

4차산업혁명센터

Master's Insight

AI, 인류의 구원인가 파멸인가? (Will AI Save the World or End it?)

‘AI의 대부’ 제프리 힌턴 교수(2024년 노벨 물리학상 수상자)와의 인터뷰

2025년 8월 18일

C4IR Korea Global Intelligence Hub (koreago.net)

본 자료는 2024년 노벨물리학상 수상자이자 ‘AI의 대부’로 불리는 제프리 힌턴(Geoffrey Hinton) 캐나다 토론토대 컴퓨터과학과 명예교수가 캐나다 TVO 방송과 2025년 4월 인터뷰한 내용을 한국(경기도) 4차산업혁명센터가 번역하여 요약 정리한 것입니다(자료원 : TVO Today 유튜브 채널). 대담자의 견해는 전적으로 대담자 개인의 의견이며, 한국(경기도) 4차산업혁명센터의 공식 입장이 아닙니다.



[Geoffrey Hinton: Will AI Save the World or End it? | The Agenda - 원본 동영상 보기](#)

- AI는 인간의 뇌를 모방한 시스템으로, 엄청난 가능성을 지닌 동시에 인류 멸종까지 초래할 수 있는 실존적 위협이기도 하다.
- AI 안전성을 확보하려면 지금 당장 글로벌 공감대를 형성하고, 대기업과 정부가 공동으로 대응과 연구에 나서야 한다.
- AI는 의료와 교육 등 삶의 질을 획기적으로 높일 수 있는 도구인 만큼, 위험을 통제하며 책임 있게 발전시켜야 한다.

요약 (Executive Summary)

주제: 인공지능의 약속과 위협, 그리고 사회적 대응 방안

출처: TVO 방송 인터뷰 (진행자: Steve Paikin)

출연자: Geoffrey Hinton (AI의 대부, 토론토대학교 명예교수, 노벨 물리학상 수상자)

1. 🧠 인공지능의 본질과 연구 동기

신경망의 원리:

인간의 뇌는 뉴런(신경세포) 간 연결 강도를 변화시키며 학습함.

AI도 이 원리를 모방한 인공신경망으로 구성되어, 데이터 기반으로 스스로 학습함.

연구를 선택한 이유:

인간을 이해하려면 뇌를 이해해야 하며, **"뇌는 아직 미지의 영역"**이라는 인식이 연구의 출발점.

2. ⚠️ AI의 위험: 단기와 장기

단기 위험 (이미 현실화):

AI를 악의적으로 활용하는 사용자들:

- 예) 선거 조작, 딥페이크, 조작 영상, 개인 정보 타겟팅

피싱 공격 1,200% 증가 (2023~2024):

- LLM 기반의 정교한 언어 표현력 탓에 구분이 어려워짐

장기 위험 (실존적 위험):

AI가 인간보다 더 똑똑해질 가능성은 대부분의 전문가들이 동의

- 시점은 1년~20년까지 다양한 의견

AI가 통제권을 장악할 가능성도 존재 (개인적으로 10~20% 정도로 확률 예상. 단 정확한 추정은 불가)

- 바이러스 설계 등 다양한 시나리오가 가능
- “우리는 그 이후 어떤 일이 벌어질지 알지 못한다”는 점이 가장 큰 문제

3. 사회적·정치적 대응의 필요성

우선순위: 사회적 합의(consensus) 형성

기후변화 대응과 유사:

“문제가 실제로 존재한다는 데 대한 대중의 공감”이 전제되어야 정책 실행 가능

대기업의 책임:

AI 기업들은 단기 이익보다 안전성 연구에 최소 1/3 이상 자원 투자 필요

현재는 턱없이 부족한 수준

국제 협력 가능성:

AI가 인류 전체를 위협하는 만큼, 국가 간 협력은 충분히 가능

- “중국도, 트럼프도 멸종을 원하진 않는다”

4. AI의 긍정적 가능성

의료 분야:

AI 주치의는 수억 명의 데이터를 기반으로 정밀한 진단 가능

의사와 AI가 협업하면 진단 오류율 현저히 감소

교육 분야:

AI 튜터는 학습자의 오해나 약점을 정확히 파악해 개인화된 학습 제공

인간 과외보다 3~4배 더 효율적인 교육 기대

향후 10년 이내 현실화 가능성 높음

5. 📖 개인사와 철학적 시사점

노벨상과 기부:

노벨 물리학상 수령액의 절반을 'Water First'에 기부

- 캐나다 원주민 지역의 안전한 식수 공급 프로젝트
- "깨끗한 물은 인간의 기본권"

일론 머스크와의 충돌:

트위터/X 상에서 갈등

- 머스크의 "과학 기관 파괴"와 "정부 인력 대량 해고"를 **"**끔찍한 일**"**이라 지적

📌 결론

Geoffrey Hinton 교수는 AI의 놀라운 잠재력을 인정하면서도,

그 실존적 위협에 대한 국제적 경계와 윤리적 통제의 필요성을 강하게 호소한다.

