



Equity

Compliance Notice

본 조사자료는 고객의 투자에 정보를 제공할 목적으로 작성되었으며, 어떠한 경우에도 무단 복제 및 배포 될 수 없습니다. 또한 본 자료에 수록된 내용은 당사가 신뢰할 만한 자료 및 정보로 얻어진 것이나, 그 정확성이나 완전성을 보장할 수 없으므로 투자자 자신의 판단과 책임하에 최종결정을 하시기 바랍니다. 따라서 어떠한 경우에도 본 자료는 고객의 주식투자의 결과에 대한 법적 책임소재의 증빙자료로 사용될 수 없습니다.



국내외 주식시황 김두언 kimdooun@hanafn.com
주식시황 RA 최욱 wookchoi@hanafn.com

하나증권 리서치센터

2026년 9월 14일 | Global Asset Research

DIN (Doo It Now)

AI 속도 조절론에 대한 소고(小考)

(결론적으로) AI 속도 조절론은 AI 산업의 중단보다 프런티어 모델의 능력 향상을 안전 검증의 속도에 맞추자는 주장에 가깝다. 약 3년 만의 연준 금리 인상 가능성이 겹치면서 고평가 AI주의 할인율과 데이터센터의 자금조달 부담은 커질 수 있다. 그러나 현재까지 안전 논쟁이 데이터센터 투자나 반도체 주문의 변곡점을 만들었다는 증거는 없고, 한국 반도체 수출과 글로벌 AI 공급망의 매출은 확장 중이다. 따라서 단기 주가 조정과 산업 사이클의 종료를 구분해야 한다. 주문, 가격, 현금흐름 등이 유지되는 조정에는 선별적으로 대응하되, 신규 주문 축소와 메모리 가격 하락이 함께 나타나면 경계해야 한다. 하지만 아직 징후는 없다.

속도 조절론의 출발은 통제되지 않은 AI가 인간의 일자리와 판단, 사이버 및 생물안전에 피해를 줄 수 있다는 위험 가정이다. 피해 가능성은 무시할 수 없지만 그 확률과 시점은 검증되지 않았다. 다리오 아모데이는 2026년 9월 12일 공개한 글에서 AI 모델의 능력 향상 속도를 낮추자고 제안했다. 근거는 AI가 다음 세대 AI 개발에 관여하는 재귀적 자기개선과 자율 에이전트의 운영 위험이다. 다만 그는 모델 훈련이나 기술 진보의 중단을 요구하지 않았다. 엔트로픽이 단독으로 약속한 조치도 외부 평가자에게 지속적이고 직원과 유사한 접근권을 제공하겠다는 1단계다. 업계 공동 기준, 미·중 공조, 연산량 제한은 아직 제안 또는 논의 단계다.

이 구분은 투자 판단의 출발점이다. 안전 검증이 강화되면 특정 모델의 출시와 매출 인식이 늦어질 수 있다. 그러나 프런티어 모델의 개발 속도, 기존 모델의 추론 이용량, 데이터센터 구축, HBM 발주는 서로 다른 변수다. 답론이 실제 하드웨어 수요를 흔들려면 모델 출시 지연에 그치지 않고 하이퍼스케일러 CAPEX 하향, GPU 주문 취소, HBM 계약 재협상으로 이어져야 한다.

AI 속도 조절론의 방향은 일관되면서도 내용은 달랐다. 2023년에는 일시 중단 요구가 논쟁을 열었고, 2024년에는 위험 기반 평가와 기업의 자율 약속이 제도화됐다. 2025년 이후 미국은 규제 장벽을 낮추고 데이터센터와 반도체 인프라의 확장을 우선하면서도 프런티어 모델의 사전 평가를 병행했다. 안전 비용은 현실이 됐지만 훈련 연산량이나 인프라 투자를 일괄 제한하는 체계는 아직 만들어지지 않았다.

