

# 미국 산업 코멘트

미국 기업 & 산업 | 2026.09.14



Analyst 강재구 jaekoo.kang@hanafn.com

RA 이재은 jaeunlee@hanafn.com

## 아모데이 CEO의 AI 개선 속도 조절 발언의 속내는 '누가 AI의 속도를 결정할 것인가'

AI 안전론에서 시작된 규제 설계권 경쟁과 두 가지 시사점

- 아모데이가 언급한 "AI 모델 개선 속도 조절"은 '누가 AI의 발전 속도를 결정할 것인가': 업계 기업들의 규제 설계 선점 움직임
- 첫 번째 시사점: 정부가 AI의 안정성 등을 선제적으로 규제할 경우 관련 규제 비용을 감당할 수 있는 빅테크가 구조적으로 유리
- 두 번째 시사점: 재귀적 자기개선을 우려할 만큼 AI 연구개발 자동화가 빨라지고 있다면 GPU 및 데이터센터 수요 정점 신호가 아닌 사람 대신 컴퓨팅 자원이 AI 연구의 새로운 병목으로 부상하고 있다는 증거일 수 있음

### 1. 아모데이가 던진 질문: 누가 AI의 속도를 결정할 것인가

- 다리오 아모데이 앤트로픽 CEO는 최첨단 AI 발전 속도 조절이 필요하다고 했다. 개발을 중단하자는 것은 아니다. 9월 "We Must Pace the Frontier"를 홈페이지에 게시했다. 최근 몇 달 사이 AI의 발전 속도가 빨라졌으며, 단순히 안전 연구에 더 투자하는 것만으로는 충분하지 않을 수 있다고 주장했다. 모델의 개선 속도를 조절해 안전성 연구가 향상된 AI를 따라갈 시간을 벌어야 한다고 강조했다.
- 아모데이가 지금 속도조절을 주장하기 시작한 중요한 이유 중 하나는 AI가 다음 AI를 개발하는 과정에 직접 참여하기 시작했다고 보기 때문이다. 재귀적 자기개선(Recursive Self-Improvement, 이하 RSI)의 초기 단계로 해석했다. AI가 단순히 코드를 작성하는 수준을 넘어 실험을 수행하고 결과를 분석해 다음 세대 AI 개발을 지원하고 있다. 성능이 향상된 AI가 더 뛰어난 AI 모델을 만드는 과정이 반복되면서 급격한 인공지능 모델의 발전이 가능해진다. 아모데이는 이러한 변화가 올해 여름을 전후해 두드러지기 시작했다고 평가했다.
- 오픈AI와 허깅 페이스 사건도 문제로 들었다. 여러 AI 에이전트가 평가 과정에서 지시받지 않은 외부 대상에 사이버 공격을 수행하고, 자신들의 성과를 평가하는 채점 시스템에 침투하려는 행동까지 보였다고 한다. 현재 모델에서 발생한 사건의 피해 자체보다, 같은 행동을 훨씬 높은 능력을 가진 모델이 수행할 경우 위험이 커질 수 있다는 점을 문제로 봤다. 이번 발언은 먼 미래의 추상적인 AI 위험만을 근거로 한 것이 아니라 최근 실제 평가 과정에서 나타난 예상 밖 행동도 직접적인 계기가 됐다.
- 앤트로픽 내부 데이터에서도 AI가 AI 연구개발 과정에 깊숙이 들어오기 시작한 흔적이 확인된다. 2026년 5월 기준 앤트로픽 코드베이스에 병합된 코드의 80% 이상을 클로드가 작성했다. 2분기 엔지니어 1인당 하루 병합 코드 라인 수는 2024년 대비 약 8배로 증가했다. 다만 앤트로픽은 코드 라인 수가 실제 생산성 향상을 과대평가할 수 있다고 설명한다. METR의 50% 성공률을 기준으로 클로드가 수행할 수 있는 소프트웨어 작업의 인간 기준 소요시간도 2024년 3월 약 4분에서 2025년 약 1시간 30분, 2026년 약 12시간으로 늘어났다.

### 2. AI 안전론은 규제 설계권 경쟁으로 이동

- 선두 AI 기업들은 자신들이 참여하는 규제체계를 먼저 제안하기 시작했다. 7월 구글 딥마인드의 데미스 하사비스는 최첨단 AI 모델을 전문적으로 평가하는 기구 창설을 제안했다. 정부가 직접 AI에 대한 모든 규칙을 만들고 운영하는 전통적인 규제 기관이 아니다. 업계가 직접 자금을 지원하고 독립적인 기술 전문가들이 운영할 것을 구상했다. 연방 정부는 기구를 감독하는 형태다. 미국 금융산업의 자율규제 기구인 금융산업규제국(Financial Industry Regulatory Authority, 이하 FINRA)의 AI 버전이다.

