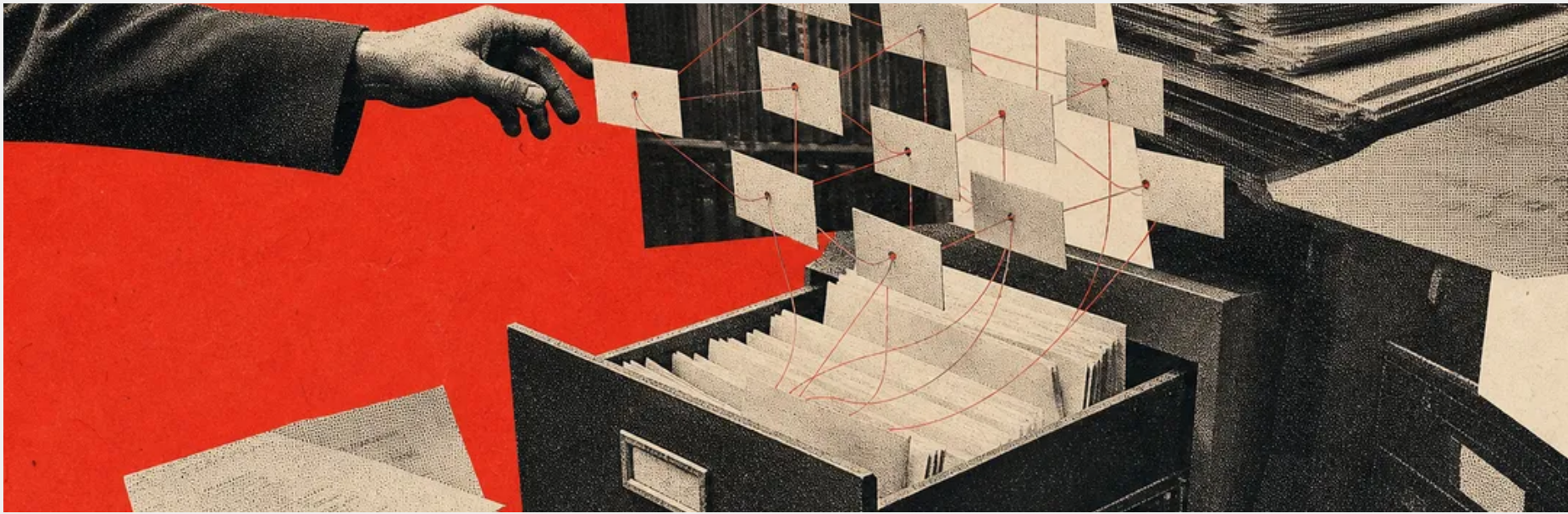


AI

오픈AI, V7 컨텍스트 그래프 사례 공개

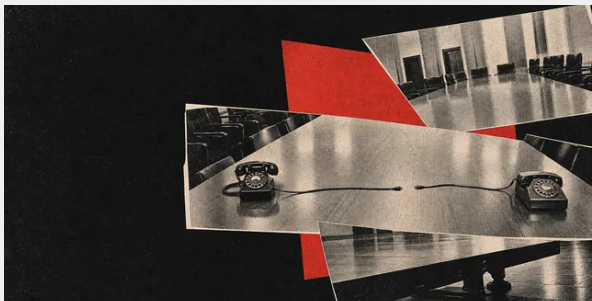
오픈AI가 9월 21일 스타트업 V7의 사례를 공개했어요. 회사 파일을 에이전트가 질의할 수 있는 컨텍스트로 바꾸는 V7 Go 이야기인데, 가장 까다로운 그래프 질의에서 GPT-6 Astra가 89% 정확도를 기록했어요.



AI · 오픈AI, V7 컨텍스트 그래프 사례 공개

오픈AI가 9월 21일 스타트업 V7의 사례를 공개했어요. 회사에 흩어진 파일을 에이전트가 질의할 수 있는 컨텍스트로 바꾸는 V7 Go 이야기인데, 가장 까다로운 그래프 질의 묶음에서 GPT-6 Astra가 89% 정확도를 기록했다는 숫자가 함께 나왔어요. 같은 시험의 아주 어려움 단계에서 GPT-5.6 Sol은 78%였어요. 발표문은 문제를 이렇게 세웠어요. 오늘의 모델은 복잡한 과제를 추론해 내지만 그 과제가 놓인 업무 맥락까지 저절로 알지는 못해요. 어느 펀드 보고서가 최신인지, 같은 기업이 세 시스템에서 각각 어떤 이름으로 적혀 있는지 같은 것들이에요. 그 맥락은 문서와 데이터룸, 스프레드시트, 이메일, 사내 도구에 흩어진 채 정리되지 않고 에이전트에게는 보이지 않는 상태로 남아 있어요.