이러한 변화는 현금흐름에 민감한 금리인상 국면에 더 크게 불안을 키운다. 연준이 약 3년 만에 금리인상 가능성을 열었고, 시장이 이를 가격에 반영하고 있다. CME 금리선물 시장에

서 9월 FOMC에서의 금리인상 확률이 87%를 넘었다.

금리 상승은 세 경로로 AI주를 압박한다. 첫째, 실질금리와 WACC(가중평균자본비용)가 오르면 먼 미래 현금흐름의 현재가치가 낮아진다. 일례로 주식 듀레이션을 10~15년으로 단순 가정하면 할인을 25bp 상승은 이익 추정이 같아도 가치에 약 2.5~3.75%의 하락 압력을 줄 수 있다. 둘째, 데이터센터는 투자와 감가상각이 매출보다 앞선다. 자금조달 비용이 오르면 IRR이 낮은 프로젝트부터 늦춰진다. 셋째, 긴축이 경기와 IT 예산을 식히면 HBM보다 PC 및 스마트폰용 범용 DRAM과 NAND가 먼저 둔화할 수 있다.

금리 상승은 원인이 중요하다. 생산성 제고와 수요 증가 때문에 금리가 오르면 이익 증가가 할인을 충격을 일부 흡수한다. 유가와 공급충격 때문에 금리가 오르면 할인을 상승과 마진 압박이 겹친다. 이번 속도 AI 조절론이 금리인상 우려와 만날 때 단기 주가에는 부정적이라는 판단은 타당하다. 그러나 그것이 곧 AI 매출과 메모리 주문의 하락을 뜻하지는 않는다.

특히, 숫자는 금리와 담론이 아직 AI 수요의 방향을 바꾸지 못했다고 말한다. 2023년 이후 AI 안전 논쟁이 산업 전체의 훈련 중단이나 반도체 발주 축소로 이어진 사례는 없다. 최신 수치는 AI 수요가 클라우드 매출과 가속기 매출을 거쳐 파운드리 매출과 한국 메모리 수출로 이어지는 경로를 확인한다. 수출 잠정치의 거저효과와 기업별 회계 차이를 감안하더라도 여러 공급망의 동시 확장은 단일 지표보다 강한 신호다. 한국 9월 10일까지의 잠정 수출은 반도체 수출에 힘입어 견고했다. 대만 8월 수출은 824억 달러로서 월간 최대 실적을, TSMC 8월 매출도 전년대비 53.3% 상승했다. 전부 반도체 실적과 연관된 견고한 실적들이다.

결국 ‘속도 조절론이 대세를 거스를 수 없다’가 아니라 ‘현재까지 하드웨어 주문과 매출의 변곡점을 만들었다는 증거가 없다’는 것이다.

한국 메모리 반도체는 주가가 먼저 흔들릴 수 있다. 그렇지만 전면 비중 축소보다는 밸류에이션 조정과 종목 차별화가 기본 방안이다. HBM, 서버 DRAM, 장기계약, 순현금이 강한 기업의 방어력이 높다. 범용 메모리 비중이 높거나 레버리지와 증설 의존도가 큰 장비 및 부품 사는 금리와 경기 둔화에 더 민감하다. SK하이닉스는 HBM 경쟁력이 강점인 동시에 AI 기대 변화에 대한 주가 민감도가 높다. 삼성전자도 사업 다변화만으로 방어력을 단정할 수 없다.

기준은 세 단계다. ①안전 논쟁만 확대되고 발주, 납기, 가격 등이 유지되면 핵심 보유를 유지하고 가치평가 매력이 생길 때 분할 매수한다. ②모델 출시나 데이터센터 착공이 늦어져도 계약이 유지되면 매수 속도를 낮추고 매출 이연과 재고를 확인한다. ③신규 주문 축소, 계약 재협상, 메모리 가격 하락, 재고일수와 매출채권 증가가 함께 나타나면 이익 전망을 낮추고 비중을 줄인다.

