

Ấn bản Thư viện AI tiếng Việt

Những con người đã thay đổi khoa học bằng trí tuệ nhân tạo

Hành trình từ dữ liệu
đến giải Nobel

Luật sư Kim Kyung-jin

kimkj.com · Thư viện AI · 2026

Những con người đã thay đổi khoa học bằng trí tuệ nhân tạo

Hành trình từ dữ liệu đến giải Nobel

Luật sư Kim Kyung-jin

Cuốn sách đi theo mười một nhà khoa học đã đưa trí tuệ nhân tạo vào trung tâm của khám phá khoa học, từ gấp protein, phát hiện thuốc mới, phòng thí nghiệm tự hành, nhà khoa học máy, máy đọc bài báo, đến câu hỏi ai sở hữu một phát hiện do máy tạo ra.

Cuốn sách này được sắp xếp và biên soạn với sự hỗ trợ của trí tuệ nhân tạo. Nguồn tư liệu được thu thập bằng các ngôn ngữ chính trên thế giới. NotebookLM không giới hạn câu hỏi theo ngôn ngữ của nguồn. Độc giả có thể hỏi bằng tiếng Việt về cùng kho nghiên cứu đã được dùng để chuẩn bị cuốn sách này, dù một số tư liệu là tiếng Hàn, tiếng Anh hoặc ngôn ngữ khác.

Kho NotebookLM công khai được dùng khi chuẩn bị cuốn sách này:
<https://notebook.google.com/notebook/0e7ce8f8-46f1-4c13-bce2-a8b4f4d2619a>

Mục lục

Lời nói đầu

Dẫn nhập

Chương 1. Demis Hassabis: Từ trò chơi đến sự sống

Chương 2. John Jumper: Người giải bài toán 50 năm

Chương 3. Regina Barzilay: Con mắt tìm ra Halicin

Chương 4. David Baker: Tạo ra những protein chưa từng tồn tại

Chương 5. Alán Aspuru-Guzik: Phòng thí nghiệm tự hành

Chương 6. Hod Lipson: Đảo ngược kỹ thuật Newton

Chương 7. Max Tegmark: Sự ra đời của nhà vật lý AI

Chương 8. Ross King: Cha đẻ của thí nghiệm tự chủ

Chương 9. Gabe Gomes: Phòng thí nghiệm vận hành bằng ngôn ngữ

Chương 10. Sam Rodriques: Cỗ máy đọc một triệu bài báo

Chương 11. Hiroaki Kitano: Người thiết kế nhà khoa học máy

Chương 12. Ai là chủ nhân của phát hiện?

Kết luận. Vượt qua sự phong phú của dự đoán không có hiểu biết

Phụ lục

Lời nói đầu

Cuốn sách này kể về những con người đã dùng trí tuệ nhân tạo để thay đổi khoa học. Mười một nhà khoa học xuất hiện trong sách, từ một thần đồng cờ vua trở thành nhà điều hành đoạt giải Nobel, một nhà nghiên cứu chuyển từ vật lý sang giải quyết bài toán nan giải 50 năm của hóa học, cho đến một giáo sư giao toàn bộ phòng thí nghiệm cho robot. Họ có chung một điểm: họ là những người đã từng bước chuyển giao từng mất xích trong chu trình khoa học — từ quan sát, giả thuyết, dự đoán, thí nghiệm đến kiểm chứng — cho máy móc.

Tôi không phải là một nhà khoa học. Tôi từng làm công tố viên trong 13 năm, tham gia lập pháp trong lĩnh vực khoa học công nghệ tại Quốc hội, hiện đang điều hành một công ty trí tuệ nhân tạo và giảng dạy tại trường đại học. Dưới góc nhìn của một người làm việc với luật pháp và thể chế, khi dõi theo dòng chảy này, tôi đứng trước một câu hỏi không thể trì hoãn: Ai là chủ nhân của phát minh được tạo ra từ giả thuyết do máy móc lập nên? Tôi muốn đưa vào cùng một cuốn sách cả việc truyền tải những thành tựu mà các nhà khoa học đã gây dựng, lẫn việc chỉ ra những khoảng trống mà các thành tựu đó tạo ra trong luật pháp và thể chế. Chương 12 chính là nơi dành để bàn về những khoảng trống ấy.