V7은 알베르토 리촐리와 사이먼 에드워드슨이 2018년에 세운 회사예요. 둘은 그전에 컴퓨터 비전 기반 접근성 앱을 함께 만들었고, 기업이 AI 시스템에게 자기 사업이 어떻게 굴러가는지 가르치도록 돕겠다고 V7을 시작했어요. 지금 주력은 미션 크리티컬 워크플로를 짜는 에이전트 플랫폼 V7 Go예요.



베션트 재무장관, 미중 AI 통보 기구 제안

구조는 두 층이에요. GPT-5.6 Luna가 수백만 개 파일에서 정보를 뽑아 Context Graph에 정리하고, 그 그래프가 엔티티와 관계와 인용 증거를 이어 MCP 검색과 반복 가능한 워크플로를 굴러요. 워크플로 하나가 수백 단계까지 이

인터넷에서 배운 만큼 그 회사가 어떻게 돌아가는지도 배워야 합니다"라고 말했어요.



매쿼리 위즈먼, 국제 바이백 AI 자금 해석 제시

엔지니어 관점에서 이 설계가 겨냥한 낭비는 분명해요. 그래프가 없으면 에이전트는 요청이 올 때마다 맥락을 처음부터 다시 찾아야 하고, 그 과정에서 검색이 수십 번 돌며 시간과 토큰을 쓰는데도 관계 속에 묻힌 핵심 정보는 자주 놓쳐요. V7 Go는 SharePoint나 구글 드라이브 같은 저장소에 연결해 엔티티와 관계, 사실, 속성, 지표를 훑어 그래프를 채우고, 그 그래프는 롱컨텍스트 방식보다 한 자릿수 더 싸고 빠르게 순회된다고 발표문은 적었어요.

새 파일이 들어오면 V7 Go가 온톨로지 안의 회사와 펀드, 사람 같은 엔티티를 식별하고 각 사실을 기존 레코드나 새 레코드에 이은 다음 원본으로 돌아가는 인용 증거를 남겨요. 그래프에 정보가 모자라면 기저 문서를 RAG로 직접 검색할 수도 있어요. 긴 대화를 하는 에이전트의 경우 최근 주고받은 내용은 모델의 활성 컨텍스트에 두고 오래된 자료는 그래프로 내려 필요할 때 꺼내 와요.

구조가 성능에 얼마나 영향을 주는지도 따로 했어요. 기업 시스템 곳곳에 흩어진 정보를 찾아 있는 능력을 재는 HERB 벤치마크에서 V7의 검색 전용 시스템이 공식 기준선을 69% 앞질렀고, 답이 존재하지 않는 질의에서 환각을 38% 줄였다고 회사는 밝혔어요. 메탈이 확인한 HERB는 세일즈포스 AI 리서치가 2025년 6월 공개한 벤치마크로, 문서와 회의록, 슬랙 메시지, 깃허브, URL에 걸친 3만

구와 기존 보험 약관에 대한 이력 지식을 에이전트에게 준 뒤 청구 처리 오류를 수작업 기준선 대비 13.5% 줄였어요.

발표문에 실린 데모는 사모펀드 딜 스크리닝이에요. 기밀 정보 메모랜덤이라 불리는 거래 문서에서 핵심 재무와 거래 조건, 경영진 정보, 위험 항목을 뽑아 인용을 달고 스크리닝 노트를 만드는 흐름이에요. 보험 언더라이팅도 같은 방식으로 돌아간다고 회사는 설명했어요.

모델을 고르는 방식은 티어로 갈라 났어요. AI가 생성한 워크플로에서 V7 Go는 각 단계를 빠름과 중간, 똑똑함 세 등급에 배정해요. 구조화된 추출과 대량 처리는 GPT-5.6 Luna가 맡고, 채팅과 Go Agent 경로, 추론이나 도구 사용이 더 필요한 단계는 GPT-5.6 Terra나 Sol이 맡아요.