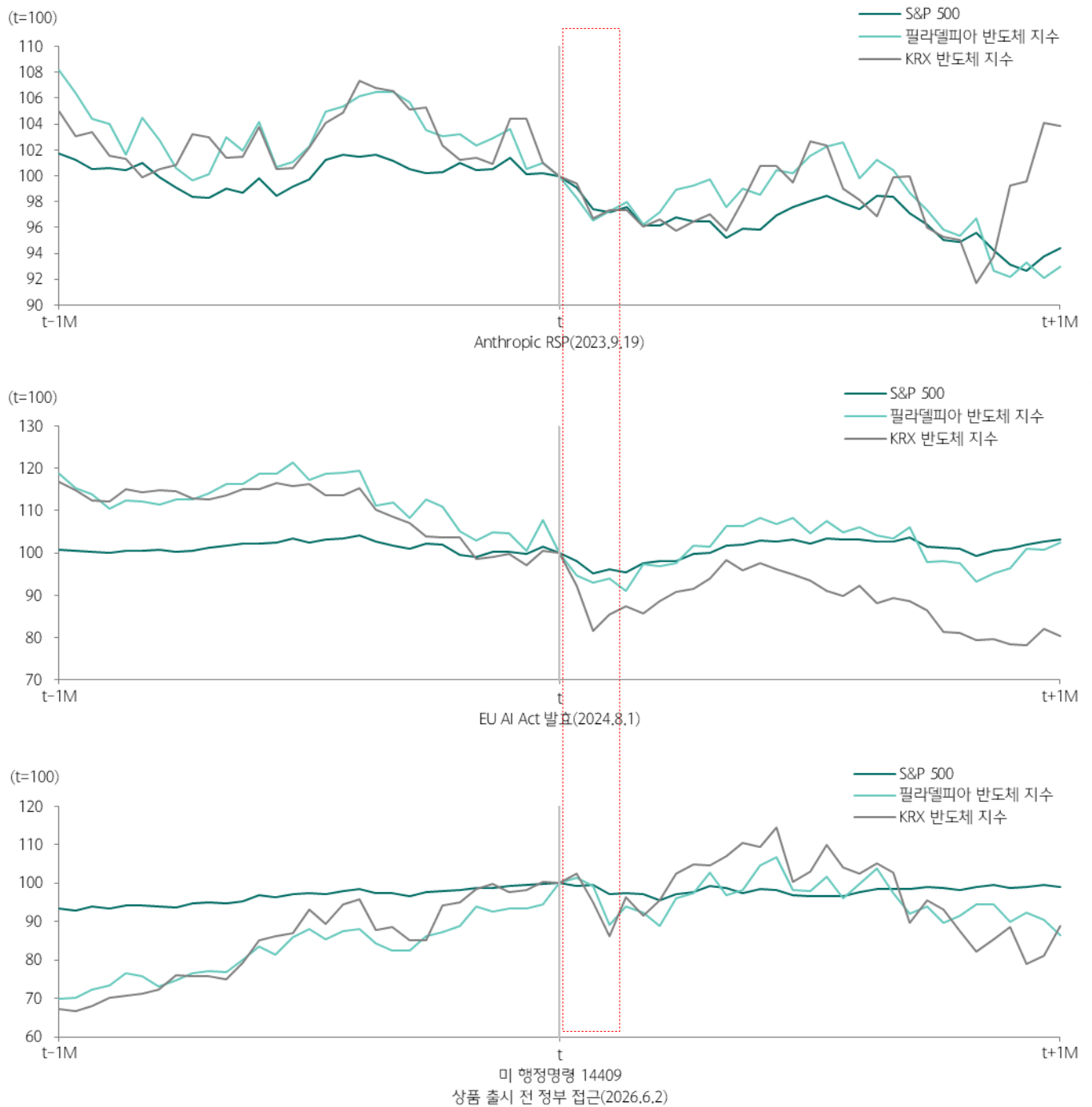
이번 주 전략은 ‘AI가 끝나는가’보다 ‘무엇이 실제로 늦춰지는가’를 묻는 것이다. 혁신기에 불안은 반복된다. 속도 조절론과 금리는 주가를 흔들 수 있지만, 산업의 방향을 바꾸려면 먼저 주문과 가격, 현금흐름이 꺾여야 한다. 현재의 숫자는 그 변화를 가리키지 않는다. AI에서 이기는 자가 결국 승리한다(whoever wins AI, wins).