AI의 발전은 더 이상 기술자의 문제가 아니라,

인류 전체가 함께 논의하고 책임져야 할 정치적·사회적 과제라는 것이 그의 핵심 메시지다.

인터뷰 전체 번역본

(지엽적 내용 일부 제외하고 최대한 전체 번역)

("AI의 대부" Geoffrey Hinton 교수, TVO 출연)

진행자 스티브 (Steve):

다음 초대 손님은 "AI의 대부"로 불리는 인물입니다. 최근 인공지능 분야에서의 선구적인 연구로 노벨 물리학상을 수상하셨죠. 일론 머스크의 심기를 건드리기도 한 인물이기도 합니다. 지금 스튜디오에 함께한 분은 토론토대학교 컴퓨터 과학 명예교수인 제프리 힌튼(Geoffrey Hinton)입니다. 오늘은 그와 함께 고도화된 인공지능의 가능성과 위험에 대해 이야기 나눠보겠습니다. 다시 만나 뵙게 되어 반갑습니다.

Geoffrey:

초대해 주셔서 감사합니다, 스티브.

🎥 스톡홀름 노벨상 수상 장면

Steve:

처음부터 아주 인상적인 장면 하나를 보여드릴게요. 제 생각엔 당신 인생에서 가장 멋진 날 중 하나였을 것 같네요. 쉘든, 영상 틀어주시죠.

(영상: 스웨덴의 칼 16세 국왕이 힌턴 교수에게 노벨상을 수여하는 장면)

Steve:

영상 속 인물은 스웨덴의 칼 16세 국왕이고, 당신은 스톡홀름에서 노벨상을 받고 있었습니다. 그 감격이 가신 건가요?

Geoffrey:

그렇게 되면 알려드릴게요.

Steve:

아직 완전히 가시지 않았군요?

Geoffrey:

네, 완전히는 아니에요.

Steve:

그 날은 얼마나 대단했나요?

Geoffrey:

정말 놀라운 경험이었죠. 특히 저는 물리학자가 아닌데 물리학상을 받았다는 점에서 더 놀라웠습니다.

Steve:

그건 어떻게 된 건가요?

Geoffrey:

제 생각엔, 인공지능 분야에서의 발전에 대해 노벨상을 주고 싶었던 것 같아요. 요즘 과학계에서 가장 큰 흥미를 끄는 분야니까요. 그래서 물리학상이라는 틀을 빌려 제게 수여한 것 같습니다.

Steve:

그들에게 "저 물리학 안 하는데요?"라고 말하진 않았나요? 혹시 그들이 "선물 말고는 따지지 마라"라는 반응을 보였는지 궁금하네요.

그런데 그 메달, 어디에 보관하세요?

Geoffrey:

그건 비밀입니다.

Steve:

제가 훔치려는 건 아니에요, 힌턴 교수님.

Geoffrey:

하지만 다른 사람이 노릴 수도 있죠. 금으로 만들어졌고 무게는 6온스, 약 15,000달러 정도 가치가 있거든요.

 "신경망"은 어떻게 작동하나?

Steve:

제가 공식 수상 이유를 읽어드리겠습니다:

“인공 신경망을 통한 머신러닝의 기반적 발견과 발명”으로 수상하셨습니다.

수백만 번은 들으셨겠지만, 이게 무슨 뜻인지 설명해주시죠.

Geoffrey:

우리 뇌에는 뉴런이라는 세포가 아주 많이 있어요. 이 뉴런들은 서로 연결돼 있고, 우리가 무언가를 배울 때는 이 연결의 강도가 변화하죠.

그러니까 뇌가 어떻게 작동하는지 알려면 이 연결의 강도를 어떻게 바꾸는지를 알아야 합니다.

그게 핵심이에요. 연결을 더 강하게 할지, 약하게 할지 결정하는 뇌의 방식을 말이죠.

