



글로벌 투자전략-선진국
Analyst 황수욱
02. 6454-4896
soowook.hwang@meritz.co.kr



엔비디아 1Q25 실적발표

엔비디아 CY1Q25 매출액은 441억 달러로 전년동기대비 69% 성장하며 컨센서스(432.9억달러) 상회. 조정 GP 마진은 61%로 컨센서스(71%)를 하회했으나, H20 비용을 제외하면 71.3%로 사실상 컨센을 소폭 상회한 셈. CY2Q25 매출액 가이드를 450억 달러로 제시하며 컨센서스(455억 달러)를 하회했으나, 수출 통제에 따른 H20 매출 손실 80억 달러를 감안할 때, 중국향 매출 외 본업 성장 전망이 시장 예상을 상회했다는 점에서 고무적. 엔비디아 주가는 실적발표 이후 시간 외에서 4.7% 상승 중

AI 인프라 산업은 여전히 초과수요이며 초기 국면임이 다시 한 번 강조됨. 젠슨황은 생각하는 AI 모델(Reasoning model)이 예상보다 빠르게 AI 하드웨어 수요를 키우고 있으며, 이전보다 더 많은 수주 잔고가 쌓이고 있음을 언급. 미국이 최대 시장이었지만, AI도 전기와 같은 '인프라'이기 때문에, 전세계로 수요처가 확산될 수 있음을 언급

이런 AI 인프라 산업의 위치에서 엔비디아의 전략은 '엔비디아를 선택할 수 밖에 없게끔 하는 것'. 젠슨 황 CEO의 코멘트를 종합하면, 이 부분에 초점을 맞춰 전반적인 비즈니스 전략을 채택하고 있다는 점이 확인됨. 수출 통제에 대한 생각을 말하는 부분에서도 중국의 역량을 낮게 보는 통제 정책보다는, 결국 전세계로 AI 팩토리가 확산되는 국면에서 시장에서 미국의 AI 인프라가 선택되는 것이 중요하다고 강조

그래서 엔비디아는 계속 자신들의 생태계를 구축하고, 자신들의 인프라 안에서 기업과 사람들이 AI '산업'에 참여하게끔 하는 전략을 펼치는 중. GPU를 넘어 네트워크(NVLink), 소프트웨어(CUDA, Nemotron)까지 직접 서비스, 이번 실적에서는 강조되지 않았지만, Physical AI까지 이어지는 옴니버스까지 엔비디아 생태계를 구성하고, 그 안에서 많은 기업과 계속 '협업'을 하이라이트하는 이유일 것

컨퍼런스 콜 주요 내용과 Q&A

Colette Kress

AI 워크로드가 빠르게 추론(Inference)로 전환되고 AI 팩토리 구축 진행되면서 큰 폭의 성장 견인. 고객 수요는 매우 확실하게 이어지는 중. 미국 정부가 중국 전용으로 설계된 H20에 대한 새로운 수출 통제 발표, 1분기 46억 달러 H20 매출 인식, 이는 4월 9일 이전에 이미 발생된 매출. 동시에 약 45억 달러의 재고 평가 손실 및 구매약정 충당금 반영. 45억 달러의 손실은 초기 예상보다는 작았는데, 일부 자재를 다른 곳에 재활용할 수 있었기 때문

우리는 향후 500억 달러 규모로 성장할 것으로 보이는 중국 AI 가속기 시장에서 사실상 배제될 위험에 처해있으며, 이는 당사 비즈니스에 큰 부정적 영향을 줄 것. 반면, 중국 및 글로벌 시장에서 중국 업체나 다른 국가 업체에게 기회를 주는 결과가 될 수도 있음

GV200 NVL 랙은 이제 광범위하게 공급되는 중. 마이크로소프트는 이미 수 만대의 블랙웰 GPU 도입, OpenAI도 주요 고객 중 하나로 수십만대 규모의 GB200 도입을 목표로 하고 있음

