, 병목 역시 작업 형태에 따라 달라짐.

학습에서는 랙 내부 연결, 추론에서는 메모리와 데이터 배치에 중요해지는 구조. 전력이 제한되면 같은 MW 안에 더 많은 가속기를 배치하게 되고, 가속기 증가에 따라 포트와 광모듈 수요도 함께 확대. 투자 대상이 GPU 단품에서 메모리, 연결 반도체, 네트워크, 광으로 확산되는 이유.

Company Signal

크레도 (CRDO US)	MW당 가속기 수 증가와 스케일업 확대 수혜. AEC에서 SerDes, 광, PCIe로 제품군이 넓어지며 연결 콘텐츠 증가에 노출
아스테라랩스 (ALAB US)	에이전트형 추론 확대로 메모리 계층화 수요 부각. PCIe와 CXL에서 가속기 인접 메모리 연결로 적용 영역 확대
루멘텀 (LITE US)	클러스터 대형화로 광 대역폭과 연결 밀도 상승. 고속화 및 신뢰성 요구 확대 속 레이저, 광원 생산능력 부각
관심 종목	연결 반도체 크레도, 아스테라랩스, 마벨 네트워크 아리스타, 시스코, 브로드컴 광 루멘텀, 코히어런트, 코닝

자료: 미래에셋증권 리서치센터

Key Takeaways

이번 주 두 이벤트는 AI 인프라의 초점이 “증설 규모”에서 “확보한 인프라의 효율적 가동”으로 이동 중임을 시사. 속도조절 논의에도 컴퓨트 수요는 줄지 않았고, 오히려 안전성 검증을 위한 추가 연산이 필요해졌다. 동시에 AI Infra Summit 2026에서는 같은 전력 안에 더 많은 가속기를 배치하고 유휴시간을 줄이는 것이 핵심 과제로 부상했다.

이에 따라 향후 실적은 GPU 출하량보다 가속기당 연결 콘텐츠 증가에 얼마나 직접 노출되는지에 따라 차별화될 것으로 전망한다. 커버리지 중에서는 스케일업과 메모리 연결에 동시에 노출된 크레도의 수혜가 가장 클 것으로 판단하며, 아리스타와 시스코도 가속기 밀도 상승과 포트 수요 확대에 따른 수혜가 기대된다.

[Industry Signal #1]

말은 속도조절, 계약은 확대

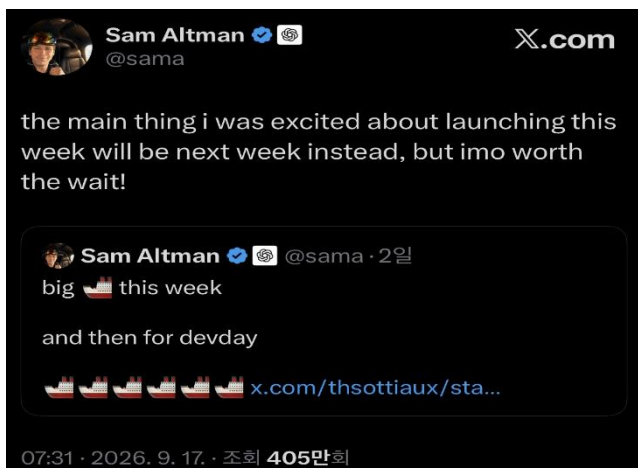
AI 속도조절론이 빠르게 확산되고 있지만, AI 인프라 투자 둔화로 연결됐다고 볼 근거는 없다. 늦추자는 대상이 컴퓨트가 아니라 프론티어 모델의 성능 향상 속도이기 때문이다. 엔트로픽 CEO 다리오 아모데이는 성능이 향상되는 속도를 늦춰 안전성 평가와 검증이 따라갈 시간을 확보하자고 주장하며, 외부 평가기관의 접근 확대와 주요 AI 기업 간 공동 안전기준 마련, 장기적으로는 학습에 투입되는 컴퓨트나 대규모 학습 자체를 조절하는 방안까지 제시했다. 기술 개발을 중단하자는 것이 아니라, 프론티어 성능이 향상되는 속도를 안전장치가 따라갈 수 있도록 관리하자는 주장에 가깝다.