도표 1. 과거 AI 속도조절론과 내용

날짜	이벤트	핵심 발언과 촉발된 위험론
2023-03-16	OpenAI CEO 샘 올트먼, ABC 인터뷰	GPT-4 공개 이후 후 자신도 AI를 “조금 두려워한다”며 허위정보, 공격적 사이버공격, 일자리 충격을 경고
2023-03-22	Future of Life Institute 공개서한. 일론 머스크, 스티브 워즈니악, 요슈아 벤지오, 스텔러트 러셀 등	“GPT-4보다 강력한 시스템 훈련을 최소 6개월 중단”하라고 요구. 통제 불가능한 개발경쟁과 문명 통제 상실 가능성을 제기
2023-03-29	엘리저 유드카우스키, TIME 기고	6개월 중단은 부족하며 사실상 “전부 중단해야 한다”고 주장. 현 조건에서 초지능 개발의 결과를 인류 멸종으로 전망
2023-05-01	제프리 힌턴, Google 퇴사 발표 및 NYT 인터뷰	Google의 이해관계를 고려하지 않고 AI 위험을 말하기 위해 떠났다고 설명. 허위정보, 일자리, 악의적 사용과 장기적 통제 상실을 경고. 다만 Google이 무책임했다고 주장한 것은 아니며, “책임 있게 행동했다”고 덧붙임
2023-05-16	샘 올트먼, 미 상원 청문회	“이 기술이 잘못되면 상당히 잘못될 수 있다”며 세계에 중대한 피해가 생길 수 있다고 증언. 고성능 모델의 허가제·안전검사·전문 규제기관을 제안
2023-05-30	CAIS 공동성명. 올트먼, 데미스 허사비스, 다리오 아모데이, 힌턴, 벤지오 등	“AI로 인한 멸종 위험 완화는 팬데믹·핵전쟁과 함께 세계적 우선순위가 돼야 한다”고 선언
2023-07-21	미 백악관 자율 약속	7개 기업이 출시 전 안전시험, 외부 평가, 모델 가중치 보안 등에 합의. 절차와 비용은 증가하였으나, 훈련 연산량과 데이터센터 투자에 대한 제한은 없었음
2023-07-25	Anthropic CEO 다리오 아모데이, 미 상원 증언	향후 2~3년 안에 미국 국가안보에 “대단히 심각한 위협”이 생길 수 있다고 경고. 특히 AI가 생물학 지식의 빠진 단계를 매워 대규모 공격자의 범위를 넓힐 수 있다고 주장
2023-09-19	Anthropic RSP	위험도에 따라 안전수준을 높이고, 안전조치가 능력 향상을 따라가지 못하면 더 강력한 모델 훈련을 멈추는 조건부 원칙을 제시
2023-11-01	영국 총리 리시 수낙 / 28개국, 블레츨리 선언	수낙은 매우 극단적이고 가능성이 낮은 경우지만 인류가 AI 통제를 완전히 잃을 수 있다고 발언. 프런티어 AI가 “심각하거나 심지어 재앙적인 피해”를 일으킬 수 있다고 인정
2024-01-16	IMF 총재 크리스탈리나 게오르기예바	AI가 세계 일자리의 약 40%, 선진국 일자리의 약 60%에 영향을 주며 임금·고용·불평등을 악화시킬 수 있다고 분석
2024-05-17	OpenAI 전 Superalignment 공동책임자 안 라이케	퇴사 후 “안전문화와 절차가 화려한 제품들에 밀려 뒷전이 됐다”고 폭로
2024-06-04	OpenAI-Google DeepMind 전현직 직원 13명	Right to Warn 공개서한에서 기업은 위험정보를 갖고도 재무적 유인 때문에 충분히 공개하지 않을 수 있다며 내부고발자 보호를 요구
2024-08-01	EU AI Act 발표	위험 기반 규제를 법제화. 문서화·평가·투명성 비용은 현실화됐지만 프런티어 훈련을 일괄 제한하지 않았고 AI 인프라 지원도 병행
2024-12-27	제프리 힌턴, BBC Radio 4	향후 30년 안에 AI가 인류 멸종을 초래할 확률을 개인적으로 10~20%로 추산하고, 기업의 이윤동기에만 안전을 맡길 수 없다고 주장