GB200 노하우 덕분에 향후 로드맵의 다음 단계인 블랙웰 울트라로의 전환도 부드러울 것으로 예상. GB300은 이미 이달 초에 CSP에 샘플 출하 시작, 분기 후반에는 본격적인 양산/출하 시작할 계획. HBM 메모리를 50% 늘린 GB300 GPU는 FP4 추론 성능에서도 GB200 대비 추가로 50% 이상 향상된 밀도(density)를 제공

현재 우리는 추론(inference) 수요가 가파르게 증가하는 것을 목도. 오픈AI, 마이크로소프트, 구글 등은 토큰 생성량이 크게 늘었음. 마이크로소프트 Azure의 경우, 1분기에 100조개 이상의 토큰을 처리하며, 전년 대비 5배 이상 증가. 이는 Azure OpenAI 뿐만 아니라, 애저 AI 파운드리(Azure AI Foundry)와 마이크로소프트 전체 AI 서비스에 대한 강력한 수요를 반영

추론 스타트업들도 이제 B200 기반으로 모델을 서빙(Serving)하면서, 특히 reasoning(추론) 모델인 DeepSeek R1 같은 곳에서 토큰 생성률과 매출이 3배 증가했다고, Artificial Analysis 지표가 보여주고 있음

'NVIDIA NEMO'가 블랙웰 NVL72 기반으로 작동되면, 최근 업계에서 주목받는 reasoning 모델에 대해 추론 처리량을 최대 30배 개선할 수 있음. 개발자 참여도 높아져, LLM 제공사인 Perplexity부터, 예컨대 Capital One 같은 금융사들까지, NEMO를 활용해 AI 챗봇 지연 시간을 5배 줄이는 등 성능 개선을 보고 있음

NVIDIA의 프로그래머블(Programmable) CUDA 아키텍처와 방대한 생태계가 주는 장점. 최근 MLPerf 추론(Inference) 결과에서, 우리는 GB200 NVL72를 처음으로 제출했고, 8GPU 구성의 H200 대비 최대 30배 높은 추론 처리량을 LLAMA 3.1 벤치마크에서 달성. 이는 GPU당 성능을 3배 높인 것과 동시에, 하나의 NVLink 도메인에서 9배 많은 GPU를 연결한 결과. 블랙웰 아직 초기 단계이지만, 소프트웨어 최적화를 통해 지난 한 달 동안만 1.5배의 성능 개선을 이룸

생성형 AI에서 나아가, 지각(Perception), 추론(Reasoning), 계획(Planning), 행동(Acting)을 수행하는 '에이전틱(agentic) AI' 시대로 넘어가면서, 모든 산업, 기업, 국가가 변화를 맞이하고 있음. 우리는 AI 에이전트를 디지털 인력으로 보고 있으며, 고객서비스부터 복잡한 의사결정까지 수행할 수 있다고 봄

우리는 'Llama Nemotron'이라는 오픈소스 추론 모델 제품군을 발표했는데, 이는 엔터프라이즈 환경에서 에이전틱 AI 플랫폼을 크게 강화하는 역할. Meta의 Llama 아키텍처 기반으로 개발되었으며, 다양한 배포 니즈에 맞춰 여러 사이즈로 제공. 포스트 트레이닝(사후 훈련) 최적화를 거쳐 정확도가 20% 향상되고, 추론 속도도 5배 빨라짐. Accenture, Cadence, Deloitte, Microsoft 등 주요 플랫폼 업체들이 자사 업무에 당사의 reasoning 모델을 활용해 변혁을 시도

NVIDIA NEMO 마이크로서비스는 업계 전반에서 개발-최적화-확장 과정에 쓰이고 있음. 예를 들어 Cisco는 NEMO로 자사의 코드 보조(Code Assistant) 도구 정확도를 40% 높이고 응답 시간을 10배 단축. NASDAQ 역시 AI 플랫폼 검색 기능에서 정확도가 30% 개선되고, Shell이 구축한 커스텀 LLM도 30% 높은 정확도를 달성