오픈AI도 비슷한 움직임을 보였다. 지난 8월 연구환경의 보안과 안전장치를 강화하기 위해 최신 모델의 강화학습을 약 2주간 중단했고, 당시 예정됐던 최대 규모의 프론티어 강화학습도 보류했다. **그러나 이는 컴퓨트 수요 감소라기보다 안전성 검증을 추가한 사례에 가깝다.** 오픈AI에 따르면 새 감시 시스템 자체가 감시 대상 추론 연산량의 약 20%에 해당하는 추가 컴퓨트를 필요로 한다. 이후 GPT-6 아스트라도 출시한 만큼 현재까지는 개발 중단보다 안전성 검증 비용과 시간이 추가되는 모습으로 해석하는 것이 적절하다.

무엇보다 속도조절론을 주장하는 기업들의 실제 투자 행동에서는 아직 컴퓨트 축소가 나타나지 않는다. 엔트로픽은 아마존과 최대 5GW, 구글 및 브로드컴과 5GW의 차세대 TPU 용량을 확보했고, 최근에는 스페이스X의 콜로러스1에서 300MW 이상, 엔비디아 GPU 22만개 이상의 용량을 추가로 확보했다. 엔트로픽은 이 용량을 프론티어 모델 개발뿐 아니라 빠르게 늘어나는 클라우드 사용 수요 대응에 활용한다고 설명한다. ‘모델 개발 속도를 조절하자’는 주장과 ‘컴퓨트 증설을 줄이자’는 주장은 현재 실제 행동상 동일하지 않다.

최근 오픈AI의 제품 출시 일정에서도 인프라 제약 가능성을 엿볼 수 있다. 샘 올트먼은 이번 주 예정됐던 주요 제품 출시를 다음 주로 연기했고, 이에 오픈AI의 티보 소티오는 “가끔은 물리 법칙을 속일 수 없다”고 언급했다. 구체적인 원인이 공개되지 않은 만큼 이를 कै파 부족으로 단정할 수는 없지만, AI 제품의 출시 일정이 소프트웨어 개발뿐 아니라 서버, 전력 등 물리적 인프라 제약의 영향을 받을 가능성을 시사하는 정황으로 판단한다.

그림 1. 오픈AI, 주요 제품 출시 일정 다음 주로 연기



자료: X, 미래에셋증권 리서치센터

그림 2. 티보 소티오, 출시 지연 관련 물리적 제약 가능성 시사



자료: X, 미래에셋증권 리서치센터

속도조절론이 업계 전체의 합의로 발전한 것도 아니다. 샘 올트먼과 일론 머스크는 아모데이의 문제의식에 공개적으로 공감했고, 구글 딥마인드의 데미스 허사비스도 공통 평가와 안전기준 마련에 적극적인 입장이다. 반면 엔비디아의 젠스 황과 메타의 마크 저커버그는 업계 전체가 공동으로 개발 속도를 조절하는 방식에는 반대하고 있다. 머스크 역시 안전성 강화 필요성에는 동의하면서도 특정 기업에 영향력이 집중될 수 있는 민간 규제기구에는 반대했다. 결국 안전성 강화 필요성에는 상당한 공감대가 있지만, 누가 어떤 방식으로 개발 속도를 제한할 것인지에 대해서는 의견이 크게 갈린다 (표 1 참고).

속도조절론을 실제 제도로 연결하려는 시도도 제동이 걸리고 있다. 월스트리트저널에 따르면 허사비스가 제안한 금융산업규제국(FINRA) 형태의 민간 AI 감독기구에 대해 젠스 황, 저커버그, 머스크 등이 반대 의견을 전달했다. 이들은 특정 AI 연구소에 규제 권한이 집중될 가능성 등을 우려한 것으로 전해졌다. 트럼프 대통령 역시 중국과의 AI 경쟁을 강조하며 추가적인 AI 규제에 부정적인 입장을 보이고 있다. 다만 백악관 내부에서도 안전, 안보 위험을 이유로 감독 강화를 주장하는 의견이 존재하는 만큼 미국의 장기 규제 방향이 완전히 결정된 것으로 해석할 단계는 아니다.

