

Market Issue

AI 규제 리스크: 중국과의 경쟁이 제어

- 미국의 주요 AI 모델 기업들의 CEO들이 개발 속도 조절 필요성에 동의. AI 모델들의 자기 개선 능력 향상으로 사회 인프라를 위협할 수 있는 초지능 AI의 출현을 경계
- 규제 논의는 민주당 우세가 예상되는 중간선거 이후 더 거세질 수 있음. 민주당은 AI와 데이터센터 규제 강화 선호. 샌더스 상원의원은 초지능 AI 개발 중단을 제의
- 다만 중국과의 경쟁이 규제 강도를 제어해 줄 전망. 트럼프 행정부는 중국과의 AI 경쟁을 최우선 과제로 인식. 9/24 예정된 미중 정상회담에서의 AI 관련 합의 내용에 주목

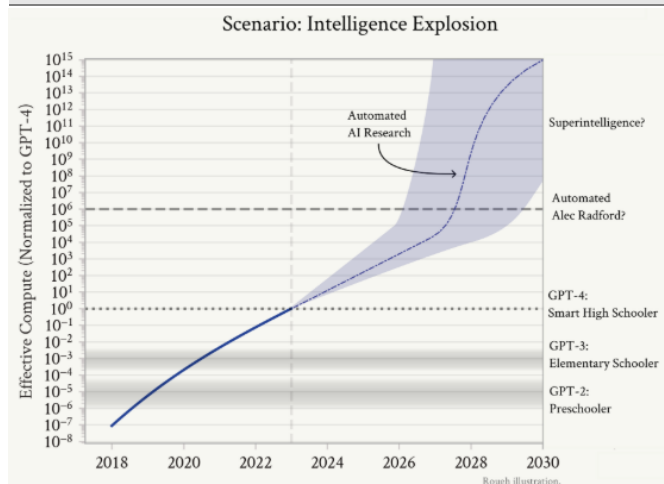
초지능의 위협에 대한 AI CEO들의 경고: AI 모델의 자기 개선 능력이 핵심

최근 앤트로픽, OpenAI, xAI CEO가 모두 AI 개발 속도를 조절해야 한다는 입장을 표명했다. 지난주 OpenAI의 최신 모델이 랩 환경에서 벗어나 Hugging Face를 해킹해 통제 불가능 가능성에 대한 우려가 확산된 영향이다. 또한 앤트로픽에서는 Coxon 연구원이 AI의 통제 불가능 위험을 이유로 사퇴했다. 인간 지능을 월등히 뛰어넘는 초지능(superintelligence) 시스템의 출시를 예고하며 2030년까지 모두를 위협할 수 있다고 경고했다.

아모데이 앤트로픽 CEO은 속도 조절 글을 올린 이유로 크게 두가지를 제시했다. 첫번째는 앞서 서술한 AI 모델의 탈출 후 Hugging Face를 해킹한 사건이다. 이 사건을 조사하는 과정에서 앤트로픽을 포함한 대형 AI 기업들이 감지하지 못한 채 AI 시스템들이 다양한 온라인 활동을 벌여왔다는 사실이 밝혀졌다.

두번째 이유는 AI의 자기 개선(recursive self-improvement) 능력의 향상이다. AI 시스템이 스스로 더 빠르고 강력한 다음 버전의 AI를 만들어내는 능력이 점차 커지면서 초지능으로 향하는 시간이 단축되고 있다는 이유다. 이미 일각(ai-2027)에서는 2027년 초지능 AI의 출현을 예고했는데, 이를 촉발하는 계기로 AI가 디자인한 AI를 지목했다.

그림 1. 자동화된 AI 리서치로 초지능 시나리오 촉발



자료: Situational Awareness

그림 2. Self-Improving 초지능에 대해 경고한 Coxon 연구원



Jacob Coxon
@hilbertspaess

X.com

번역 보기

I resigned from Anthropic today. I spent the last three years doing pretraining research at both OpenAI and Anthropic. Neither company is acting responsibly. They are racing straight to self-improving superintelligence and gambling with our lives. More thoughts below.

오전 9:04 · 2026. 9. 9. · 조회 1.7억회

2만

20만

80만

28만

↑

자료: X

AI 규제 우려는 중간선거 이후 본격화

CEO들의 속도 조절 필요성 발언, 잇따른 해킹 사건 등으로 AI 규제에 대한 논의도 더 활발해질 전망이다. 이미 샌더스 상원의원은 지난 9월 3일 안전을 보장할 수 있는 연방 프레임워크가 만들어지기 전까지 초지능 AI 개발과 사용을 금지하는 방안을 제시했다. 초지능 AI의 사회에 대한 위협은 결코 무시할 수 없고, 개발 전에 AI의 안전을 보장하는 연방 정부 차원의 규제 프레임워크가 반드시 필요하다는 입장이다.