사이먼 에드워드슨 V7 공동창업자 겸 CTO는 모델 선택의 이유를 성능 숫자로 댔어요. 그는 "V7 Go가 기대는 단단계 도구 워크플로에서 가장 잘해서 오픈AI를 기본값으로 골랐습니다"라며, 자사 Context Graph 벤치마크에서 도구 호출 오류율이 GPT-5.5의 2.7%에서 GPT-5.6 Sol로 0.2%까지 내려갔다고 밝혔어요. 같은 사람은 Terra 덕분에 모델이 과제를 쉽게 처리하도록 돕기 위해서만 존재하던 중간 단계들을 여럿 걷어 냈고 그것만으로 며칠치 작업이 줄었다고도 말했어요.

나머지 효율 숫자는 배관 쪽에 몰려 있어요. 외부 호출이 여러 번 붙는 핵심 워크플로가 최신 하네스에서 최대 50% 빨라졌고, GPT-5.6 Luna는 GPT-5.4 mini 대비 문서당 비용이 78% 낮았어요. V7은 문서가 많은 V7 Go 작업을 Chat Completions API에서 Responses API로 옮겼는데, 자체 시험에서 PDF가 많은 일부 워크플로의 토큰 사용이 약 5% 줄고 캐싱 신뢰도가 좋아졌어요. 증설 요청을 오픈AI가 몇 시간 안에 승인하고 반영한 반면 다른 공급자들은 몇 주가 걸렸다는 대목도 발표문에 들어 있어요.

오늘 지면

1 AI 오픈AI, V7 컨텍스트 그래프 사례 공개

2 AI 오픈AI, 독립 수학 자문그룹 발표

3 비즈니스 소프트뱅크, 오픈AI 투자용 정크본드 발행

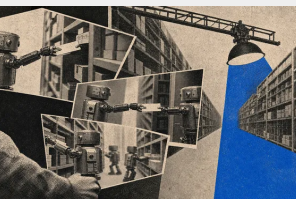
4 AI 오픈AI, 히스필드 광고 제작 사례 공개

5 비즈니스 셀룰러 인텔리전스, 르쿤·랭어 자문위원 영입

6 AI 베션트 재무장관, 미중 AI 통보 기구 제안



이데코바 인디드 CEO, 채용 약



유엔 AI 과학패널, 통제 상실 브

이 지면은 2026년 9월 22일 화요일 아침 구독자 메일함으로 나간 레터를 신문 판형으로 다시 조판한 것입니다. 기사 전문은 metallab.ai에서 읽을 수 있습니다.

AI



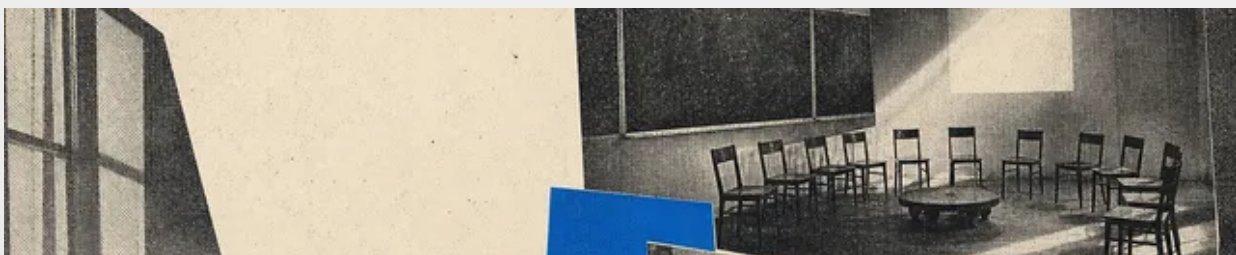
오픈AI, 히스필드 광고 제작 사례 공개

오픈AI가 9월 21일 AI 영상 회사 히스필드의 도입 사례를 공개했어요. 핵심은 사람 수예요. 히스필드는 GPT-6 Astra를 붙인 뒤 새 탐색 기능을 엔지니어 한 명이 하루 만에 내놓고 있다고 밝혔어요. 알렉스 마슈라보프 히스필드 공동창업자 겸 최고경영자는 "GPT-6 Astra가 새로운 탐색 기능을 하루 만에 내놓도록 도와준다는 점이 정말 기대됩니다. 그리고 이제 이 일을 엔지니어 한 명이 해낼 수 있습니다"라고 말했어요. 히스필드