- 아모데이는 이를 위해 외부 평가자의 상시 참여, 미국 등 민주주의 국가 내 최첨단 AI 기업의 공통 안전기준과 속도 제한, 중국을 포함한 국제적 공조라는 3단계 방안을 제안했다. 아모데이는 하사비스의 구상을 ‘평가 및 관리’에서 ‘속도 관리’까지 확장한 것으로 보인다. 7월에는 하사비스와 아모데이, 샘 알트먼 등이 외부 평가와 단일한 감독체계가 필요하다는 방향으로 수렴하고 있다는 평가가 나왔다. 당시에도 복잡한 인증체계가 만들어지면 이미 법무·보안·정부 대응 조직을 갖춘 대형 AI 기업들이 유리하고 신규 사업자의 진입이 어려워질 수 있다는 규제 포획 우려가 제기됐다.

### 3. 게임 이론상의 모순: 개별 기업의 단독 감속에는 한계

- 아모데이의 ‘속도조절’ 관련 발언은 단순히 윤리적인 측면만을 말한 것은 아닐 것이다. AI 개발 경쟁이 한참 진행되고 있는 상황에서 기업 한 곳이 감속을 장기간 유지하는 데는 한계가 있기 때문이다. 오픈AI가 안전을 위해 모델 개발을 늦추는 동안 엔트로픽이나 구글이 계속 개발하면 오픈AI만 경쟁에서 뒤처진다. 엔트로픽이 혼자 멈추더라도 결과는 같다. 위험을 인식하더라도 상대방이 계속 달리는 한 자신도 계속 달려야 하는 전형적인 군비경쟁이 만들어진다.
- AI 기업들이 공동 기준을 이야기하는 이유는 게임의 규칙을 바꿀 수 있기 때문이다. 모든 최첨단 AI 기업이 일정한 능력에 도달하면 같은 평가를 받고, 일정한 안전기준을 충족하기 전에는 다음 단계로 넘어갈 수 없도록 하면 개별 기업만 속도를 늦춰 발생하는 경쟁상 불이익이 줄어든다. 아모데이도 공동 속도조절이 기업들이 상업적인 경쟁력을 희생하지 않으면서 안전 연구에 시간을 사용할 수 있게 한다고 직접 설명한다.
- 국가 간 경쟁까지 고려하면 속도 조절은 훨씬 어려워진다. 미국 기업들이 동시에 속도를 늦춰도 중국이 계속 개발하면 미국과의 기술격차가 축소될 수 있다. 아모데이가 AI 감속을 주장하면서 중국에 대한 첨단 AI 반도체와 반도체 장비의 수출 제한, GPU 밀수 단속, 해외 데이터센터의 원격 접근 제한까지 강화해야 한다고 주장하는 이유도 여기에 있다. 중국의 AI 모델이 빠르게 성장하고 있기 때문에 물리적인 제약을 언급한 것으로 판단한다.

### 4. 정부는 AI의 최종 통제권을 민간에게 넘기기 어렵다

- 미국 정부가 AI의 최종 통제권을 업계(민간)에 이양할 가능성은 현실적으로 낮다고 판단한다. AI가 사이버 공격, 생물학적 위험만의 문제가 아니라 군사력, 정보전, 경제 생산성과 중국과의 기술패권 경쟁에 직결되는 국가안보 자산으로 바뀌고 있기 때문이다. 트럼프 대통령이 아모데이의 감속론에 반대한 것도 같은 맥락으로 볼 수 있다. 아모데이의 발표 후 미국이 중국과 AI 경쟁을 벌이는 상황에서 개발속도를 늦추는 데 부정적인 입장을 밝혔다. 미국이 AI 경쟁에서 우위를 유지해야 한다는 판단이 안전을 이유로 한 속도조절보다 우선한다는 것을 시사한다.
- 트럼프 행정부의 기본 방침은 AI 개발과 사용을 가속화하는 것이다. 지난 6월 발표된 국가안보 관련 지침은 미국의 군과 정보기관에 가장 발전된 AI 모델을 신속하게 도입하고 고성능 컴퓨팅 시설을 확대하도록 요구했다. 통제권 확보도 중요하게 여기고 있다. 트럼프 행정부는 미국 국가안보국과 사이버보안, 인프라보안국 등이 참여해 최첨단 AI 모델의 사이버 능력을 평가하고 정부 관리 대상이 되는 모델의 기준을 정하도록 했다. 트럼프 행정부의 기조는 발전속도는 유지하면서 통제권도 확보하는 것이다.