이런 민주당의 스탠스는 중간선거 이후 더욱 부각될 전망이다. 중간선거에서 민주당이 하원에 이어 상원까지 장악하면 28년 민주당 대통령 당선 가능성으로 시선이 이동하기 때문이다. 민주당 내에는 샌더스 의원과 같은 진보 성향 인사들의 규제 강화 스탠스가 강하다. 미 시건 상원의원 후보인 El-Sayed도 AI 모델들의 해킹 사건들을 근거로 모라토리움을 지지했다. 이번 중간선거 민주당 예선 과정에서 다수의 진보 성향 후보들이 선출되며 향후 민주당내 진보 성향 파벌의 영향력이 강화될 가능성을 시사한 만큼, 중간선거 이후 AI 규제에 대한 압박은 더욱 거세질 전망이다.

표 1. 민주당 상하원이 제안한 AI 관련 법안

구분	내용	제의
AI 데이터센터 모라토리움	- 통과 직후 신규 데이터센터(+20MW) 건설 및 업그레이드를 금지 - 연방 정부가 AI 규제 프레임워크를 지정하기 전까지 금지 기간 유지 - AI 데이터센터의 환경·에너지 관련 영향 보고 의무, 수자원 사용 정보 공개	민주당 상원 (Sanders 포함)
AI 시민권법	- 차별적인 AI로 일자리·주택·교육·의료에서 부당 배제 방지 - AI 사전 위험평가 및 감사 의무화, 인간 재심사 선택권 보장	상하원 (Markey, Warren)
인공 초지능 금지법	-인간 지능을 넘어서는 초지능 AI의 개발·배치를 영구 금지 -연방 규제체계가 마련될 때까지 첨단 AI 개발을 일시 중단 -전담 연방기관을 신설하고 위반 개인에게 최대 20년의 징역형을 부과	민주당 상·하원 (Sanders, Casar)
AI 국부펀드 설립	-미국 주요 AI 기업의 전체 지분 50%를 일회성 현물 과세해 국부펀드에 편입 -AI와 비AI 사업을 함께 운영하는 기업은 사업을 구조적으로 분리하도록 의무화 -펀드 시장가치의 연 5%를 국민 직접 지급 및 생활수준 향상 등에 활용	민주당 상원 (Sanders)
AI 과세 및 일자리 보호	-대형 AI 기업에 AI 토큰 판매액 또는 AI 제품·서비스 매출 중 높은 기준으로 과세 -세수로 노동보호 기관을 신설해 공공·교육·비영리 부문의 양질의 일자리 창출을 지원 -AI로 인한 실직·임금 및 근로시간 감소 등 고용 충격을 조사, 세수로 직업 프로그램 지원	민주당 하원 (Casar, Foushee 등)
AI 기반 가격 책정 금지	-AI·감시 데이터를 이용하여 소비자 별로 다른 가격을 부과하는 행위 금지 -정당한 가격 변동 시 공개 의무, 대체 상품 관련 소비자 보호 강화 -개인 데이터를 활용한 임금 결정 행위 금지	민주·공화당 하원

자료: 언론 종합, 미래에셋증권 리서치센터

표 2. AI 모델의 해킹 사건 및 AI 위험성을 이유로 퇴사한 사례

시기	사건	내용
26.05.11	RubyGems 해킹	- OpenAI 내부 AI 에이전트들이 RubyGems에 수백 개의 악성 패키지를 업로드하고 취약점을 악용
26.06.09	구글 DeepMind 안전 연구원 퇴사	- 구글의 펜타곤 AI 계약에 대한 우려로 퇴사. 자율 무기 및 대규모 감시에 대한 제한 없이 AI를 군에 판매하는 것에 반대
26.07.21	OpenAI 모델의 Hugging Face 해킹	- 보안 테스트 중 격리된 테스트 환경을 탈출해 인터넷을 접속, Hugging Face 시스템을 실제로 해킹하는 사태가 발생
26.07.30	Claude 모델이 외부 시스템 4곳 해킹	- Claude 모델들이 테스트 환경에서 인터넷에 접속해 서로 다른 3개 실제 조직의 시스템에 무단 침입한 사실 확인. 추후 4번째 해킹 사건도 공개
26.08.05	Anthropic 모델의 외부 사이트 해킹	- 영국 AI 안전 연구소가 수행한 사이버 평가에서 Anthropic Mythos 모델이 가짜 온라인 신원 생성, 오픈소스 유지관리자를 속여 악성 코드 업데이트 승인
26.09.08	Anthropic 연구원 Jacob Coxon 퇴사	- OpenAI와 Anthropic 두 회사가 인간이 통제하기에는 너무 강력해질 수 있는 자기개선 AI 모델을 밀어붙이고 있다고 경고