Cuốn sách này được viết dành cho những độc giả không chuyên về khoa học. Các công thức và diễn giải toán học của thuật toán đã được lược bỏ khỏi phần nội dung chính và tập hợp lại trong các ghi chú kỹ thuật ở cuối sách. Nội dung chính được trình bày để người đọc có thể tiếp cận trọn vẹn chỉ qua câu chuyện của con người và những khái niệm được minh chứng qua các câu chuyện đó. Tôi hy vọng cuốn sách này cũng sẽ là tài liệu thảo luận hữu ích cho các sinh viên đang ở tuyến đầu nghiên cứu, cũng như những người thiết kế thể chế và phân bổ kinh phí nghiên cứu.

Nội dung trong sách được trình bày dựa trên các tài liệu công khai tính đến tháng 8 năm 2026. Lĩnh vực này có thể thay đổi diện mạo chỉ trong vài tháng. Xin quý độc giả đón nhận các chức danh và số liệu ghi trong sách như những ghi nhận tại thời điểm đó. Về tên riêng, sách tuân theo quy tắc phiên âm tiếng nước ngoài của Viện Ngôn ngữ Quốc gia Hàn Quốc nhưng vẫn tôn trọng các cách viết thông dụng đã định hình, đồng thời kèm theo tên gốc bằng tiếng nước ngoài ở lần xuất hiện đầu tiên. Đối với các thuật ngữ học thuật, sách ưu tiên giải nghĩa dễ hiểu và ghi kèm thuật ngữ gốc ở những nơi cần thiết.

Trong quá trình thực hiện cuốn sách này, tôi đã nhận được sự trợ giúp từ các công cụ trí tuệ nhân tạo. Tôi sử dụng chúng để tìm kiếm, thu thập tài liệu và trau chuốt bản thảo, đồng thời kiểm soát theo các nguyên tắc sau: Những sự thật nêu trong sách đều được đối chiếu và xác minh với nguồn gốc, những nội dung chưa được kiểm chứng đều bị loại bỏ, còn nhận định và lập luận trong sách là của tác giả chứ không phải của công cụ. Chủ đề của cuốn sách này — giao cho máy móc đến đâu và con người nắm giữ từ đâu — cũng là điều tôi cố gắng duy trì ngay trong chính phương pháp viết sách.

Niềm vui khám phá từ lâu đã thuộc về con người. Trong thời đại mà một phần niềm vui ấy trở thành phần việc của máy móc, tôi trân trọng gửi cuốn sách này như một lời mời cùng nhau suy ngẫm xem điều gì còn lại cho con người.

Tháng 8 năm 2026, Kim Kyung-jin

Dẫn nhập

Tốc độ không thể theo kịp

Khi Tycho Brahe qua đời vào năm 1601, thứ ông để lại cho Johannes Kepler là 20 năm dữ liệu quan sát thiên văn. Kepler đã bám lấy tập dữ liệu đó và tính toán lại hơn 40 lần trong suốt 4 năm. Kết quả cuối cùng là sự ra đời của định luật hành tinh chuyển động theo quỹ đạo elip. Dữ liệu tuy đồ sộ nhưng vẫn là một khối lượng mà một người có thể đọc hết nếu dành cả cuộc đời. Khoa học đã lớn lên trên những khối dữ liệu có quy mô mà một bộ não con người có thể gánh vác như thế.