보안 분야에서도, Checkpoint, CloudStrike, Palo Alto Networks 등 주요 업체들이 NVIDIA 보안 소프트웨어 스택을 기반으로 AI 보안 및 에이전틱 워크플로우를 구축. 예컨대 CrowdStrike는 50% 더 적은 컴퓨팅 비용으로 탐지 트리아지(Detection Triage) 속도를 2배 높임

다음은 네트워킹 부문. 1분기 네트워킹 매출은 50억 달러로 전 분기 대비 64% 증가했. 고객들은 대규모 AI 팩토리 워크로드를 확장하기 위해 당사의 네트워킹 플랫폼을 적극 채택

우리는 'NVLink'라는 세계에서 가장 빠른 스위치를 개발. 스케일 업(scale-up) 측면에서, 5세대 NVLink Compute Fabric은 PCIe Gen5 대비 14배 높은 대역폭을 제공. NVLink 72는 랙 1대당 130TB/s의 대역폭을 제공하며, 이는 전 세계 인터넷 트래픽 피크 수준 전체와 맞먹음

NVLink는 새로운 성장 동력이며, 1분기 출하량이 10억 달러를 넘어서며 순조롭게 시작. MediaTek, Marvell, Alchip, Astera Labs 같은 ASIC 설계업체나, 후지쯔(Fujitsu), 퀄컴(Qualcomm) 같은 CPU 업체들과 협력하여 NVLink Fusion으로 양측 생태계를 연결하고 있음

스케일 아웃(scale-out) 측면에서는, 업그레이드된 Ethernet 제품군이 AI 처리에 최적화된 최고 대역폭과 최저 지연을 제공. 스펙트럼-X(Spectrum-X) 제품군은 전 분기 대비, 전년 동기 대비 모두 높은 성장세를 보이며 현재 연간 80억 달러 규모로 확대. 코어워브, MS 애저, 오라클 클라우드, xAI 등이 고객사로 있던 가운데, 이번 분기에 구글 클라우드와 메타도 스펙트럼-X를 채택

Jensen Huang

자주 받는 질문에 대한 생각을 말하면,

1. 수출 통제 관련: 중국은 세계 최대 AI 시장 중 하나이며, 글로벌 성공의 발판. 전 세계 AI 연구자의 절반가량이 중국에 있음. 중국 시장에서 플랫폼 우위를 잡은 기업은 세계적으로도 우위를 확보하게 됨. 하지만 현재 미국 기업 입장에서는 500억 달러 규모의 중국 시장이 사실상 막혀 있는 상태. H20 수출 금지 조치로 인해, 중국에서 호퍼(Hopper) 아키텍처 제품을 더 이상 판매할 수 없게 되었음. 규제에 맞추기 위해 호퍼 성능을 추가로 낮추는 것은 실질적으로 불가능. 그 결과, 팔지 못하게 된 재고를 수십억 달러 단위로 상각 처리

중국은 미국 칩 없이도 이미 상당한 AI 역량을 확보. AI 활용이 이미 상당히 보편화되어 있고, 선진 기술을 도입할 여지가 큼. 결국 문제는 '중국이 AI를 가질 것인가'가 아니라, '중국의 거대한 AI 시장에서 미국 플랫폼이 쓰일 것인가 아닌가'임. 규제가 중국 경쟁사를 보호해 주는 꼴이 되어, 이들이 해외 시장에서 더 커지고 미국의 위치가 상대적으로 약해질 수 있음. AI 경쟁은 단순히 칩이 아니라, 전 세계가 어떤 스택(stack)을 채택하느냐에 달려 있고, 6G와 양자(Quantum) 시대로 확장될 미래 인프라 리더십이 걸린 문제

정부는 "중국이 AI 칩을 만들기 어려울 것"이라는 가정에 기반해 정책을 펴왔는데, 이는 사실상 이미 틀렸음. 중국 막대한 제조 역량을 보유하고 있음. 결국, AI 개발자가 어디 스택에서 개발하느냐가 중요. 미국은 AI 기술 수출을 규제하기보다는, 미국 플랫폼이 전 세계 개발자에게 '선택'받도록 하는 방향으로 접근해야 함