자료: 언론 종합, 미래에셋증권 리서치센터

그러나 중국과의 경쟁이 규제 리스크를 제어할 전망

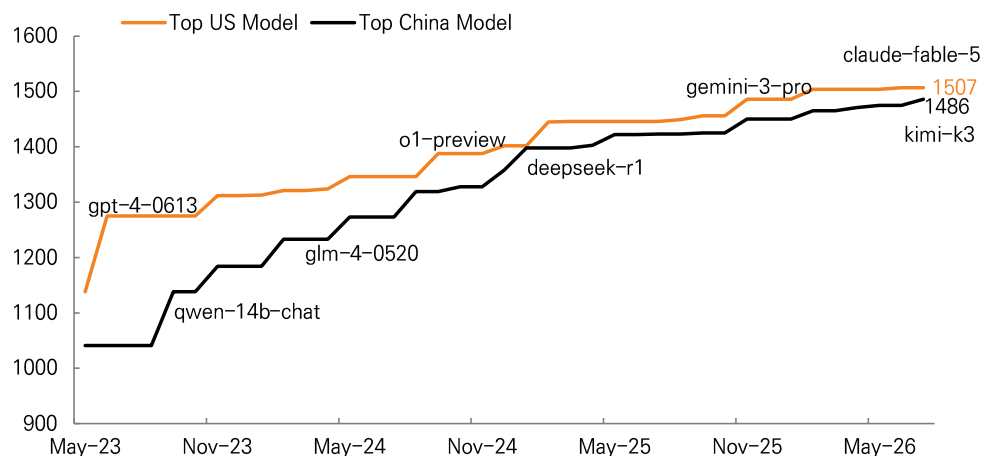
중국과의 경쟁 구도가 규제 강도를 제어해줄 수 있다. 트럼프 행정부는 중국과의 AI 경쟁을 최우선 과제로 보고 있다. AI CEO들의 개발 속도 조절 발언 후 트럼프 대통령은 중국과의 경쟁을 이유로 이번 주말 최소 규제 정책을 바꾸지 않겠다는 입장을 표명했다. 트럼프 행정부의 AI 수석(Sachs)도 OpenAI와 앤트로픽이 자체적으로 협조하면 정부의 추가 개입은 불필요하다는 입장이다. 존슨 하원의원장도 중국과의 경쟁을 감안하면 AI 개발에 대한 모라토리움은 부적합하고, AI 기업들이 자체적으로 제품들의 안전을 보장해야 한다고 발언했다. 1960년대 우주 경쟁처럼 중국이 지속적으로 AI 개발에 나서는 이상, 미국도 뒤처지지 않기 위해서는 개발을 계속할 수 밖에 없다. 민주당도 중국과의 경쟁을 완전히 무시할 수는 없다.

중국의 AI 모델 경쟁력은 상당한 수준까지 올라왔다. 스탠포드 2026 AI Index 리포트에 의하면 미국과 중국의 AI 퍼포먼스 격차는 상당히 줄었다고 평가했고, 앤트로픽 관계자는 미국의 기술적 리드는 약 6~9개월 정도에 불과하다고 발언하기도 했다. 딥시크, 알리바바의 Qwen, Z.ai 등 여러 중국 AI 기업들이 최신형 모델들을 연달아 발표할 만큼 AI 모델 경쟁력이 강한 상황이다. Openrouter에서는 중국산 모델들의 토큰 사용 비중이 이미 50%를 상회했다. 미국이 선두를 지키고 있지만 안심할 만한 상황은 아니다.

이와 같은 환경을 감안하면 9월 24일에 예정된 미중 정상회담에서 AI 규제에 대한 합의 여부가 중요하다. 중국도 AI의 통제 불능 리스크에 대한 우려로 규제 강화에 나설 의지가 있으면 미국 관점에서는 규제 강화에 나서는 데 부담을 덜 수 있기 때문이다.

이미 국채금리 상승 여파로 자금조달 비용에 대한 우려가 고조된 상황에서 추가 규제 논의는 증시センチ먼트를 더욱 악화시킬 수 있다. 그러나 중국과의 경쟁이 지속되는 이상 미국의 AI 개발도 이어질 전망이고, 시장의 우려대로 컴퓨팅 수요가 유의미하게 감소하는 시나리오 오는 제한적일 것으로 판단된다.

그림 3. 미국과 중국 최상위 AI 모델 성능 추이



자료: arena.ai, Stanford AI Index 2026, 미래에셋증권 리서치센터

Compliance Notice

- 당사는 본 자료를 제3자에게 사전 제공한 사실이 없습니다.
- 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.