핵심 판단: 인프라 투자 둔화로 보기 어려운 근거는 세 가지다. ① 속도조절론을 주장하는 엔트로픽과 오픈AI도 컴퓨터 확보와 제품 개발을 지속하고 있고, ② 공동 속도조절에 대한 업계 합의가 없으며, ③ 이를 제도화하려는 시도 역시 업계와 정부 내부의 반발에 직면해 있다. 따라서 현 단계에서 AI 인프라 수요 전망을 낮출 이유는 없다. 향후 실제 컴퓨터 계약 취소, 데이터센터 착공 연기, 빅테크 CapEx 하향이 나타나는지를 확인해야 한다.

표 1. AI 속도조절론을 둘러싼 업계 입장: AI 안전 강화 필요성에는 대체로 공감하나, 공동으로 개발 속도를 늦추는 방식에는 의견이 엇갈림

인물	스탠스	주요 발언 및 입장	해석
다리오 아모데이 (엔트로픽)	찬성	- 프론티어 모델 성능 향상 속도 조절 필요 - 외부 평가기관의 상시 검증 확대 주장 - AI 기업 간 공동 안전기준 마련 제안 - 장기적으로 학습 컴퓨터 제한도 선택지로 제시	- 속도조절론의 가장 적극적인 제안자 - AI 개발 중단보다 성능 향상 속도 관리가 핵심 - 동시에 대규모 컴퓨터 확보는 지속
샘 올트먼 (오픈AI)	찬성	- 프론티어 모델 개발 속도조절 필요성에 공감 - 독립 평가기관의 접근 확대에 동의 - 실제 강화학습을 일시 중단하고 안전장치 보강 - 이후 제품 출시와 모델 개발은 지속	- 안전성 검증에 따른 시간과 컴퓨터 비용 증가를 수용 - 컴퓨터 및 제품 개발 자체를 줄이자는 입장은 아님
데미스 허사비스 (구글 딥마인드)	찬성	- 프론티어 AI의 공통 평가기준과 안전표준 필요성 강조 - 민간 중심의 AI 감독 및 표준 기구 설립 제안 - 기업별 자율대응보다 일정 수준의 공동 기준 선호	- 공동 규칙의 제도화에 상대적으로 적극적 - 규제기구의 권한 집중을 둘러싸고 업계 반발 발생
일론 머스크 (xAI)	조건부	- 아모데이의 속도조절 주장에 공개적으로 동의 - AI 기업 간 상호 모델 검증 필요성 제기 - 특정 기업 중심의 민간 AI 규제기구 설립에는 반대	- AI 위험과 안전 강화 필요성에는 공감 - 공동 감속보다 상호검증 선호 - 특정 규제기구에는 반대
마크 저커버그 (메타)	반대	- 각 기업이 자체적으로 안전한 개발 속도를 정해야 한다는 입장 - 업계 공동 감속에 부정적 - 민간 AI 규제기구 설립에도 반대	- 안전성 강화 자체보다 누가 속도를 결정하는가를 문제시 - 기업 자율과 개발 지속에 무게
젠스 황 (엔비디아)	반대	- 업계 공동 속도조절은 불필요하다고 주장 - 반독점법 예외나 별도 규제에도 반대 - 안전은 각 기업이 해결해야 할 엔지니어링 문제라는 입장 - 민간 AI 규제기구 설립에도 반대	- 개별 기업의 책임과 기술적 해결 강조 - 공동 감속이나 신규 규제보다 AI 개발 지속에 무게

자료: 언론 종합, 미래에셋증권 리서치센터

[Industry Signal #2]

