



사전 시청 자료 · 두 번째 찾자리

알파고 너머의 세상

AI는 어떻게 배우는가,
그리고 우리 아이는 어떻게 배우는가

2026. 7. 15. 수요일 · 카인드오브썸머

넷플릭스 「데미스 하사비스: 알파고 너머의 세상」
약 65분 · 모임 전 시청 부탁드립니다

여는 글

이 다큐를, 엄마의 눈으로 보면 좋겠습니다

지난 첫 자리는 이런 질문에서 시작했습니다. "데미스 하사비스가 한국에 왔는데 왜 이렇게 세상이 조용하지?" 답은 간단했어요. **누군지 모르니까.**

이번에는 한 발 더 들어가려 합니다. 이 다큐는 AI가 얼마나 똑똑해졌는지를 말하는 것처럼 보이지만, 제가 보기에 진짜 이야기는 다른 데 있습니다. **어떻게 배우는가.** 그리고 그 이야기는 이상하게도, 아이를 키우는 우리 이야기와 자꾸 겹칩니다.

보면서 이 세 가지를 눈여겨봐 주세요

하나. 정답을 알려주는 것과, 스스로 넘어지게 두는 것

AI를 가르치는 방법은 둘입니다. 정답을 일일이 주입하거나, 틀리게 두고 스스로 찾게 하거나. 딥마인드는 후자를 택했고, 게임 하나 이기게 만드는 데 **3년**이 걸렸습니다. 우리는 아이에게 어느 쪽을 하고 있을까요.

둘. 인간이 가르쳐준 걸 지웠더니, 더 강해졌습니다

알파고 제로는 인간의 기보를 **전부 삭제**하고 백지에서 다시 배웠습니다. 그리고 더 강해졌어요. 제가 물려주는 "요령"과 "정답"이, 혹시 아이의 천장이 되고 있진 않을까요.

셋. 게임에 빠진 아이를, 아무도 말리지 않았습니다

여덟 살 하사비스는 체스 상금으로 컴퓨터를 샀습니다. 부모는 말리지 않았어요. 그 아이는 훗날 노벨상을 받습니다. 우리 집 거실에서 게임하는 아이를 저는 어떤 눈으로 보고 있을까요.

물론 저도 답을 모릅니다. 첫 자리에서도 말씀드렸듯이, 제가 가진 지식이 완벽하지 않아 오히려 헛갈리게 하는 건 아닐까 걱정도 됩니다. 그래도 **모르는 채로 지나가는 것보다는, 모르는 걸 같이 이야기하는 게 낫다**고 믿습니다. 차 한 잔 앞에 두고, 90분 동안. 그거면 충분할 것 같아요.

차와 달리기와 AI는 닮았습니다.
기다림, 과정, 지속, 탐색, 그리고 자율성.

박지영 · 키즈러닝랩

한 장 요약

10년 만에 돌아온 사람

2016년 알파고가 이세돌을 이겼습니다. 10년이 지나 하사비스가 다시 한국에 왔습니다. 그 사이 그는 **노벨 화학상**을 받았고, AI는 바둑판을 떠나 **단백질과 물리 세계**로 갔습니다.



딥마인드 본사가 있는 런던 킹스크로스

언제	무슨 일이
2016년 3월	알파고 vs 이세돌 · 알파고 4승 1패
2국 · 37수	3천 년간 아무도 두지 않던 수. 해설자도 "실수"라 여겼다
4국 · 78수	이세돌의 승리. 인간이 AI를 이긴 마지막 순간
그날 이후	서울에서 돌아온 날, 바로 알파폴드를 시작했다
2024년	단백질 구조를 푼 공로로 노벨 화학상
2026년 4월	10년 만에 방한. 이세돌과 재회

이 다큐의 진짜 질문

"AI가 얼마나 세졌나"가 아니라
"어떻게 배우는가"입니다.

정답을 주입받은 AI는 정답이 있는 문제만 풀었습니다.

정답을 지우고 스스로 넘어지게 두었더니, 인간이 3천 년간 못 본 수를 두었습니다.

01 · 한 아이가 어떻게 자랐나

여덟 살이 상금으로 컴퓨터를 샀습니다

4살	8살	11살	17살	18살
체스 시작	상금으로 컴퓨터 구입	독학 코딩 AI 프로그래밍	게임 「테마파크」 AI 개발	케임브리지 입학

부모는 말리지 않았습다. 그 컴퓨터로 책을 사서 **혼자 코딩을 배우고**, 오셀로 게임을 직접 만들었습니다. 하사비스는 그 경험이 인생을 바꿔놓았다고 말합니다.