그리고 만약 이 과정을 컴퓨터에서 흉내 낼 수 있다면 어떻게 될까요?

우리는 이미 알고 있습니다. 굉장히 똑똑해지죠.

왜 이 연구를 선택했나?

Steve:

수천, 수만 가지 과학 분야 중 왜 하필 이것을 선택하셨나요?

Geoffrey:

이게 가장 흥미로운 분야니까요.

Steve:

당신에게 있어서?

Geoffrey:

아니요, 모두에게 가장 흥미로운 분야라고 생각해요.

사람을 이해하려면 결국 뇌가 어떻게 작동하는지를 이해해야 하니까요.

그런데 우리는 아직 뇌를 제대로 이해하지 못했죠.

예전보다는 나아졌지만, 여전히 갈 길이 멉니다.

AI의 위험: 단기와 장기

Steve:

노벨상 이전에도 유명하셨지만, 수상 이후에는 그야말로 세계적인 주목을 받게 되었죠.

이후 AI의 위험성을 경고하시면서 구글도 떠나셨는데요.

이 주제를 조금 나눠서 살펴보죠.

당장 우리가 직면한 인공지능의 단기적 위험은 무엇인가요?

Geoffrey:

위험은 두 가지 유형이 있어요.

첫째는 **악의적인 사용자**가 AI를 오용하는 경우입니다. 이걸 이미 벌어지고 있고, 단기적이고 즉각적인 위협이죠.

둘째는 **AI가 인간보다 똑똑해졌을 때** 생기는 전혀 다른 수준의 위험이에요.

그때 AI가 인간을 무시하고 스스로 통제권을 가지려고 한다면?

더 똑똑한 존재가 덜 똑똑한 존재에게 통제받는 사례를 본 적 있으신가요?

Steve:

거의 없죠.

Geoffrey:

물론 약간 더 똑똑한 사람이 덜 똑똑한 사람에게 통제받는 경우는 있어요.

하지만 그건 지능의 차이가 미미할 때의 이야기입니다.

Steve:

그럼 첫 번째, 악의적인 사용 사례부터 봅시다. 어떤 예가 있을까요?

Geoffrey:

사람들에 대한 방대한 데이터를 모아서, 그들에게 **투표를 포기하게 만들기 위한 AI 조작 영상을 타겟팅**해 보내는 겁니다.

이건 실제로 벌어질 수 있는 일이고, 지금도 벌어지고 있습니다.

예를 들어, 피싱 공격이 대표적인 사례죠.

2023년에서 2024년 사이에 피싱 공격은 무려 12배 증가했어요.

그 이유는 대규모 언어모델 덕분에 피싱 메시지가 너무 정교해졌기 때문이죠.

이전에는 맞춤법도 틀리고 번역체 문장도 많았기 때문에 쉽게 구별이 가능했는데,

지금은 **완벽한 영어**로 되어 있어서 구분하기가 어려워졌습니다.

Steve:

점점 정교해지네요.

그럼 두 번째 유형의 위험, 즉 더 똑똑한 존재가 우리를 통제하는 사례는요?

Geoffrey:

제가 아는 유일한 예는 **아기와 엄마**입니다.

진화는 아기의 울음소리를 엄마에게 참을 수 없게 만들어서 엄마를 제어하도록 했죠.

그게 거의 유일한 예입니다.

Steve:

그럼 더 장기적인 위험에 대해 얘기해 보죠. 언제쯤 그런 위험이 현실화된다고 보시나요?

Geoffrey:

장기적인 위험은 바로 **AI가 인간보다 똑똑해지는 것**입니다.

대부분의 선도적인 연구자들이 이 점에 동의하고 있어요.

다만, **언제가 될 것**이냐에 대해 의견이 갈릴 뿐입니다.

어떤 사람들은 **20년 후**라고 하고, 어떤 이들은 **3년 후**,

심지어 몇몇은 **1년 안에도 가능**하다고 생각해요.

우리는 모두 언젠가 AI가 인간보다 똑똑해질 것이라는 데 동의합니다.

하지만 그다음에 무슨 일이 벌어질지에 대해서는 아무도 모릅니다.

의견은 있겠지만, 정확한 확률을 추정할 수 있는 기반은 없습니다.
저는 대략 10~20% 확률로 AI가 통제권을 장악할 거라고 추정해요.
하지만 솔직히 말해 정확한 수치는 아닙니다.
1%보다는 크고, 99%보다는 작다는 정도죠.