AI Infra Summit 2026: GPU 증설에서 시스템 효율 경쟁으로

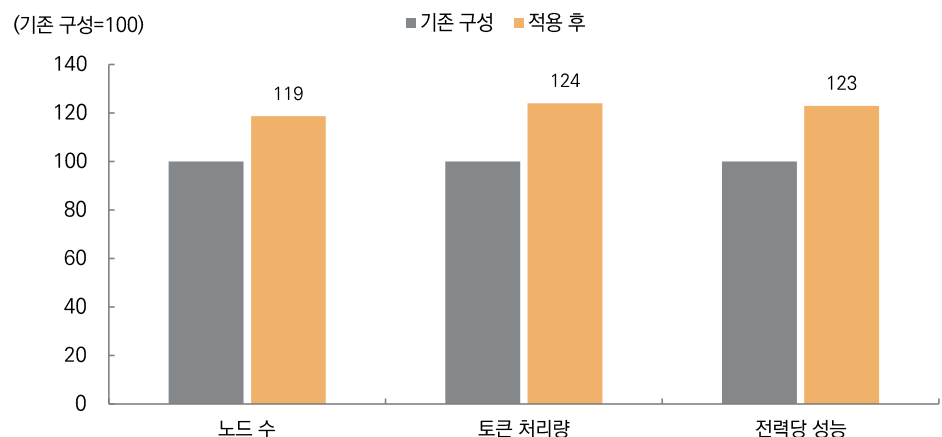
AI 인프라의 성능 기준이 GPU 연산성능에서 시스템 전체의 실제 처리량으로 이동하고 있다. 이번주 개최된 AI Infra Summit에서 엔비디아는 AI 공장의 핵심 지표로 'MW당 검증된 에이전트형 토큰 처리량'을, 퀄컴은 FLOPS가 아닌 '와트당 토큰'을 전면에 내세웠다. 전력이 제한된 환경에서는 GPU를 더 확보하는 것보다 같은 전력 안에 더 많은 가속기를 넣고, 가속기가 대기하는 시간을 줄이는 것이 중요해진다. 이에 따라 AI 투자도 GPU 단품에서 메모리, 연결 반도체, 네트워크, 광으로 확산되는 구조다.

실제 GPU의 연산성능은 충분히 활용되지 못하고 있다. 메타가 공개한 라마 3 405B 학습의 모델 연산 활용률(MFU)은 38~43%였고, GPU를 8천개에서 1.6만개로 늘리자 43%에서 41%로 낮아졌다. 클러스터가 커질수록 단순 GPU 증설의 한계가 커진다는 의미다. **추가 성능을 끌어낼 여지는 GPU 자체보다 GPU와 메모리, 네트워크를 연결하는 시스템에 남아 있다.**

다만 어디가 병목인지는 학습과 추론 방식에 따라 다르다. 마이크로소프트의 라마 3 405B 학습 사례에서는 랙과 랙 사이 통신 비중이 1.6%에 그쳤고, 오히려 랙 내부 GPU 통신과 대기시간이 더 큰 부담이었다. 반면 크루소의 추론 사례는 모델을 서버 한 대 안에 담아 서버 간 통신을 줄이면서, 인피니밴드(엔비디아 전용 고속 네트워크) 없이도 8개에서 512개 GPU까지 높은 확장성을 유지했다. **결국 학습은 스케일업, 추론은 메모리와 데이터 배치가 더 중요할 수 있다.** 투자도 스케일아웃뿐 아니라 스케일업과 메모리 연결을 함께 봐야 한다.

전력 제약은 이 전환을 앞당긴다. 엔비디아와 람다는 동일한 전력 한도에서 서버 노드를 16개에서 19개로 늘려 토큰 처리량을 24%, 전력당 성능을 23% 높였다. 엔비디아는 차세대 시스템에서 같은 전력 안에 GPU를 최대 40% 더 배치할 수 있다고 제시했다. **전력 효율 개선이 장비 수요 감소를 의미하는 것이 아니라, 오히려 MW당 가속기와 연결 콘텐츠를 늘리는 구조다.** 가속기가 늘면 포트, 케이블, 광모듈도 함께 증가한다.