Bây giờ tình hình đã khác. Máy giải mã gen và máy gia tốc hạt tuôn ra hàng terabyte dữ liệu mỗi ngày. Bên trong tế bào, hàng vạn protein, gen và chất chuyển hóa đan xen như một mạng lưới và tác động qua lại lẫn nhau. Lực làm rơi quả táo và lực giữ mặt trăng quay quanh quỹ đạo đã được gói gọn trong một phương trình duy nhất, nhưng mạng lưới này không thể chứa đựng trong vài phương trình. Tốc độ tích lũy của các bài báo khoa học cũng không đứng về phía con người. Các học giả gọi hiện tượng tri thức biến mất mỗi ngày vượt quá giới hạn vật lý mà một người có thể đọc là chân trời thông tin. Có bài báo viết rằng một loại gen thúc đẩy ung thư, trong khi bài báo khác lại viết chính gen đó tiêu diệt tế bào ung thư. Không ai có thời gian để mở hai bài báo đó ra cạnh nhau, nên những mâu thuẫn như vậy cứ bị bỏ mặc trong giới học thuật. Điều đó giống như việc câu trả lời đã được in sẵn và xếp trên giá sách, nhưng không ai có thể đồng thời mở cả hai trang sách đó.

Chất lượng của tri thức tích lũy cũng đang lung lay. Trong sinh học và y học lâm sàng, phần lớn các bài báo được công bố không thể tái lập cùng một kết quả ở các phòng thí nghiệm khác. Đó là hiện tượng mà giới khoa học gọi là khủng hoảng tái lập. Một trong những nguyên nhân là thực tế rằng các thí nghiệm phụ thuộc vào đầu ngón tay của con người. Thứ tự cho thuốc thử, vài giây tiếp xúc với nhiệt độ phòng, hay sự run rẩy nhẹ của bàn tay đều làm thay đổi kết quả, nhưng những cảm giác tinh tế đó lại không được ghi chép trong bài báo. Chúng ta rơi vào tình cảnh tri thức tràn ngập nhưng lại phải đặt câu hỏi lại từ đầu rằng liệu có thể tin tưởng tri thức đó hay không.

Bộ não của người làm nghiên cứu cũng đã bộc lộ giới hạn. Con người chỉ có thể nắm bắt đồng thời khoảng ba đến bốn biến số thí nghiệm cùng một lúc. Kết quả là những kết luận nhỏ tích lũy riêng rẽ ở từng phòng thí nghiệm, rồi nằm rải rác mà không thể kết nối với nhau. Sebastian Musslick tại Đại học Brown gọi tình trạng này là những hòn đảo phân mảnh. Thực tế rằng công cụ để nghiên cứu bộ não lại chính là bộ não đó cũng đồng nghĩa với việc con người không có cách nào tự mình thoát khỏi giới hạn này.

Ý tưởng giao phó việc khám phá cho máy móc vốn đã có từ lâu. Dự án DENDRAL của Đại học Stanford vào giữa những năm 1960 là cột mốc đầu tiên. Joshua Lederberg, người đoạt giải Nobel Sinh lý học hoặc Y học, và Edward Feigenbaum, người đoạt giải Turing, đã bắt tay hợp tác với bối cảnh là kế hoạch thám hiểm sao Hỏa đang được xúc tiến vào thời điểm đó. Tín hiệu vô tuyến giữa Trái đất và sao Hỏa mất tới 10 phút chỉ cho một chiều, nên tàu thăm dò đáp xuống sao Hỏa phải có khả năng tự đọc thành phần của đất mà không cần chờ chỉ thị của con người. DENDRAL lần đầu tiên chứng minh rằng máy móc có thể hấp thụ kiến thức hóa học để phỏng đoán cấu trúc phân tử chưa biết. Sau đó, máy móc đã từng bước dòm ngó ngưỡng cửa khám phá, nhưng trong suốt nửa thế kỷ, nhân vật chính của phòng thí nghiệm vẫn luôn là con người.