Steve:

"통제권을 장악한다"는 건 곧 인류가 멸종할 수도 있다는 의미로도 들립니다.

Geoffrey:

그게 현실화된다면, 그렇게 될 수도 있죠.

Steve:

그게 언제쯤일지 시간대도 예측 가능할까요?

Geoffrey:

아니요. 그런 예측을 할 수 있는 신뢰도 있는 방법은 없어요.
하지만 지금 우리가 뭔가 조치를 취하지 않는다면, 그 시간이 오래 걸리지 않을 수도 있습니다.
현재는 그래도 "초지능 AI를 안전하게 개발할 수 있는 가능성"이 열려 있는 시점이에요.
아직 우리는 그 방법을 모르고, 가능한지조차 확실치 않지만,
가능하다고 믿고 연구를 집중해야 할 시기입니다.

Steve:

그렇다면 AI가 인류를 멸종시킨다는 시나리오가 실제로 일어난다면, 어떤 방식이 될까요?

Geoffrey:

방법은 정말 많아요. 그래서 구체적인 시나리오를 상상해보는 건 의미가 없다고 생각해요.
다만 영화 <터미네이터> 같은 방식은 아닐 것 같고,
예를 들어 치명적인 바이러스를 생성해 인류를 전멸시키는 것도 가능하겠죠.



정부와 기업은 무엇을 해야 하나?

Steve:

그렇다면 우리는 지금 AI의 안전 문제를 통제하고 있다고 볼 수 있을까요?

Geoffrey:

그랬으면 좋겠어요.

실제로 AI 안전성과 존재적 위협에 대한 연구는 진행 중입니다.

하지만 연구량이 턱없이 부족합니다.

대형 기술 기업들은 단기 수익에 몰두해 있고,

정부가 이를 강제하지 않으면 스스로 자원 투입을 하려고 하진 않죠.

우리는 정부에게 요구해야 해요.

대기업들이 전체 자원의 최소 3분의 1 정도를 AI 안전 연구에 투자하도록 말이죠.

Steve:

그 요구는 잘 작동하고 있나요?

Geoffrey:

사람들의 인식은 점점 높아지고 있고, 정치인들도 조금씩 관심을 보이고 있어요.
하지만 최근 미국에선 **한 발짝 후퇴한 상황**입니다.

Steve:

그건 어떤 상황을 말씀하시는 건가요?

Geoffrey:

바이든 행정부는 AI 안전에 관심이 있었고, 관련 **행정명령도** 발표했죠.
하지만 지금은 그 명령들이 **사라지고** 있습니다.
트럼프가 돌아오면서, 바이든의 여러 행정명령들이 모두 사라지고 있는 셈이죠.

Steve:

그런 변화가 생긴 이유는 현재 미국을 이끄는 이들이 **기술기업들과 가까운 관계**이기 때문일까요?

Geoffrey:

안타깝게도... 그렇습니다.

Steve:

이 문제를 해결하려면 결국 우리가 뭘 해야 하나요?

Geoffrey:

첫 번째는 “사회적 합의”를 만드는 일입니다.
이건 과학 소설이 아니에요. **실제의 심각한 문제**입니다.
그리고 더 많은 안전성 연구가 필요하다는 공감대를 만들어야 해요.
기후변화를 떠올려보세요.
먼저 "기후변화는 실재하며, 방치하면 끔찍한 일이 벌어진다"는 공감대가 필요했죠.
그다음에야 비로소 정부가 행동하기 시작했어요.
아직 충분하진 않지만, 그래도 **시작은** 됐습니다.
AI도 마찬가지입니다.

합의부터 시작해야 해요.

그리고 한 가지 희망적인 점은,

AI의 존재적 위협이라는 문제는 전 세계 국가들이 함께 협력할 수 있는 사안이라는 겁니다.
중국 지도자들도 멸망을 원하진 않을 테니까요.
이 문제만큼은 중국과도 협력할 수 있어야 합니다.

Steve:

"우리"라고 하셨는데, 지금은 캐나다와 미국이 더는 하나의 '우리'가 아니라고 느끼시나요?

Geoffrey:

맞아요. 예전엔 캐나다와 미국이 하나라고 생각했지만, 지금은 더는 그렇게 느껴지지 않아요.

Steve:

그렇다면 이런 국제 협력을 주도할 수 있는 **국제기구**가 있나요?