17살에 만든 게임 「테마파크」는 전 세계에서 **수백만 장**이 팔렸습니다. 놀이공원을 유저가 설계하는 게임인데, 여기서 이미 AI를 썼습니다. 손님들이 화를 내고, 배고파하고, 반응합니다. 햄버거를 먹고 바로 놀이기구를 타면 속이 안 좋아지니까, 동선까지 계산해서 설계해야 했죠.

"인공지능을 기반으로 스스로 생각하는 존재들이 돌아다니는 그런 게임을 추구했던 것 같아요."

다큐 내레이션

게임 회사가 거액의 연봉으로 붙잡으려 했지만 실패했습니다. 케임브리지 입학허가를 받았거든요.

그가 파고든 질문

"생각이란 도대체 뭘까." 하사비스가 찾은 답은 **기억과 상상**이었습니다. 인간의 지능을 닮은 시를 만들려면 반드시 풀어야 할 조건이었죠.

"미래를 시뮬레이션으로 그려볼 수 있는 것.

무슨 일이 일어날지에 대한 시나리오를 계획하는 것."

데미스 하사비스

02 · 이 다큐의 핵심 개념

AI를 가르치는 두 가지 방법

방법 1

지도학습

고양이 사진과 강아지 사진을 주고 어떤 게 고양이인지 정답을 알려준다. 수만 번 반복하면 구분해낸다.

한계: 정답을 모르는 문제는 아예 못 푼다.

방법 2

강화학습

정답을 알려주지 않는다. 틀리면 경고, 잘하면 점수. 이기는 법을 스스로 찾게 한다.

대가: 오래 걸린다. 게임 하나에 3년.

"인간 아이가 배우는 방식과 비슷합니다. 걸으려다 넘어지면 아프죠. 그러면 다르게 해봐야겠다고 알게 됩니다. 그게 바로 강화학습 신호입니다. 세상과 부딪히면서 배우는 거죠."

토레 그레페 · 딥마인드 리서치 그룹 리더

알파고는 이렇게 만들어졌습니다

STEP 1

인간 기보
10만 개

지도학습

STEP 2

시끼리
서로 대국

강화학습 시작

STEP 3

자가대국
3천만 번

인간이 평생 두는
양을 아득히 넘음

가장 충격적인 대목 — 알파고 제로

인간의 지식을 지웠더니, 더 강해졌습니다.

하사비스는 이미 인간을 이긴 알파고에서 인간 기보를 전부 삭제했습니다. 백지에서 자기 자신하고만 두게 했죠. 그 과정을 18번 반복해 나온 것이 알파고 제로. 나아가 체스와 장기까지 스스로 배우는 알파 제로가 되었습니다.

02-1 · 우리가 이미 알고 있던 것

그런데 이걸, 우리가 달리기에서 하던 겁니다

강화학습 이야기를 들으며 저는 계속 달리기가 떠올랐습니다. 딥마인드가 어렵게 발견한 그 방법을, 우리는 이미 아이들과 트랙에서 하고 있었거든요.

"이 세상에 있는 대부분의 문제는, 우리가 **정답을 모르면** 지도학습을 쓸 수가 없어요."

다큐 인터뷰

강화학습은	달리기도
정답을 알려주지 않는다	뛰는 법은 말로 가르칠 수 없다 . 직접 뛰어봐야 안다
넘어져야 신호를 얻는다	숨이 차봐야 자기 페이스를 안다
보상이 늦게 온다	오늘 뒀던 것은 3개월 뒤에 몸에 나타난다
3천만 번의 반복	매일의 누적 . 하루로는 아무것도 안 바뀐다
시뮬레이터에서 만 번 넘어진다	대회 전에 연습에서 실패해본다

한 문장으로

**강화학습은 결국,
몸으로 배우는 법입니다.**

AI에게 정답을 주입하면 정답이 있는 문제만 풀립니다. 아이에게 요령을 주입하면 요령이 통하는 데까지만 갑니다. **넘어질 자유**를 줘야 그 너머로 갑니다. 달리기가 아이에게 주는 것도 정확히 그것이었습시다.