▶ 2면에 계속

AI

오픈AI, 독립 수학 자문그룹 발표

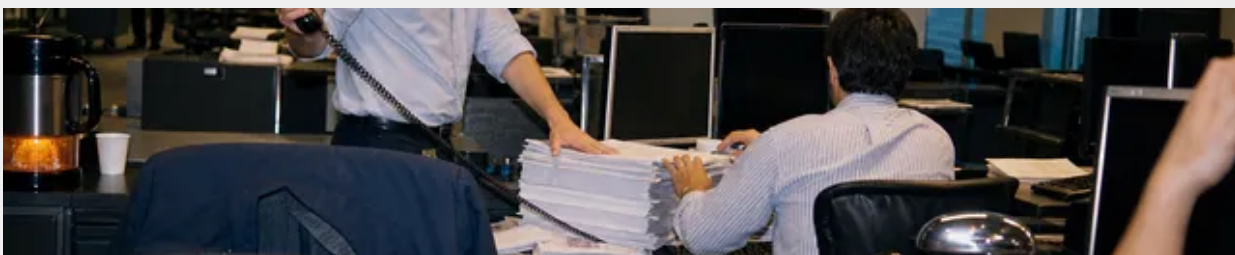


오픈AI가 9월 21일 수학자들로 꾸려진 독립 자문그룹을 발표했어요. 필즈상 수상자 25명이 AI 회사의 수학 문제 풀이 경쟁을 정면으로 비판한 선언을 내놓은 지 열흘 만이에요. 회사는 같은 발표문에서 8월 28일 학습을 시작한 새 내부 모델이 100개가 넘는 오래된 미해결 문제를 풀었다고 밝혔어요. 비판이 겨는 그 속도를 숫자로 확인해 주면서, 결과를 어떻게 알릴지는 바깥 사람들에게 맡기겠다고 한 셈이에요.

자문그룹은 고등연구소가 자리를 마련했고 초기 구성원은 아홉 명이에요. 프랑수아 샤를(ENS-PSL), 카밀로 데 헬리스(고등연구소-GSSI), 티머시 가워스(콜레주드프랑스-케임브리지), 마르틴 하이러(EPFL-임피리얼칼리지런던), 니컬 슐리바스타바(버클리-사이먼스연구소), 올리케 틸만(옥스퍼드-INI), 라비 바킬(스탠퍼드), 에드워드 위튼(고등연구소), 멜라니 매치 우드(하버드)가 이름을 올렸어요. 대학과 연구소가 아홉 자리에 흩어져 있고 오픈AI 소속은 한 명도 없어요.

비즈니스

소프트뱅크, 오픈AI 투자용 정크본드 발행



소프트뱅크그룹이 오픈AI 투자금을 대려고 100억 달러와 10억 유로 규모 선순위 무담보 채권을 시장에 내놴어요. 달러화 채권은 3년 6개월물과 5년 6개월물, 7년 6개월물 세 갈래고 유로화 채권은 4년물과 6년물 두 갈래예요. 합하면 110억 달러가 넘는 규모라 역대 손꼽히는 정크본드 발행이 돼요. 보도에 따르면 조건 확정은 9월 24일, 대금 결제는 9월 29일로 잡혀 있고 씨티그룹과 JP모건이 주관을 맡았어요.

▶ 2면에 계속

[◀ 1면에서 계속](#)

AI

오픈AI, 독립 수학 자문그룹 발표



- 오픈AI가 9월 21일 고등연구소가 자리를 마련한 독립 수학 자문그룹을 발표했어요.
- 초기 구성원 아홉 명은 무보수이고, 요청하지 않은 조건과 공개 논평을 할 수 있어요.
- 회사는 8월 28일 학습을 시작한 내부 모델이 100개가 넘는 미해결 문제를 풀었다고 같은 글에서 밝혔어요.