Geoffrey:

여러 단체들이 주도권을 잡으려고 하고 있지만, 아직 하나로 통일되진 않았습니니다.

Steve:

UN이 그 역할을 해야 하지 않을까요?

Geoffrey:

UN은... 좀 무력해요.

Steve:

그럼 누가 할 수 있을까요?

Geoffrey:

지금 가장 실질적인 역량을 가진 곳은 **대형 기술 기업들**이에요.

최신 AI 모델을 연구하려면 막대한 자원이 필요한데, 그걸 감당할 수 있는 건 결국 이들뿐이니까요.



일론 머스크와의 갈등, 그리고 과학의 위기

Steve:

세계에서 가장 부자인 사람, 일론 머스크에 대해 이야기해 볼까요?

그가 지금은 당신의 크리스마스 카드 목록엔 없는 것 같던데요.

Geoffrey:

저는 그와 일부 사안에선 의견을 같이합니다.

예컨대 **AI의 존재적 위협**에 대해 진지하게 받아들이는 점은 저도 동의해요.

그가 전기차나 우크라이나의 통신 인프라를 위한 **Starlink** 제공 같은 좋은 일들도 했죠.

하지만 지금 **X(옛 트위터)**에서 하고 있는 일들은 정말 끔찍하다고 생각해요.

그는 지금 **정부 기관에서 일하는 사람들을 무작위로 해고**하고 있어요.

그들은 성실하게 자신의 일을 해온 사람들인데,

그들을 향해 “무능하고, 부패하고, 게으르다”며 **모욕**을 주고 자르죠.

이건 정말 **끔찍한 결과**를 초래할 겁니다.

그리고 그는 그에 대해 **전혀 신경**을 쓰지 않는 듯해요.

그런데 그가 유일하게 신경 썼던 순간이 **제가 그를 비판했을 때**였어요.

그는 그걸 두고 “잔인하다”고 말했죠.

Steve:

당신은 그가 운영하는 X(트위터)에 직접 올리셨죠.

이렇게요:

“일론 머스크는 **영국 왕립학회에서 제명**되어야 한다.

그가 **음모론을 퍼뜨리고 나치식 경례를 했기 때문이 아니라,

미국의 과학 기관에 막대한 해악을 끼치고 있기 때문이다.

이제 그가 정말 표현의 자유를 믿는지 두고 보자.”

이에 대해 머스크는 직접 반응했죠.

“상장이나 학회 따위에 신경 쓰는 건 비겁하고 불안한 자들뿐이다.

역사는 진정한 심판자다.

당신의 발언은 무지하고, 잔인하며, 거짓이다.

그렇다면 내가 어떤 행동을 바꿔야 하는지 말해보라.

나는 실수를 하더라도 최대한 빨리 고치려고 노력할 것이다.”

이 트윗에 대해 어떤 반응을 보이셨나요?

Geoffrey:

그와 장기적인 말다툼을 벌이고 싶지는 않았어요.

왜냐하면 미국에 입국할 수 있는 권한을 유지하고 싶었거든요.

그래서 **제 친구 얀 르쿤(Yann LeCun)**이 대신 그 트윗에 답했습니다.

Steve:

그 답변은 어디서 볼 수 있죠?

Geoffrey:

트위터에서 볼 수 있습니다.

Steve:

그럼 머스크와 직접 소통한 건 그게 전부인가요?

Geoffrey:

아니요, 몇 년 전 그가 직접 저에게 전화를 요청했어요.

AI의 존재적 위협에 대해 이야기하고 싶다는,

사실은 자신의 X AI 회사의 고문이 되어달라는 제안이었죠.

Steve:

뭐라고 답하셨나요?

Geoffrey:

우리는 처음에 그 위협에 대해 잠깐 얘기했고,

그가 저의 제자 중 한 명이 X에서 기술 책임자로 일하고 있다는 걸 강조하면서 저를 설득하려 했어요.

그런데 갑자기 그가 횡설수설하기 시작해서,

“죄송합니다, 다른 미팅이 있어서 끊어야겠네요” 하고는 전화를 끊었습니다.

Steve:

그러니까 그게 끝이었군요.

머스크가 무대에서 한 행동을 나치 경례라고 표현하셨는데,

일부 사람들은 그저 “군중에게 손을 흔든 것”이라고 말하기도 해요.

그 해석에 동의하지 않으시는 거죠?