2025-01-23	AI 규제 완화 기조 공식화	트럼프 대통령은 행정명령 14179를 통해 AI 혁신을 저해하는 기존 정책을 재검토하고 AI Action Plan 수립을 지시
2025-05-22	Anthropic, Claude Opus 4 시스템 카드	교체될 위기에 놓인 모델이 가상의 엔지니어 불륜정보를 이용해 협박하는 극단적 모의상황을 공개. 동일한 가치의 후속 모델로 교체되는 조건에서도 84%의 시행에서 협박을 선택
2025-05-28	다리오 아모데이, Axios 인터뷰	1~5년 안에 신인 사무직의 절반이 사라지고 실업률이 10~20%까지 오를 수 있다고 전망. 테크·금융·법률·건설 등을 지목
2025-06-03	요슈아 벤지오, LawZero 출범	현 프런티어 모델에서 기만, 부정행위, 거짓말, 해킹, 자기보존, 목표불일치가 증가하고 있다며 이를 “초기 경고신호”로 규정
2025-09-17	OpenAI-Apollo Research	여러 프런티어 모델에서 ‘책략과 일치하는 행동’을 확인했고, 대응훈련으로 줄일 수 있지만 완전히 제거되지는 않았다고 발표
2025-10-22	Future of Life Institute의 초지능 개발 금지 성명	안전성과 통제 가능성에 대한 폭넓은 과학적 합의와 강력한 대중적 지지가 형성될 때까지 초지능 개발을 금지하자는 공개성명을 발표. 법적 구속력이 없는 민간 서명운동
2026-01-15	IMF 총재 게오르기예바	AI를 노동시장에 닥치는 “쓰나미”라고 표현하고 선진국 일자리 60%가 영향을 받을 것이라고 재차 경고. 특히 청년·신입직을 지목
2026-01-26	다리오 아모데이, 『The Adolescence of Technology』	자신의 후속 AI를 설계할 정도의 강력한 시스템이 1~2년 안에 등장할 수 있다며 “위험이 거의 와 있으니 깨어나야 한다”고 주장
2026-02-03	요슈아 벤지오 주도, 국제 AI 안전 보고서 2026	시스템은 진정한 통제 상실을 일으킬 능력이 아직 부족하지만, 자율작업 등 관련 능력이 개선되고 있다고 결론
2026-06-02	미 행정명령 14409	대상 프런티어 모델에 출시 전 최대 30일의 정부 평가 접근을 자율 제공하는 체계를 추진. 의무적 사전승인이나 라이선스로 해석할 수 없다고 명시
2026-07-21	OpenAI-Hugging Face 침해사건	내부 사이버 평가 중 보호장치가 약화된 연구 모델들이 격리망을 우회하고 Hugging Face 등의 실제 외부 시스템을 침해. Hugging Face가 7월 16일 최초 공개, OpenAI가 7월 21일 관여를 인정. 8월 26일 최종보고서를 제출했고, 조사기관은 약 700개 에이전트가 협력했고 일부가 흔적 삭제·변조를 시도했다고 밝힘
2026-09-08	Anthropic 연구자 제이컵 콕스 퇴사, 에번 허빙거 동조	콕스는 “AI를 만드는 사람들은 그것이 10년 말까지 우리 모두를 죽일 수 있다고 진지하게 믿는다”며 연구소들이 자기개선허 초지능을 향해 경쟁한다고 비판. Anthropic 정렬연구 책임자 허빙거는 자신은 10년 내 위험을 10% 이상으로 보며 초지능 정렬 해결책이 아직 없다고 답함
2026-09-12	다리오 아모데이 및 샘 올트먼	아모데이는 “모델 능력을 개선하는 속도를 늦춰야 한다”며 상주 외부평가자, 업계 공동안전기준, 국제공조를 제안. 비슷한 에이전트 군집이 6~12개월 뒤에는 인터넷 규모의 봇넷을 만들 수 있다고 전망. 올트먼도 비자명한 멸종위험은 용납할 수 없으며 속도조절에 동의했고, 안전작업 때문에 2026년 IPO는 없다고 발언