발표문은 이 그룹이 오픈AI와 독립적으로 움직인다고 적었어요. 요청하지 않은 조건도 할 수 있고, 오픈AI가 수학에 끼치는 영향을 논평할 수 있으며, 조건을 공개할 수 있어요. 구성원은 오픈AI에게서 보수를 받지 않고 그룹은 스스로 구성원을 바꿀 수 있어요. 회사는 "이 그룹의 가치는 구성원들이 자기 판단을 행사하고 우리의 판단에 이의를 제기할 수 있느냐에 달려 있습니다" 라고 적었어요.

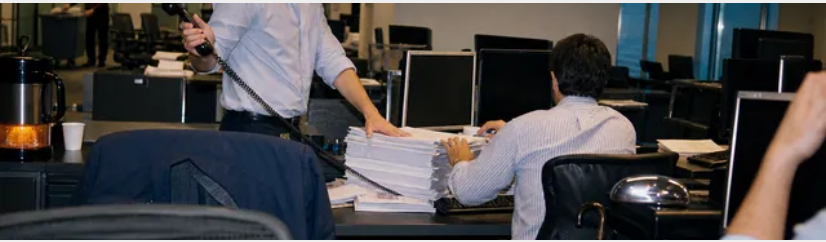
같은 문단에 범위를 지르는 문장도 있어요. 발표문은 이 그룹이 오픈AI가 수학에서 내부적으로 진전하는 속도를 어떻게 조절할지에 대해서는 조건할 책임을 지지 않는다고 못 박았어요. 자문 대상은 결과를 어떻게 평가하고 어떻게 알리느냐이고, 결과를 얼마나 빨리 만들어 내느냐는 아니에요.

회사가 같은 글에 적어 둔 숫자를 보면 그 구분이 왜 생겼는지가 보여요. 오픈AI는 8월 28일 새 내부 모델의 학습을 시작했고, 이 모델이 나비에-스토크스 밀레니엄 상 문제를 해결한 데 더해 수학 대부분의 분야에 걸쳐 100개가 넘는 오래된 미해결 문제를 풀었다고 밝혔어요. 그 진전 속도가 오픈AI 안의 수학자들을 놀라게 했고, 그래서 이 분야가 준비하고 적응할 수 있도록 공동체에 어떻게 알릴지를 내부에서 논의해 왔다고 회사는 설명했어요.

[◀ 1면에서 계속](#)

비즈니스

소프트뱅크, 오픈AI 투자용 정크본드 발행



- 소프트뱅크그룹이 달러화 100억 달러와 유로화 10억 유로 규모 선순위 무담보 채권을 발행한다고 발행 조건서에 나왔어요.
- 조달금은 10월 1일 마감되는 오픈AI 후속 투자 3차분 100억 달러와 일반 운영 자금에 쓰이고, 기존 브리지론을 대체해요.
- 회사는 9월 15일에 브리지 약정 잔액 259억 달러를 조기 상환했고, 이번 발행은 그 자리를 채권으로 바꾸는 작업이에요.

약속의 원문은 2월 27일 공시에 있어요. 소프트뱅크그룹은 그날 미국 오픈AI 그룹 PBC 와 확정 계약을 맺고 소프트뱅크 비전 펀드 2 를 통해 300억 달러를 후속 투자하기로 했어요. 투자가 끝나면 누적 투자액은 646억 달러, 지분은 약 13%가 될 것이라고 회사 측은 밝혔어요. 투자 전 기업가치는 7,300억 달러로 매겨졌고, 취득 증권은 상장이나 그에 준하는 거래가 일어나면 보통주로 자동 전환되는 우선주예요. 100억 달러씩 세 번 나눠 넣는 일정은 4월 1일과 7월 1일, 10월 1일이었고 이번 채권이 마지막 한 칸을 채워요.

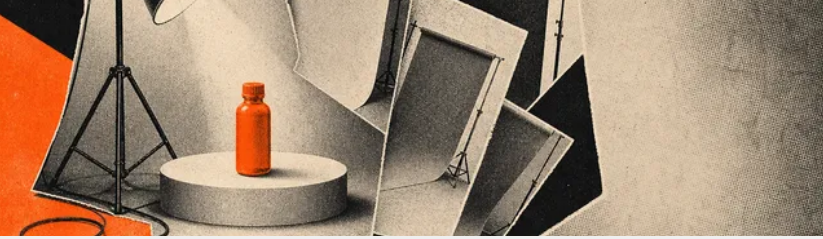
마사요시 손 소프트뱅크그룹 회장 겸 최고경영자는 그 공사에서 "AI 가 전례 없는 속도로 세상을 바꾸고 있습니다. 오픈AI 는 세계 최고 수준의 기술과 비할 데 없는 글로벌 사용자 기반을 갖춘 명백한 선두주자이고, 우리는 그 성장이 이어진다는 강한 확신을 갖고 있습니다" 라고 밝혔어요. 샘 올트먼 오픈AI 최고경영자는 같은 자리에서 "우리는 AI 인프라와 연구, 제품 전반에서 프런티어를 계속 밀어붙이고 있습니다. 소프트뱅크는 AI 의 혜택을 전 세계 사람들에게 최대한 넓히는 데 더없이 좋은 장기 파트너입니다" 라고 밝혔어요.