Geoffrey:

네. 특히 그의 과거 이력과 가족의 정치적 성향 등을 보면, 저는 믿지 않습니다.

Steve:

그는 특정 세력과 지나치게 가까운 행보를 보이기도 하죠.

하지만 그가 보낸 반응 중에는 그래도 구체적인 조언을 요청하는 부분도 있었잖아요?

Geoffrey:

그건 긍정적인 부분이죠.

하지만 저는 다른 사람이 답변하도록 했습니다.

Steve:

그가 어떤 행동을 바꿔야 한다고 생각하시나요?

Geoffrey:

그가 지금 추진하는 정책의 본질부터 알아야 합니다.

그는 초부유층을 위한 4조 달러 규모의 감세를 원하고 있어요.

그걸 실현하려면 정부 지출을 줄이거나 빚을 늘려야 하죠.

아니면 관세를 부과해서 일반 서민들에게 세금을 전가할 수도 있어요.

Steve:

관세는 일종의 서민세니까요.

Geoffrey:

맞아요. 관세는 역진적 세금입니다.

일반 사람들에게 모든 물가를 더 비싸게 만드는 것이죠.

결국 국민은 4조 달러 감세를 위해 더 많은 물건값을 내게 될 겁니다.

그건 정말 끔찍하고 역겨운 정책입니다.

Steve:

당신은 지금 미국 정부 정책이 과학 기관을 해치고 있다고 했는데,

구체적으로 어떤 사례를 말씀하시는 건가요?

Geoffrey:

예를 들어 RFK Jr. 같은 사람이 보건 시스템의 책임자가 되는 일이에요.

Steve:

그는 백신과 제약회사에 대한 음모론을 퍼뜨리는 인물인데,

혹시 일부 공감하는 점도 있나요?

Geoffrey:

그가 "건강한 식단이 중요하다"고 강조하는 점은 동의합니다.

물론 씨앗 기름이 독이다 같은 주장은 너무 과한 면이 있지만요.

그래도 건강한 식단이 건강을 향상시킨다는 건 중요한 사실이죠.

하지만 그 외의 주장들 대부분은 말도 안 되는 소리입니다.

Steve:

백신이 자폐증을 유발한다는 주장은 전혀 신빙성이 없다는 말씀이지요?

Geoffrey:

네, 그 문제에 대해서는 이미 많은 과학적 연구가 이루어졌습니다.

이런 주장을 펼치는 사람들 중 많은 수가 사실은 자기 제품을 팔려는 의도가 있습니다.

실제로 자신은 백신을 맞추면서도 대중에게 "백신은 위험하다"고 말하죠.

자신의 자녀들에게 백신을 맞혔다는 사실이 이를 말해줍니다.



AI의 희망, 교육과 의료의 미래, 그리고 기부와 은퇴

Steve:

AI의 위험성에 대해서는 충분히 이야기했습니다.

그렇다면 우리에게 희망을 줄 수 있는 무언가가 있을까요?

AI 덕분에 상황이 좋아질 수 있다는 점을 들려주실 수 있을까요?

Geoffrey:

AI 개발이 멈추지 않는 이유는,

그만큼 좋은 결과들도 무수히 많이 기대되기 때문입니다.

예를 들어, 의료 분야에서 AI는 정말 엄청난 발전을 이룰 것입니다.

앞으로는 당신의 가정의가

1억 명의 환자를 진료한 경험을 가진 AI일 수 있어요.

그 AI는 당신과 가족의 모든 진료 기록, 검사 결과를 기억하고,

훨씬 더 정확한 진단을 내릴 수 있죠.

이미 지금도 AI와 의사가 함께 진단할 경우,

단독으로 진단할 때보다 복잡한 질병에 대한 오류율이 훨씬 줄어들었습니다.

그리고 앞으로는 더 정교해질 겁니다.

교육 분야도 마찬가지예요.

개별 튜터가 붙은 아이는 학습 속도가 두 배나 빨라지죠.

왜냐하면 튜터는 아이가 무엇을 오해하고 있는지를 알아채기 때문입니다.

지금 AI는 아직 그 수준에 도달하진 못했지만,

10년 안에는 아주 정교한 튜터 역할이 가능해질 것입니다.

그 AI는 수백만 명의 학생을 이미 관찰해 왔기에

당신 아이가 어떤 오해를 하고 있는지 정확히 파악하고,

그걸 해결할 최적의 예시를 줄 수 있죠.