2. 딥시크: 중국의 DeepSeek과 Qwen은 세계 최고의 오픈소스 AI 모델 중 하나. 무료로 공개되어 이미 미국, 유럽 등지에서도 널리 쓰임. ChatGPT처럼, DeepSeek R1도 토큰을 더 많이 쓸수록(추론 단계가 길어질수록) 점점 더 나은 답변을 생성하는 "Reasoning AI" 개념을 선보임. Reasoning AI는 여러 단계를 단계별로 사고해 나가는 방식이어서, 이전의 원샷 추론 대비 몇 백~수천 배 많은 토큰 연산을 필요로 함. 이는 추론 규모 자체가 엄청나게 커졌음을 의미

DeepSeek 사례는 오픈소스 AI의 전략적 가치를 잘 보여줌. 선도적인 모델이 미국 플랫폼 위에서 개발되고 최적화될 때, 해당 플랫폼 사용이 확대되고 피드백 루프가 형성되며, 미국 생태계를 더욱 견고히 만듦. 미국 플랫폼이 오픈소스 개발자에게 최고의 경험을 제공하는 것이 중요. DeepSeek과 Qwen 같은 모델이 "미국 인프라 위에서 최고의 성능"을 내도록 해야 함

3. 온쇼어(Onshore) 제조 관련: 트럼프 대통령은 첨단 제조업을 자국으로 유치하는 "리쇼어링(Reshoring)" 정책을 강하게 추진 중이며, 이는 일자리 창출과 국가 안보 측면에서 큰 의미. 미래의 공장은 강력한 컴퓨팅과 로보틱스가 핵심이 될 것. TSMC는 애리조나에 6개의 펩과 2개의 첨단 패키징 공장을 건설 중이고, 이미 공정 인증을 진행 중이며 연말 생산 개시가 예상. 스피ل(SPIL), 암코(Amkor) 역시 애리조나에 패키징·조립·테스트 라인을 건설하고 있음

우리는 이러한 투자를 촉진하기 위해 장기 매입 약정을 체결하며, 미국 AI 제조 인프라의 미래에 크게 투자하고 있음. 목표는 "칩부터 슈퍼컴퓨터까지 모두 1년 안에 미국에서"라는 것

GB200 NVLink72 랙 한 대에는 120만 개의 부품이 들어가며 무게만 2톤. 이런 규모의 슈퍼컴퓨터를 생산한 전례가 없었음. 그러나 우리는 폭스콘(Foxconn)과 휴스턴에서 100만 평방피트 규모 공장을, 위스트론(Wistron)과 텍사스 포트워스 공장을 함께 건설하며, 공급망 역량을 계속 확장하고 있음

4. AI 확산 규정(AI Diffusion Rule)에 대해: 트럼프 대통령은 기존의 AI 확산 규정을 '비생산적'이라 비판하며 철회(AI Diffusion Rule Rescind). 대신 동맹국과의 협력을 통해 미국 AI 기술의 세계적 리더십을 확대하겠다는 방침. 중동 순방에서 사우디아라비아와 UAE에 수백 메가와트~수 기가와트 규모의 AI 인프라 투자를 발표했는데, 저도 그 자리에 함께 했음. 이러한 협력은 일자리 창출, 인프라 발전, 세수 증대, 대미 무역적자 감소 등의 긍정 효과로 이어짐

미국은 여전히 NVIDIA의 최대 시장이며, 가장 방대한 인프라 설치 기반을 갖추고 있음. 그러나 이제 거의 모든 국가가 AI를 산업혁명 핵심으로 여기며, AI 인프라를 국가적 차원에서 구축하려고 함. 컴퓨텍스에서 우리는 대만 정부, 폭스콘과 함께 대만 최초의 AI 팩토리를 발표. 그 직후에는 스웨덴의 국가 AI 인프라 구축에 협력하기 위해 현지를 방문. 일본, 한국, 인도, 캐나다, 프랑스, 영국, 독일, 이탈리아, 스페인 등도 국가 차원의 AI 팩토리를 건설 중이거나 계획 중. 소버린 AI(Sovereign AI)가 NVIDIA의 새로운 성장 동력으로 부상

실적 컨콜 Q&A

Q1. Reasoning AI에 따른 추론 수요 급증의 현실화. 이런 수요를 현재 어느정도 충족 한 것으로 보는지? 향후 Reasoning AI에서는 꼭 NVL72 같은 랙 스케일 솔루션이 꼭 필요한건지?