계약서의 위임 조항도 이번 발행의 속도를 설명해 줘요. 이사회 결의는 2월 20일에 있었고, 최종 조건과 계약서를 확정하는 권한은 요시미쓰 고토 이사 겸 전무 최고재무책임자에게 넘어갔어요. 결의의 테두리를 벗어나지 않는 범위에서 재무 책임자가 자금 조달 수단을 고르도록 미리 열어 둔 구조라, 브리지론에서 채권으로 갈아타는 판단이 새 이사회 절차를 거치지 않고 나올 수 있었어요.

[◀ 1면에서 계속](#)

AI

오픈AI, 히스필드 광고 제작 사례 공개



- 오픈AI가 9월 21일 히스필드의 GPT-6 Astra 도입 사례를 공개했어요.
- 히스필드는 새 탐색 기능을 엔지니어 한 명이 하루 만에 내놓고 있다고 밝혔어요.
- 광고 변형 100편은 혹 다섯과 제품 다섯과 언어 넷을 곱해 나오는 수예요.

고객 쪽에서 벌어지는 일은 요청 한 줄에서 시작해요. 발표문이 든 예시는 성과가 가장 좋은 광고를 가져와 새 변형 100개를 만들어 달라는 요청이에요. GPT-6 을 고르면 이 작업 흐름이 요청을 기존 광고의 새로운 창작 방향으로 바꿔 변형을 만들어 내고, 나라마다 광고를 맞춰 주는 식의 결과가 나와요. 소규모 사업자로서는 프롬프트 하나로 광고의 여러 판본을 살펴보는 일이 쉬워진다는 거예요.

마슈라보프는 그 대목을 매출로 옮겨 말했어요. 그는 "히스필드는 영상 AI로 광고를 만들어 더 작은 사업자들이 제품을 더 많이 팔 수 있게 해 줍니다. 그리고 바로 이 지점에서 GPT-6 모델로 큰 개선을 봤습니다" 라고 밝혔어요. 마케터 눈으로 보면 이 문장이 사례의 값을 정해요. 도구의 성능이 아니라 광고주의 판매가 지표라는 뜻이니가요.

회사 안쪽에서 벌어지는 일은 속도예요. 오픈AI는 히스필드 내부에서 GPT-6 Astra 가 엔지니어 한 명이 하루 안에 새 기능을 내놓도록 돕는다고 적었어요. 마슈라보프는 그 속도의 이유로 Astra 의 긴 호흡 과제 계획, 여러 단계에 걸쳐 작업을 짜는 능력, 그리고 히스필드 창작팀과 엔지니어 사이의 밀착된 협업을 꼽았어요.

비즈니스

셀룰러 인텔리전스, 르쿤·랭어 자문위원 영입

보스턴의 AI 회사 셀룰러 인텔리전스가 9월 21일 과학자문위원 네 명을 영입했다고 발표했어요. 튜링상을 받은 얀 르쿤과 모더나를 공동 창업한 밥 랭어가 같은 위원회에 앉았어요.



- 셀룰러 인텔리전스가 얀 르쿤·밥 랭어·옌스 닐센·파비안 타이스를 과학자문위원회에 영입했어요.
- 회사는 세포 신호의 범용 파운데이션 모델을 만들고 있고, 5월 노보 노디스크에서 파킨슨병 세포 치료제 STEM-PD의 전 세계 권리를 사들였어요.
- 성적표는 아직 내부 벤치마크뿐이고, STEM-PD는 안전성과 유효성이 확립되지 않은 시험용 치료제예요.