지금의 인간 튜터가 2배 더 낫다고 본다면,

AI 튜터는 3배, 4배 더 뛰어날 수도 있습니다.

Steve:

대학에겐 별로 좋은 소식은 아니겠군요?

Geoffrey:

그럴 수도 있죠.

대학이 더 이상 필요 없어질 수도 있습니다.

물론 대학원 수준의 연구 훈련은 여전히 필요해요.

연구는 결국 견습 과정이고,

“이 문제는 이렇게 접근해봐”처럼 감각을 익히는 일이니까요.

공식이 정해진 건 아니죠.

Steve:

지금 많은 젊은이들이 "코딩을 배우기 위해" 대학에 가고 있는데,
그 선택이 이제 잘못된 결정일 수도 있겠군요?

Geoffrey:

그럴 수도 있습니다.

물론 컴퓨터 과학은 단순히 코딩만 배우는 건 아니기 때문에,
그 안에는 여전히 배울 가치가 많은 부분이 있어요.

Steve:

사람들은 당신을 "AI의 대부(Godfather)"라고 부르는데,
그 별칭 마음에 드시나요?

Geoffrey:

꽤 맘에 듭니다.

사실 호의적인 의미로 붙은 건 아니었어요.

한 번 회의에서 제가 자꾸 끼어들자

Andrew Ng(앤드류 응)이 저를 **"Godfather"**라고 부르기 시작했거든요.

영국 원저에서 열린 회의였어요.

Steve:

그러니까 사람들이 말을 못하게 해서요?

Geoffrey:

제가 그 방에서 가장 나이가 많았고,

사람들을 밀어붙이고 있었죠.

Steve:

노벨상 상금의 절반이 대략 50만 달러쯤 되죠? (*번역자 주: 2024년 노벨 물리학상처럼 2명이
공동 수상자일 경우 일반적으로 수상자들이 노벨상 상금을 절반씩 나누어 받음)

그 중 35만 달러를 Water First에 기부하셨다고 들었어요. 맞나요?

Geoffrey:

맞습니다.

전체 상금이 약 100만 달러고,

그중 **25만 달러(미국 달러 기준)**를 기부했는데,

캐나다 달러로는 35만 달러쯤 됩니다.

Steve:

Water First는 어떤 기관인가요?

Geoffrey:

Water First는 캐나다 원주민 공동체 사람들에게

물 정화 기술을 교육하는 단체입니다.

그 지역 주민들이 스스로 안전한 식수를 확보할 수 있게 도와주는 거죠.

Steve:

왜 그 단체를 선택하셨나요?

Geoffrey:

예전에 페루에서 아이를 입양하면서,
그곳에서 두 달 정도 살았는데요,
그때 수도물을 마실 수 없었어요. 마시면 정말 위험했죠.
아기가 있는 집에서 안전한 물이 없다는 건,
매일 하루 종일 아이가 아프지 않게 하려 애쓰는 일입니다.
그건 정말 말도 안 되는 부담이에요.
그리고 저는 캐나다처럼 부유한 나라에 살면서,
여전히 수도물을 마실 수 없는 원주민 공동체가 존재한다는 건
정말 부끄럽고 부당한 일이라고 생각합니다.
온타리오주만 해도, 전체 원주민 공동체의 20%는 안전한 물을 못 마셔요.

Steve:

당신이 만족하실지는 모르겠지만,
그래도 10년 전보다는 상황이 나아졌어요.

Geoffrey:

그럴 수도 있겠죠.

Steve:

그래도 여전히 모든 공동체에 안전한 식수를 제공해야 합니다.

Geoffrey:

당연하죠. 모든 사람이 깨끗한 물을 누릴 자격이 있습니다.

Steve:

그럼 이제 앞으로의 계획은 어떻게 되시나요?

Geoffrey:

은퇴하려고 하는데요...

그게 잘 안 되고 있어요.

Steve:

실례가 안 된다면 나이가 어떻게 되시나요?

Geoffrey:

77살입니다.

Steve:

그 나이면 아직 은퇴하기 너무 젊으신데요.

Geoffrey:

저는 75세에 구글을 떠나면서 은퇴하려고 했어요.

Steve:

아직 앞날이 창창하신 것 같아요.

앞으로도 한두 챕터는 더 남아 있으실 것 같아요.

이상으로 Geoffrey Hinton 교수님 인터뷰를 마칩니다.

함께해 주셔서 감사합니다.