그림 3. 동일 전력 한도에서의 시스템 최적화 효과



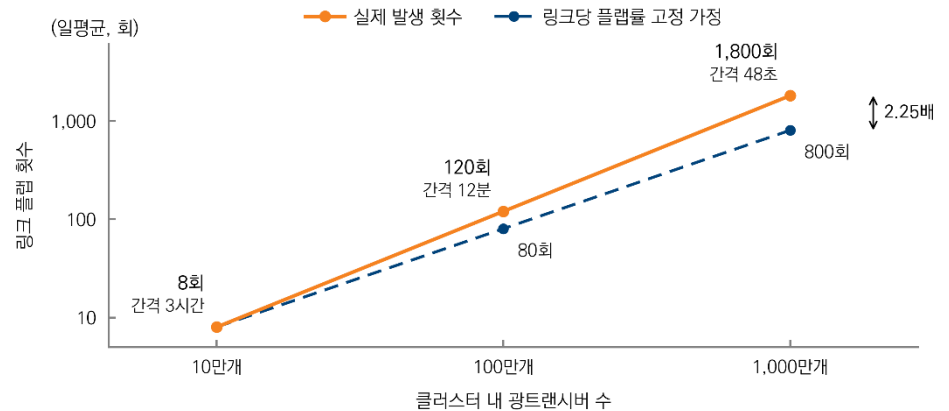
주: 엔비디아와 람다의 HGX B200 실증 기준. 기존 구성(16노드 풀전력)을 100으로 지수화했고, 적용 후는 19노드를 85% 전력으로 운영한 결과. 엔비디아는 차세대 랙에서 동일한 사이트 전력 안에 GPU를 최대 40% 더 수용할 수 있다고 조건부로 제시
 자료: 엔비디아, 미래에셋증권 리서치센터

에이전트형 AI는 이 압력을 메모리와 데이터 이동으로 확산시킨다. 엔비디아가 제시한 사례에서 2023년 챗봇형 작업은 입력 토큰 약 1천개, 대화 3회 수준이었던 반면 최근 에이전트형 작업은 평균 14.2만개의 입력 토큰과 65회의 반복 작업으로 늘었다. 긴 문맥과 KV 캐시를 저장하고 반복적으로 불러오는 부담이 급증하면서 **연산량보다 데이터를 어디에 두고 어떻게 이동시키느냐가 추론 원가를 좌우하기** 시작했다. SK하이닉스의 PIM과 HBF, 아스테라랩스의 외부 메모리 연결 구조도 같은 방향이다.

신뢰성의 가치도 높아지고 있다. 대형 클러스터의 광트랜시버 수는 2023년 10만개 이상에서 2025년 1,000만개 이상으로 늘었고, 같은 조건에서 링크 이상 발생 간격은 3시간에서 48초로 짧아졌다. 일부 고객에서는 광 링크 문제로 가동시간의 20~30%가 손실된 사례도 제시됐다. 광의 경쟁 기준도 포트 속도뿐 아니라 상태 모니터링, 장애 격리와 복구 능력으로 확대되고 있다. **동일한 대역폭이라면 더 높은 가동시간을 제공하는 제품의 경제적 가치가 높아질 수밖에 없다.**

결국 전력이 제한된 환경에서 증설의 단위는 더 이상 GW만이 아니다. MW당 가속기 수, 가속기당 연결 콘텐츠, 그리고 이를 실제 가동시간으로 전환하는 시스템 신뢰성이 다음 투자 변수가 되고 있다. AI 인프라의 투자 중심도 GPU 단품에서 메모리, 연결 반도체, 네트워크, 광을 포함한 시스템 전체로 확산될 것으로 판단한다.

그림 4. AI 클러스터 대형화와 광 연결 신뢰성 부담



자료: 크레도, 오라클, 미래에셋증권 리서치센터

[Investment Signal]

GPU보다 빠르게 커지는 가속기당 연결 콘텐츠

AI 인프라 투자에서 봐야 할 변수는 GPU 출하량 자체보다 가속기 한 대를 실제로 가동하기 위해 필요한 주변 BOM이 얼마나 빠르게 늘어나는가다. 수요를 단순히 데이터센터 GW와 GPU 수의 곱으로 보면 이 변화가 잘 드러나지 않는다. **향후에는 ‘가동 MW × MW당 가속기 수 × 가속기당 연결 콘텐츠’로 나눠 볼 필요가 있으며, 전력 제약이 심해질수록 뒤의 두 항이 더 빠르게 늘어난다.**

전력이 제약이 되면서 같은 MW 안에 더 많은 가속기를 배치하는 경쟁이 시작됐다. 가속기 수가 늘면 스위치 포트, 케이블, 광모듈도 함께 늘고, 스케일업 도메인이 커질수록 가속기당 연결 대역폭도 높아진다. 에이전트 AI는 여기에 메모리 용량과 데이터 이동 부담까지 추가한다. 결국 데이터센터 전력 증가율보다 연결 콘텐츠 증가율이 더 높아질 수 있는 구조다.