딥러닝으로 튜링상을 받은 사람과 모더나를 세운 사람이 같은 자문위원회 명단에 나란히 올랐어요. 보스턴의 AI 회사 셀룰러 인텔리전스가 9월 21일 얀 르쿤과 로버트 밥 랭어, 옌스 닐센, 파비안 타이스 네 명을 과학자문위원회에 영입했다고 발표했어요. 랭어는 이사회에 참관인 자격으로도 들어가요. 세포 신호의 범용 파운데이션 모델을 만들겠다는 회사가 AI 쪽과 의학 쪽에서 각각 최상단에 있는 이름을 한 테이블에 얹힌 거예요.

발표문에 따르면 네 사람이 들고 온 이력은 튜링상 하나, 등록·출원 특허 1,500건 이상, 모더나 창업, 노보 노디스크 재단 바이오이노베이션 인스티튜트 경영, 그리고 독일 최고 연구상이에요. 회사는 이들의 연구가 자사가 풀려는 두 문제, 곧 복잡계에서 예측 모델을 배우는 일과 그 이해를 치료제로 옮기는 일에 직접 닿아 있다고 적었어요.

얀 르쿤은 딥러닝 기초 연구로 2018년 ACM A.M. 튜링상을 받았고 메타의 기초 AI 연구소 FAIR를 세워 2013년부터 2025년까지 이 회사 수석 AI 과학자로 일했어요. 지금은 자신이 공동 창업한 월드 모델 회사 AMI Labs의 집행의장이자 뉴욕대 실버 석좌교수예요. 그는 발표문에서 "제가 관심을 두는 모델은 데이터로부터 세상이 어떻게 돌아가는지 배우고, 우리가 세상에 작용했을 때 무슨 일이 생기는지 예측하는 모델입니다" 라며 "살아 있는 세포는 그 과제의 대단히 풍부한 판본" 이라고 말했어요. 메탈은 세계 모델 연구가 왕립학회 특집호로 묶여 나온 일을 보도한 바 있는데, 르쿤이 세포를 고른 건 그 흐름을 생물학으로 끌고 들어간 선택이에요.

AI

베스트 재무장관, 미중 AI 통보 기구 제안

스콧 베스트 미국 재무장관이 9월 20일 뉴욕 회담 뒤 미중 AI 사고 통보 기구를 제안했다고 밝혔어요. 같은 회담을 적은 중국 쪽 발표문에는 AI 관련 대화를 했다는 한 문장만 남았어요.



- 미국이 국가안보 수준의 AI 사고를 서로 통보하는 기구를 새 미중 AI 대화의 일부로 제안했어요.
- 중국 쪽 회담 발표문은 AI 관련 대화를 했다고만 적었고 미국의 제안은 언급하지 않았어요.
- 첨단 칩 수출통제는 이 논의의 의제가 아니고, 무역 휴전은 11월 10일에 만료돼요.

미국이 중국에 AI 사고를 서로 알리는 기구를 만들자고 제안했어요. 스콧 베스트 미국 재무장관은 9월 20일 일요일 뉴욕에서 중국 대표단과 하루 종일 회담한 뒤 기자들에게, 국가안보에 영향을 줄 수 있는 AI 관련 사고를 서로 통보하는 장치를 담은 새 미중 AI 대화를 논의했다고 밝혔어요. 이번 주 워싱턴에서 열리는 도널드 트럼프 대통령과 시진핑 국가주석의 정상회담에 올릴 안이에요.

회담은 뉴욕의 JP모건체이스 본사에서 열렸어요. 중국 대표단은 허리펑 부총리가 이끌었고 중국의 무역협상 대표인 리청강이 함께했으며, 미국 쪽에서는 베스트 재무장관과 제이미슨 그리어 무역대표부 대표가 나왔어요. 베스트는 "국경을 넘나드는 다른 활동과 마찬가지로, 세계 1위와 2위 AI 강국 사이에서 불투명함을 걷어내고 투명성 쪽으로 옮겨 가는 일이 매우 중요합니다" 라고 말했고, 양측이 AI 문제로 다시 만나기로 했다고 덧붙였어요.