다만, 한 가지 차이

알파고는 **3천만 번** 넘어질 수 있었습니다. 지치지도, 상처받지도 않으니까요. 우리 아이는 그렇지 않습니다. 그래서 우리에게서 AI에게 없는 일이 하나 더 있습니다. **넘어진 아이 옆에 있어 주는 일**. 이세돌도 3연패 하고 새벽 5시까지 복기할 때 혼자가 아니었습니다. 동료 기사들이 방에 함께 있었고, 다음 날 78수를 뒀습니다.

차와 달리기와 AI는 닮았습니다.

기다림, 과정, 지속, 탐색, 그리고 자율성.

이번 찾자리의 주제

03 · 2016년, 서울

37수와 78수

37수 — 인간의 고정관념이 무너진 순간

2국에서 알파고가 둔 **어깨짚기**. 바둑에서는 "종지 않다"는 고정관념이 있어서 프로 기사가 그 자리에 둔 적이 없었습니다. 영어 해설자는 그 수를 판에 올려놓고 **실수라고 생각했습니다**.

"인간의 바둑에서 우리가 본 그 무엇과도 다른, 결정적인 한 수였습니다.
그리고 그것이 새로운 시대를 열었습니다."

토레 그레페 · 당시 알파고 기술팀

알고 보니 실수가 아니었습니다. 알파고는 중앙 진출에 유리한 형세를 계산하고, 승리에 꼭 필요한 수를 **미리** 둔 것이었죠.

"확률 계산이나 하는 기계에 불과하다고 생각했는데, 아니었구나.
알파고도 충분히 창의적이다."

한국 해설진

78수 — 인간이 AI에게 거둔 마지막 승리

3연패. 비관론이 팽배했습니다. 그날 밤 이세돌은 호텔 방에서 프로 기사들과 **새벽 5시까지 복기**했습니다. 분명 약점이 있을 것이다. 그게 뭘까.

4국에서 둔 78수. 알파고는 너무 많은 경우의 수를 계산하다 스스로 혼란에 빠졌고, 크게 흔들렸습니다.

"78수는 이세돌의 천재적인 수였고, 그는 여전히 공식 대국에서 알파고를 이긴 세계 유일의 인간입니다."

데미스 하사비스

10년 후, 바둑은 어떻게 됐나

AI는 적이 아니라 **훈련 파트너**가 됐습니다. 국가대표 기사들은 AI로 훈련하고, 인간의 실력은 오히려 더 높은 차원으로 올라갔습니다. "우리가 알던 바둑은 **빙산의 일각**이었다는 걸 느꼈습니다."

04 · 바둑판을 떠나 생명으로

알파폴드, 그리고 노벨상

"알파고 대국이 끝나고 서울에서 돌아온 바로 그날, 저는 알파폴드 프로젝트를 시작했습니다. 우리가 준비됐다는 걸 알았으니까요."

데미스 하사비스



단백질의 3차원 접힘 구조

우리 몸의 모든 생명현상은 단백질이 합니다. 그리고 **단백질은 생김새가 곧 기능**입니다. 산소를 나르는 헤모글로빈은 산소를 붙들기 좋은 모양이고, 항체는 병원체에 결합하기 좋은 구조죠. 원하는 모양을 만들 수 있다면 **암, 비만, 노화**의 열쇠가 됩니다.

	경우의 수
바둑	약 10의 170 제곱
단백질 접힘	약 10의 300 제곱

아미노산 100개짜리 단백질 하나만 해도, 가능한 구조를 다 탐색하려면 **우주의 나이보다 긴 시간**이 필요합니다. "바둑은 룰이 명확한 게임입니다. 단백질 구조 문제는 **룰을 모르는 거죠.**"

CASP 대회 (단백질 올림픽)	점수
2018년 이전 · 기존 연구	약 40점
알파폴드 1	약 60점
알파폴드 2	약 90점 · 거의 완벽

노벨상을 받은 진짜 이유

결과를 독점하지 않았습니다.

딥마인드는 시로 푼 단백질 구조를 **전 세계에 무료 공개**했습니다. 지금 2억 5천만 개. 서울대 백민경 교수는 여기서 힌트를 얻어 **8개월 만에** 로제타폴드를 만들었고, 한국 스타트업 갤럭스는 항체 설계에 성공해 **680억 원**을 유치했습니다.