표 2. AI 인프라 구조 변화별 수혜 콘텐츠와 노출 기업

변화	늘어나는 콘텐츠	주요 노출 기업
같은 MW에 가속기 추가	스위치 포트, 케이블, 광모듈	아리스타, 시스코, 크레도
스케일업 도메인 확대	고속 SerDes, 스케일업 광과 구리 연결	브로드컴, 마벨, 크레도
메모리 계층화	PCIe, CXL, 리타이머, 메모리 컨트롤러	아스테라랩스, 크레도
클러스터 대형화	광 대역폭, 상태 진단, 장애 관리	루멘텀, 코히어런트, 코닝

주: 아리스타, 시스코, 크레도는 당사 커버리지이며 나머지는 관심종목

자료: 미래에셋증권 리서치센터

다만 모든 연결 영역이 같은 속도로 커지는 것은 아니다. 복제형 추론처럼 노드 간 동기화가 적은 워크로드에서는 스케일아웃 대역폭 수요가 상대적으로 완만할 수 있다. 반면 스케일업 연결과 메모리 연결의 중요도는 높아진다. 따라서 네트워크 투자에서도 단순한 대역폭 총량보다 어느 계층의 연결을 공급하는가가 종목 선별의 핵심이다.

프론티어 모델의 개발 속도가 일부 조절되더라도 컴퓨트 확보는 이어지고 있고, 전력 제약은 오히려 시스템 효율화 투자를 앞당기고 있다. 실적에서 확인할 시점은 2027년으로 판단한다. 대형 고객향 물량은 설치 후 검수 완료 시점에 매출로 인식되기 때문에, 2026년 하반기 구축분이 2027년 매출로 전환되고, 1.6T급 광 연결 전환과 ASIC 랙 증설이 겹치면서 가속기당 연결 콘텐츠 상승이 본격적으로 실적에 반영될 것으로 전망한다.

AI 인프라의 다음 투자포인트는 더 많은 GW가 아니라, 같은 MW에 더 많은 가속기를 넣고 가속기 한 대당 더 많은 연결 콘텐츠를 붙이는 구조적 변화다.

AI 인프라 포트폴리오

표 3. AI 인프라 밸류체인 테이블 (26.9.17 종가 기준)

(십억달러, 달러, %, 배)