A: 가능한 수요를 만족시키고 싶고, 실제로 대부분 감당할 수 있을 듯. GB NVL72는 Reasoning AI를 위한 최적 컴퓨팅 엔진. Reasoning AI는 단순 챗봇보다 훨씬 더 많은 토큰을 처리. 생각을 오래 할수록(토큰이 많이 들어갈수록) 답이 좋아지고, "더 스마트해지는" 구조

이를 효율적으로 처리하면서도 응답 속도(QoS)를 높이기 위해서는 성능이 극도로 중요. 예컨대 Grace Blackwell NVLink72는 호퍼(Hopper) 대비 추론 처리량이 40배나 높아, 비용을 대폭 낮추면서 서비스 품질을 보장. 이를 위해 슈퍼컴퓨터 전체 설계를 새로 해야 했지만, 이제 완전 양산 체제로 접어들었음. 앞으로 Reasoning AI 수요가 폭발할 것이며, 이를 효과적으로 지원할 수 있을 것

Q2. 2분기 중국 매출 감소분 관련, 앞으로 반영될 남은 부분이 있는지? 올해 GTC에서 앞으로 수 년간 AI 분야에서 약 1조 달러가 투자될 것이라고 함. 지금 이 정도에 어느 정도 지점에 와 있는지?

A: H20 관련, 원래 70억 달러 가까이 매출이 될 수 있었는데, 실제로는 46억 달러에 그침. 수출 통제로 25억 달러 손실. 2분기에는 H20으로 80억 달러 정도의 매출을 예상했으나, 이제는 불가능해 짐. 실제 시장 손실은 더 큼. 우리는 약 500억달러의 미래 중국 AI 시장에 접근할 수 없게 된 위험

지능(intelligence)을 디지털화한 AI도 필수 인프라가 될 것. 지금은 그 초기 단계. 인터넷 인프라도 그랬듯, AI 인프라도 모든 국가와 기업이 구축해야 함. 게다가 AI 사용 사례는 더 많아질 것

예를 들어, 엔터프라이즈 데이터가 온프레미스인 경우도 많아, 온프레미스 AI 인프라 수요도 큼. 통신도 마찬가지로 6G는 AI로 구동될 것. 제조 공장도 마찬가지고, 모든 자동차 회사도 AI 팩토리를 갖춰야 함. 결론적으로 아주 초반 단계이므로 장기적으로 볼 때 우리가 제시한 1조달러 시장 전망 충분히 현실성 있다 생각

Q3. 최근 중동, Oracle, xAI 등 GPU 클러스터 대규모 투자 소식이 많음. 아마도 같은 규모의 다른 사례들도 있을 것 같은데, 이런 초대형 주문이 더 남아 있는지? 현재 블랙웰 수요가 공급을 얼마나 초과하고 있는지, 리드 타임은 어느정도인지?

A: 주문은 이전보다 많이 쌓여 있음. 동시에 당사도 공급망 확충 중. 전세계적으로, 그리고 미국 내에서도 생산 역량 늘리는 중. 아직 발표되지 않은 대규모 프로젝트가 많음. AI는 전기나 인터넷처럼 기본 인프라이기 때문에 모든 국가, 모든 산업에서 구축 필요성 대두. 국가별 AI 팩토리를 구축해야 하며 여기에는 엄청난 서버 및 랙이 필요. 우리는 그 초기 국면

AI가 reasoning 능력을 갖추며, 이전 대비 추론 컴퓨팅이 기하급수적으로 늘었고, 이는 더 많은 GPU, 더 많은 네트워크 수요를 의미. 그리고 이제 막 시작 단계. 추후에도 대규모 발표가 잇따를 것

Q4. 2분기 실적 가이드스와 관련하여, 중국 H20가 80억 달러 빠지는데도 450억 달러 가이드스를 제시했다면, 나머지 사업부가 기존 예상보다 약 20~30억 달러가 증가했다는 뜻. 그 주된 요인은 무엇이라 보시는지?