05 · 다음 10년

월드모델 — AI의 연습장

2022~

생성형 AI

정보를
알려준다

NOW

에이전틱 AI

직접 버튼 누르고
예약까지 해준다

NEXT

피지컬 AI

로봇.
물질 세계에서 실행

월드모델은 물리 세계가 어떻게 작동하는지를 배우는 AI입니다. 언어도 아니고 특정 분야도 아닌, **세상의 물리 법칙 자체**를 모델링합니다.

"제가 해온 모든 영역이 하나로 모이는 정점입니다.
컴퓨터 게임에서 시작해 이제 AI까지, 그 전부를 합치는 거죠."

데미스 하사비스

실제로 어떤 모습인가

할리우드 모션그래픽 디자이너 김규륜(에미상 수상)이 시연한 구글 월드모델은, 프롬프트로 만든 세계 안에서 **캐릭터를 실시간으로 직접 조종**할 수 있습니다. 물바다가 된 광화문에 골든 리트리버를 놓으면, 개가 물을 헤치고 계단을 타고 올라갑니다. **누가 시키지 않아도 환경과 상호작용**합니다.



왜 중요한가 — 로봇의 연습장

트럼프 카드 한 장 집기. 3~4살 아이도 쉽게 하지만 로봇에겐 어렵습니다. 로봇도 알파고처럼 **강화학습**으로 배웁니다. 즉 수없이 실패 해봐야 하는데, 실제 로봇은 한 대에 5천만 원이 넘습니다.

"시뮬레이터 안에서는 천 개, 만 개의 로봇을 돌려볼 수 있어요. 학습 효율을 극대화할 수 있죠."

박대현 교수 · 카이스트

"월드모델은 사람을 위한 도구라기보다,
기계가 세상을 이해하도록 돕는 가상세계입니다."

김규륜 : 모션그래픽 디자이너

06 · 그가 걱정하는 것

기술의 중심에는 여전히 사람이 있어야 합니다

바둑을 정복한 알파고가 모든 게임을 배우는 알파 제로로 진화했듯, 월드모델을 진화시켜 인간처럼 물리 세계를 이해하는 AGI를 만드는 것. 그것이 딥마인드가 그리는 그림입니다.

한국의 자리

하드웨어 강국(중국)과 소프트웨어 강국(미국) 사이에 끼여 있습니다. 카이스트 박대현 교수는 한국형 월드모델 개발을 정부에 제안했습니다. 10년 전 알파고 충격 이후 가속화된 피지컬 AI 경쟁에서 살아남으려면 **우리만의 월드모델이 필수 조건**이라는 겁니다.

하사비스가 우려하는 두 가지

무엇	왜 위험한가
오용 Misuse	나쁜 행위자가 이 기술을 해로운 목적 으로 바꿔 쓰는 것
정렬 실패 Misalignment	시스템이 강력해지고 자율적이 될수록, 우리가 정해진 가드레일 안에 머물게 하기가 어려워진다

딥마인드에는 이런 사람들이 있습니다

철학자, 아동심리학자, 윤리학자.

기술팀과 **함께** 일합니다. 처음부터 사회와 인문학을 염두에 두고 AI를 개발하기 위해서라고 합니다.

"우리는 AI가 **현미경이나 망원경** 같은 것이 되길 바랍니다.
우리를 둘러싼 우주를 이해하고 더 나은 미래를 짓도록 돕는.

그리고 기술의 중심에는 **여전히 사람이 있어야 합니다.**"

데미스 하사비스 · 다큐의 마지막 메시지



알파고는 **3천만 번**을 뒀습니다.
이세돌은 **밤을** 새웠습니다.
딥마인드는 **3년**을 기다렸습니다.

잘하는 사람도, 똑똑한 기계도
처음부터 잘한 것은 아니었습니다.

차를 우릴 때 기다리는 것처럼,
달리기를 할 때 숨이 차는 것처럼,
배우는 데는 시간이 걸립니다.

7월 15일, 카인드오브쌔머에서 뵈겠습니다.

저작권 및 이용 안내

본 자료의 저작권은 키즈러닝랩(박지영)에 있습니다.
찾자리 참석자 **본인의 학습 목적으로만** 사용해 주세요.

외부 공유, 재배포, 상업적 이용, SNS 게시를 금지합니다.
단톡방 전달이나 캡처 공유도 삼가 주시길 부탁드립니다.

영상 내용의 인용은 시청자의 이해를 돕기 위한 요약이며,
원 영상의 저작권은 해당 제작사에 있습니다.