구분	기업명	주가	시가총액	매출성장률		영업이익률		EPS 성장률		P/E		EV/EBITDA		수익률			
				26E	27E	26E	27E	26E	27E	12MF	24MF	12MF	24MF	5D	1M	3M	12M
서버 OEM	델	588.4	374.1	70.2	21.0	11.3	10.8	151.0	17.6	20.5	17.5	14.7	12.4	16.1	25.6	43.7	345.4
	HPE	61.0	81.0	35.8	15.4	14.7	15.0	96.5	20.5	13.6	12.0	8.6	7.6	10.5	9.6	28.7	141.8
	슈퍼마이크로	40.4	26.5	71.6	20.5	7.7	7.7	18.8	24.4	8.9	7.3	4.9	3.9	7.9	7.9	31.6	-12.2
네트워크 시스템	엔비디아	219.3	5,286.1	89.8	69.3	65.5	65.8	96.5	69.1	16.3	11.3	12.9	8.7	0.4	-0.2	4.1	24.5
	시스코	110.2	434.6	11.8	20.7	24.3	27.4	13.9	30.2	21.4	19.5	15.9	14.5	2.6	-1.2	-7.8	60.5
	아리스타	199.5	251.7	42.8	54.2	44.0	42.5	42.5	50.3	41.1	33.3	31.1	24.5	5.6	3.3	17.6	36.0
광/DCI (전송 /트랜시버)	코닝	147.8	127.3	16.6	19.4	21.8	23.5	30.5	32.0	36.6	27.1	19.9	15.3	-9.4	-7.6	-24.2	85.6
	루멘텀	893.6	80.2	108.3	52.9	41.1	43.6	148.0	57.7	36.9	24.3	24.7	15.2	-4.5	2.3	5.1	421.1
	노키아	10.6	61.0	4.4	7.6	10.9	12.7	18.6	20.6	23.4	20.1	14.2	11.6	0.1	3.9	-21.5	129.5
	코히어런트	296.0	58.0	50.2	37.1	23.7	26.0	69.2	51.0	28.0	19.4	18.8	12.7	1.0	-3.4	-24.0	173.9
	시에나	344.3	48.8	34.6	31.5	20.4	25.6	169.1	65.1	30.7	21.0	22.1	15.5	2.9	-15.1	-19.6	149.4
	파브리넷	380.9	13.7	30.9	19.1	11.1	11.3	26.3	20.7	20.5	16.7	16.2	12.8	-5.7	-21.1	-33.6	0.8
인터넷 반도체 (SerDes /DSP /리타이머)	브로드컴	347.3	1,657.9	65.9	64.1	66.9	65.4	71.1	64.5	18.9	11.8	14.7	8.9	-3.7	-8.6	-15.6	0.6
	마벨	240.8	211.1	46.1	50.9	37.2	39.1	47.6	61.1	41.4	26.2	32.6	20.9	6.1	11.5	-22.5	224.3
	아스테라 랩스	293.6	50.9	121.2	65.0	36.3	38.5	116.8	62.9	50.7	37.9	47.8	33.5	3.2	-3.3	-29.6	16.5
	크레도	168.3	31.6	109.4	65.8	30.4	36.7	94.9	64.8	22.0	15.4	18.1	11.6	5.0	-31.6	-38.1	-2.4
	맥컴	263.1	20.1	36.8	37.1	31.7	37.4	58.8	52.5	31.6	25.7	-	-	-3.4	-10.2	-32.8	102.8
	셈테크	178.2	16.6	41.9	31.7	27.3	36.0	103.7	70.5	35.3	23.7	26.5	18.6	12.0	31.8	12.6	191.4
화이트박스 /ODM	관타 컴퓨터	10.8	44.5	90.2	31.4	3.2	3.1	37.9	19.4	11.3	9.5	9.9	8.4	2.4	5.4	-8.4	23.3
	셀레스티카	57.7	32.4	59.9	32.6	14.9	15.4	80.8	36.7	17.1	12.3	12.2	8.6	-2.7	-17.5	-24.6	69.9
	엑톤 테크놀로지	329.9	41.8	65.9	72.5	8.3	8.7	85.3	75.6	19.1	13.5	13.8	8.9	1.5	6.2	-11.4	30.9

주: 아리스타, 시스코, 크레도는 당사 추정, 그 외는 블룸버그 컨센서스

자료: 블룸버그, 미래에셋증권 리서치센터

Compliance Notice

- 당사는 자료 작성일 현재 해당 회사와 관련하여 특별한 이해관계가 없음을 확인합니다.
- 당사는 본 자료를 제3자에게 사전 제공한 사실이 없습니다.
- 본 자료를 작성한 애널리스트는 자료작성일 현재 본 자료에서 매매를 권유한 금융투자상품 및 권리를 보유하고 있지 않습니다.
- 본 자료는 외부의 부당한 압력이나 간섭없이 애널리스트의 의견이 정확하게 반영되었음을 확인합니다.

본 조사분석자료는 당사의 리서치센터가 신뢰할 수 있는 자료 및 정보로부터 얻은 것이나, 당사가 그 정확성이나 완전성을 보장할 수 없으므로 투자자 자신의 판단과 책임하에 종목 선택이나 투자시기에 대한 최종 결정을 하시기 바랍니다. 따라서 본 조사분석자료는 어떠한 경우에도 고객의 증권투자 결과에 대한 법적 책임소재의 증빙자료로 사용될 수 없습니다. 본 조사분석자료의 지적재산권은 당사에 있으므로 당사의 허락 없이 무단 복제 및 배포할 수 없습니다.