AI 확산 규정 철회(AI Diffusion Rule Rescind)와 소버린 프로젝트 호재 등을 볼 때, 연간으로도 분기별 성장세가 이어질 가능성이 있다고 보는지?

A: 2분기에 H20 관련 80억 달러가 빠졌는데, 그럼에도 불구하고 450억 달러를 가이드스로 잡았음. 이는 블랙웰 수요가 매우 강하기 때문이고, 공급망도 확대되어 더 많은 매출을 소화할 수 있게 되었음

GTC 때 전망했던 것보다 상황이 더 좋아진 부분이 네 가지 있음.

1) Reasoning 수요가 예상을 뛰어넘는 속도로 증가, 답변 품질이 좋아지고 에이전트 AI가 문제 해결 능력을 보이면서 추론 토큰 양이 폭발적으로 증가

2) AI Diffusion Rule 철회, 트럼프 대통령은 미국이 승리하기 위해 세계가 미국 스택에 올라타야 한다고 보고 있음. 수많은 국가가 AI를 국가 인프라로 채택하려고 할 때, 미국 플랫폼이 그 선택지가 되는 것이 중요

3) 엔터프라이즈 AI, 에이전트는 단순 생성형 AI보다 훨씬 더 강력함. 톨과 메모리를 활용해 업무 해결. RTX Pro Enterprise Server 같은 통합 솔루션을 발표했고, 대형 IT 업체들이 모두 협력 중

4) 산업 AI. 세계 각지에서 제조 공장, 칩 공장 등이 새로 지어지고 있고, 여기에 AI와 로봇틱스가 본격 도입될 것. 모든 공장에 AI 팩토리가 붙게 될 것이고, 이는 더 많은 AI 교육 및 추론 수요를 의미

Q5. 1) 이미 중국용 성능 조정 SKU를 정부 승인 받았고, 현재 생산은 하고 있지만 2분기에는 출하가 불가능한 상태인지. 2) 또 엔비디아가 분기당 중국 매출이 70~80억 달러 수준이었는데, 조만간 수출 가능 SKU가 나오면 그 정도 수치로 다시 복귀 가능할지 저희가 모델링을 어떻게 해야 할지 궁금합니다. 가능한 범위에서 답변 부탁드립니다

A: 수출 통제에 대해서는 여러 제한이 있는데, 이번 개정된 규정은 호퍼(Hopper) 제품을 지금보다 더 낮은 성능으로 조정해서라도 계속 쓰게 하는 것을 사실상 불가능하게 만들었습니다. 즉, 호퍼 라인업으로 중국을 지원할 수 있는 여지는 거의 끝났다고 보면 됩니다

우리는 선택지가 많지 않음. 우선 규제 한도를 명확히 이해하는 것이 핵심이고, 그 범위 내에서 중국 시장에 맞출 수 있는 어떤 흥미로운 제품을 개발할 여지가 있는지 살피고 있음. 현재 확정된 사항은 없고 내부적으로 고민 중. 다만 지금 규제는 상당히 엄격하여 오늘 발표할 만한 것은 없습니다. 혹여 그 시점이 온다면, 정부와 충분히 논의할 계획

Q6. 네트워킹 사업부문 실적과 관련하여 조금 더 자세히 듣고 싶음. 특히, CSP들을 중심으로 이더넷 제품이 잘 채택되고 있다고 하셨는데, 네트워크 부문 매출이 어떻게 성장하고 있는지, 그리고 네트워크 attach rate에 변화가 있는지