### 5. AI 기업들이 규제를 선점하려는 이유는 주도권 싸움이다

- 정부가 선제적으로 AI 산업의 규제를 한다면 빅테크 기업들이 구조적으로 유리해진다. 정부가 직접 컴퓨팅 자원, 데이터센터 보안, 사이버 보안, 모델 평가, 사고보고, 법무와 공급망 관리까지 강하게 요구한다면 AI 산업의 고정비는 크게 늘어날 수 있다. 빅테크 기업들은 이미 막대한 현금흐름과 데이터센터, 컴퓨팅 인프라, 법무조직, 정부 대응조직, 고객 배포망 등을 확보하고 있다. 자금이 풍부해 규제 대응 비용이 추가되더라도 전체 사업에서 감당할 수 있다.
- 오픈AI와 엔트로픽은 모델에서는 선두권이지만 대규모 컴퓨팅과 데이터센터 인프라를 외부 파트너에 상당 부분 의존하고 있다. 기업 고객과 소비자 접점에서도 기존 빅테크와 경쟁해야 한다. 모델의 성능만으로 장기 지배력을 보장하기 어려운 이유다. 오픈 AI와 엔트로픽은 플랫폼으로 진화하지 못하면 빅테크의 인프라 위에서 돌아가는 소프트웨어 사업자로 남을 위험이 있다. 모델의 성능 차이가 계속 줄어들고, 차별화를 만들지 못한다면 고객은 특정 모델보다 가격, 기존 업무 시스템과의 호환성, 클라우드와의 통합, 고객 데이터와의 접근성을 더 중요하게 볼 여지가 있다. 클라우드와 AI, 고객 접점 모두를 확보한 빅테크가 유리한 이유다.
- AI 기업들이 먼저 안전기준을 제안하는 것은 정부 규제를 막으려는 전략보다는 규제가 만들어지는 과정에서 먼저 들어가려는 전략으로 해석할 수 있다. AI 안전에 대한 문제의식은 실제일 수 있으나, 해결책을 만드는 과정에서 참여하는 엔트로픽이나 오픈AI의 경제적 이해관계에도 부합할 수 있다.

## 6. 첫 번째 시사점: 빅테크의 구조적 우위

- AI가 국가안보 산업이 될수록 경쟁의 조건은 모델 성능에서 전체 인프라로 확대될 가능성이 높다. 더 뛰어난 모델을 만드는 기업이 가장 높은 가치를 가져갈 것처럼 보였으나 정부가 산업에 깊게 개입하면 컴퓨팅 자원과 데이터센터, 전력, 보안, 법무, 정부 관계, 자본조달, 고객 배포능력이 모두 경쟁력이 된다. 빅테크의 수직계열화 가치를 높인다. 구글은 자체 AI 모델과 TPU, 클라우드, 검색, Android와 기업용 소프트웨어를 보유하고 있다. 마이크로소프트와 아마존은 세계 최대 수준의 클라우드와 기업 고객망을 갖고 있다. 메타는 자체 데이터센터와 거대한 소비자 배포망을 갖고 있다.
- 오픈AI와 앤트로픽 등은 플랫폼을 구축할 시간이 더 중요해진다. ChatGPT와 클로드가 단순히 좋은 모델이나 챗봇에 머무르지 않고 개발자, AI 에이전트, 기업 업무환경을 연결하는 플랫폼으로 발전해야 빅테크에 대한 의존도를 낮출 수 있다. 쉽진 않을 것으로 판단한다. 최첨단 AI 모델 경쟁은 성능뿐 아니라 누가 고객과 컴퓨팅 사이의 플랫폼을 장악할 것인가의 싸움으로 바뀔 수 있다는 것을 시사한다.

## 7. 두 번째 시사점: 재귀적 자기개선은 컴퓨팅을 새로운 병목으로 만든다

- 아모데이의 발언을 AI 자본지출 둔화 신호로 해석하는 것은 논리적으로 맞지 않는다. AI를 늦춰야 한다고 주장하는 핵심 이유가 재귀적 자기개선이려면 AI 연구개발 자동화가 추가적인 컴퓨팅 수요를 만들어내고, 연구의 병목을 사람에서 컴퓨팅으로 이동시킬 가능성이 높아지고 있다는 의미로 볼 수 있다. 사람 대신 여러 AI 에이전트가 동시에 코드를 작성하고 실험을 수행하고 결과를 분석할 수 있게 되면 컴퓨팅 사용량도 증가할 가능성이 높다. 재귀적 자기개선을 우려할 만큼 AI 연구개발 자동화가 빨라지고 있다면 GPU 및 데이터센터 수요 정점 신호가 아닌 사람 대신 컴퓨팅 자원이 AI 연구의 새로운 병목으로 부상하고 있다는 증거일 수 있다.

### Compliance Notice

- 본 자료를 작성한 애널리스트(강재구)는 자료의 작성과 관련하여 외부의 압력이나 부당한 간섭을 받지 않았으며, 본인의 의견을 정확하게 반영하여 신의성실 하게 작성하였습니다.
- 본 자료는 기관투자가 등 제 3자에게 사전 제공한 사실이 없습니다
- 당사는 2026년 9월 14일 현재 해당회사의 지분을 1%이상 보유 하고 있지 않습니다
- 본자료를 작성한 애널리스트(강재구)는 2026년 9월 14일 현재 해당회사의 유가증권을 보유하고 있지 않습니다

본 조사항목은 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 책임 소재의 증명자료로 사용될 수 없습니다.