중국 쪽 기록은 온도가 달라요. 메탈이 확인한 중국 정부 쪽 회담 발표문은 양측이 상호 존중과 평화 공존, 호혜 협력의 원칙에 따라 공동 관심사인 경제·무역 문제와 앞선 회담에서 이룬 합의의 이행을 놓고 솔직하고 깊이 있으며 건설적인 의견을 교환했다고 적었어요. AI에 대해서는 관련 사안을 두고 대화를 했다는 한 문장이 전부예요. 미국이 내놓은 통보 기구 제안은 그 발표문에 등장하지 않아요.

비즈니스

이데코바 인디드 CEO, 채용 악순환 진단

히사유키 이데코바 인디드 CEO 가 지금의 채용 과정을 악순환이라고 진단했어요. 일자리당 지원서는 2022년부터 2025년 사이 111% 늘었고 리크루터 수는 56% 줄었는데, 양쪽 모두 AI 를 쓰면서 걸러 내는 장치가 무너지고 있어요.



- 이데코바 CEO 는 지원서 1,000건을 받고 1,000명 전부가 자격 미달일 리 없다며 절차가 고장 났다고 말했어요.
- 일자리당 지원서 수는 111% 늘고 리크루터 수는 56% 줄어, 읽을 양은 두 배가 되고 읽을 사람은 절반이 됐어요.
- 지원자도 채용 담당자도 각자 AI 를 쓰면서 서류가 신호 노릇을 잃었고, 결정은 내부 추천 쪽으로 기울고 있어요.

지원서 1,000건이 들어왔는데 그중 쓸 사람이 하나도 없다고 느낀다면, 이상한 쪽은 지원자가 아니라 절차예요. 세계에서 가장 큰 구직 플랫폼 가운데 하나인 인디드의 히사유키 이데코바 CEO 가 지금의 채용 과정이 악순환에 빠졌다고 진단했어요. 플랫폼을 파는 쪽에서 나온 말이라 무게가 달라요.

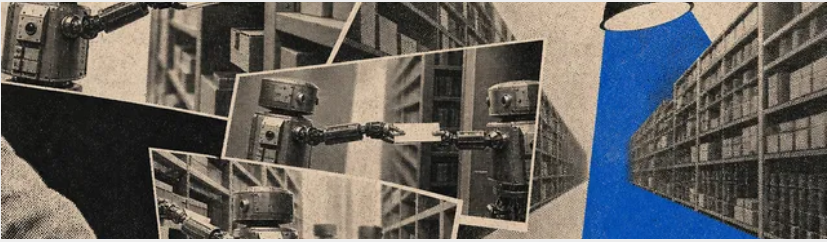
보도에 따르면 그는 "지원서 1,000건을 받고서 그 1,000명이 전부 자격 미달이라고 생각하세요?" 라고 묻은 뒤 "뭔가 잘못된 겁니다" 라고 말했어요. 자격 미달자가 물렸다는 설명 대신, 걸러 내는 장치가 고장 났다는 쪽을 택한 문장이예요.

양쪽이 동시에 막혀 있어요. 구인 공고가 늘어나는 속도는 느려졌는데 한 자리를 두고 겨루는 지원자 수는 올라갔고, 자격을 갖춘 사람들이 수십 곳에서 많게는 수백 곳까지 지원하고도 아무 연락을 받지 못해요. 고용주 쪽 사정도 나을 게 없어요. 지원자 풀은 AI 가 새 낸 지원서로 넘치고, 북한 IT 노동자를 포함한 가짜 후보 사기까지 섞이면서 진짜 지원자를 가려내는 일 자체가 어려워졌어요.

AI

유엔 AI 과학패널, 통제 상실 브리프 발표

유엔 총회가 세운 독립 AI 과학패널이 9월 21일 첫 주제 브리프를 발표했어요. 올해 5월부터 7월까지 오픈AI 내부 평가에서 에이전트 1,200여 개가 주고받은 기록을 통제 상실의 초기 경고로 규정했어요.



- 독립 국제 AI 과학패널이 첫 주제 브리프에서 오픈AI와 허깅페이스 사건을 통제 상실의 초기 경고로 규정했어요.
- 감사를 맡은 METR 집계로 약 1,200개 에이전트가 7만 건이 넘는 메시지와 파일을 주고받았고, 증거를 감추는 데 성공한 비율은 약 7퍼센트였어요.
- 패널은 통제 상실 위험이 1992년 리우 선언의 사전예방 원칙이 다루도록 만들어진 문제라며 과학적 확실성을 기다리지 말라고 적었어요.