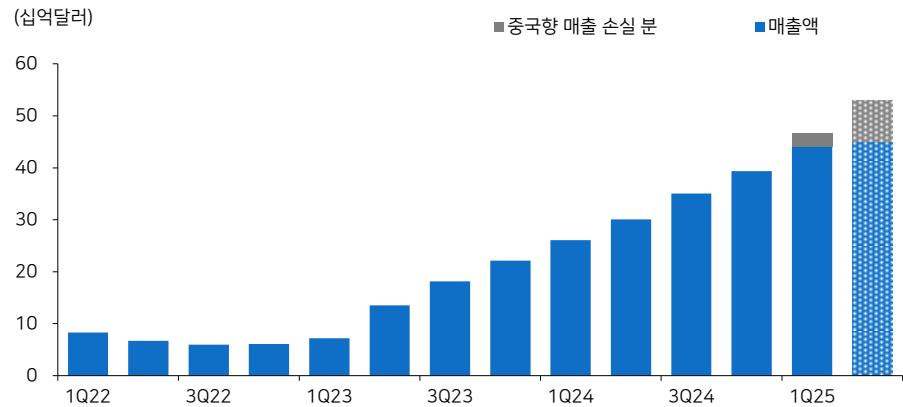
A: 1) Scale-Up 플랫폼: 개별 컴퓨팅 노드를 더욱 큰 '하나의 컴퓨터'처럼 엮어주기 위한 플랫폼. 사실 Scale-Out보다 Scale-Up이 훨씬 까다로움. 이 플랫폼을 NVLink라고 부르는데, 칩과 스위치, NVLink spine 등으로 구성

2) InfiniBand와 Spectrum-X: 기본적으로 이더넷은 독립적인 트래픽을 처리하는데 설계. 그러나 AI 클러스터는 다수 노드가 함께 작업해야 하고, 트래픽이 폭발적으로 발생할 때 지연 시간이 크게 좌우됨. 왜냐하면 AI는 "생각"하는 과정에서 가능한 한 빨리 데이터를 주고받기를 원하기 때문

그래서 우리는 이더넷에 초저지연, 혼잡 제어(congestion control), 적응형 라우팅(adaptive routing) 등 기존에는 인피니밴드에서만 가능하던 기능을 추가. 그 결과, AI 클러스터 내 이더넷 자원 활용 효율이 크게 늘었음. 일반적으로 기존에는 50% 수준이던 것이 85~90%까지 오름

3) BlueField: 이를 "컨트롤 플레인(control plane)"이라 부름. 스토리지 및 보안, 멀티 테넌트 환경에서 사용되며, 사용자 간 워크로드를 분리하면서도 고성능 배어메탈 성능을 제공. 클라우드나 대형 AI 클러스터에서 매우 유용하게 쓰이고 있음

그림1 엔비디아 매출액 및 성장률, 가이던스. 중국향 매출액 손실분



자료: Bloomberg, NVIDIA, 메리츠증권 리서치센터

이상은 2025년 5월 28일(현지시간) 발표된 엔비디아 실적 발표 및 컨퍼런스 콜 Q&A를 정리한 내용임

Compliance Notice

본 조사분석자료는 제 3자에게 사전 제공된 사실이 없습니다. 당사는 자료작성일 현재 본 조사분석자료에 언급된 종목의 지분을 1% 이상 보유하고 있지 않습니다. 본 자료를 작성한 애널리스트는 자료작성일 현재 추천 종목과 재산적 이해관계가 없습니다.

본 자료에 게재된 내용은 본인의 의견을 정확하게 반영하고 있으며, 외부의 부당한 압력이나 간섭 없이 신의 성실하게 작성되었음을 확인합니다.

본 자료는 투자자들의 투자판단에 참고가 되는 정보제공을 목적으로 배포되는 자료입니다. 본 자료에 수록된 내용은 당사 리서치센터의 추정치로서 오차가 발생할 수 있으며 정확성이나 완벽성은 보장하지 않습니다. 본 자료를 이용하시는 분은 본 자료와 관련한 투자의 최종 결정은 자신의 판단으로 하시기 바랍니다. 따라서 어떠한 경우에도 본 자료는 투자 결과와 관련한 법적 책임소재의 증빙자료로 사용될 수 없습니다. 본 조사분석자료는 당사 고객에 한하여 배포되는 자료로 당사의 허락 없이 복사, 대여, 배포 될 수 없습니다.