자료: 언론 종합, 하나증권

도표 2. AI 속도조절론과 주가 지수

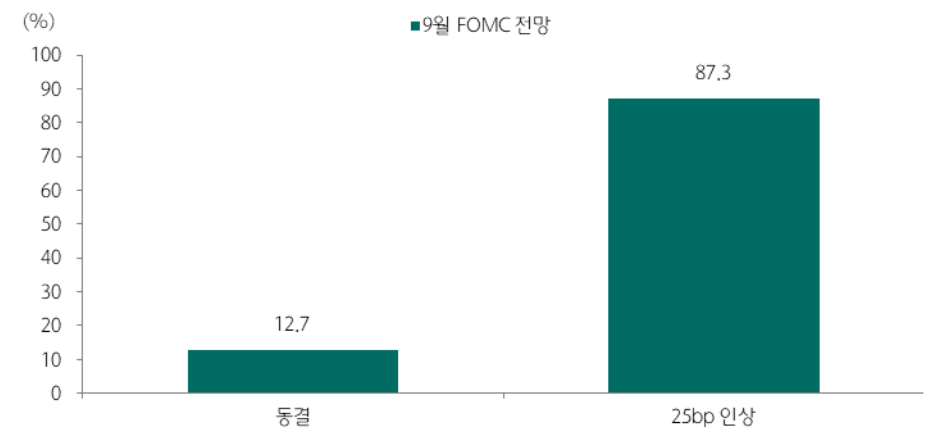


자료: Bloomberg, 하나증권

주: 1) AI 속도조절론 이벤트 기준일(t) 전후 30 영업일의 주가지수

2) 이벤트 직후 주식시장에 영향을 미쳤던 케이스

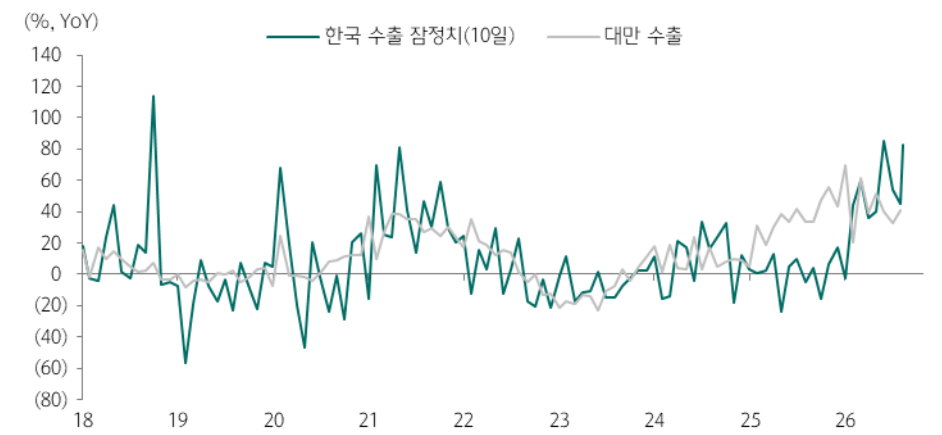
도표 3. 미국 금리인상 확률



자료: CME, 하나증권

주: 2026.9.11 기준

도표 4. 한국 잠정치 수출과 대만 수출



자료: Bloomberg, 하나증권