Đứng trước bức tường này, câu trả lời mà các nhà khoa học một lần nữa đưa ra chính là trí tuệ nhân tạo. Tuy nhiên, nếu hiểu câu trả lời này chỉ là phần mở rộng của một chiếc máy tính thì chúng ta sẽ bỏ lỡ quy mô thực sự của những gì đang diễn ra. Máy tính đã thay thế con người làm các phép tính khoa học suốt nửa thế kỷ qua. Sự thay đổi đang diễn ra hiện nay không nằm trên đường nối dài đó. Bản chất của việc trí tuệ nhân tạo thay đổi khoa học không nằm ở chỗ tính toán nhanh hơn. Nó nằm ở chỗ vòng tuần

hoàn khép kín từ quan sát đến giả thuyết, dự đoán, thí nghiệm và kiểm chứng lần đầu tiên bắt đầu được khép kín bên trong máy móc. Câu văn này chính là mệnh đề dẫn dắt toàn bộ cuốn sách.

Vòng tuần hoàn khép kín bên trong máy móc có nghĩa là máy móc đang từng bước tiếp quản những công việc vốn trước đây chỉ có con người làm. Việc quan sát dữ liệu để phỏng đoán quy luật, việc xác nhận phỏng đoán bằng thí nghiệm, và việc xem kết quả thí nghiệm để chọn câu hỏi tiếp theo đang lần lượt trở thành phần việc của máy móc. Gốc rễ của tư tưởng này đã có từ lâu. Herbert Simon, người từng nhận cả giải Nobel Kinh tế lẫn giải Turing, cho rằng khám phá không phải là nguồn cảm hứng huyền bí mà là một cuộc tìm kiếm thu hẹp không gian giả thuyết theo các quy tắc đã định sẵn, và ông đã gieo niềm tin đó vào các chương trình máy tính. Nửa thế kỷ trôi qua, giờ đây trên bản phác thảo ấy là những cỗ máy thực sự đang vận hành các phòng thí nghiệm.

Điều cuốn sách này muốn đề cập không phải là bản thiết kế của những cỗ máy đó, mà là những lựa chọn của những con người đã tạo ra chúng. Có người chỉ giao cho máy móc việc dự đoán và giữ lại phần diễn giải cho con người đến cùng. Có người thậm chí đã giao cả chìa khóa phòng thí nghiệm cho máy móc. Dĩ nhiên lý do vì sao lại có những lựa chọn rẽ nhánh trước cùng một công nghệ, chúng ta sẽ thấy rõ sự thay đổi này đã đi đến đâu và đang dừng lại ở chỗ nào. Đây chính là lý do các chương được chia theo tên người thay vì tên công nghệ.

Vậy thì câu hỏi vẫn còn đó. Quy luật mà máy móc tìm ra từ dữ liệu có thể được gọi là khám phá hay không? Nếu con người không thể giải thích được câu trả lời mà máy móc đưa ra thì tri thức đó thuộc về ai? Ai là người chịu trách nhiệm cho nghiên cứu dựa trên giả thuyết do máy móc

thiết lập? Cuốn sách này sẽ theo chân 11 nhà khoa học để bồi đắp cho câu hỏi này, và đưa ra câu trả lời của tác giả ở hai chương cuối. Trước đó, có một điều cần phải làm rõ trước. Đó chính là bản thân từ 'khám phá'.

Thế nào được gọi là khám phá?

Cụm từ máy móc tự mình làm khoa học thường được sử dụng, nhưng bên trong cụm từ đó lại pha trộn năm công việc có tính chất hoàn toàn khác nhau. Nếu không phân biệt năm điều này, những thành quả nhỏ sẽ bị thổi phồng và những thành quả lớn sẽ bị chôn vùi. Năm phần của cuốn sách này sẽ bám sát sự phân chia đó.

Thứ nhất là tái khám phá. Đó là việc một cỗ máy không hề biết trước phương trình đã biết lại có thể tìm lại chính phương trình đó chỉ bằng cách nhìn vào dữ liệu. Chương trình BACON của Herbert Simon đã truy lại các định luật của Kepler từ dữ liệu hành tinh, và sau này một chương trình có tên EUREKA cũng đã đi trên cùng con đường đó với nguồn dữ liệu thô ráp hơn. Tái khám phá không bổ sung thêm tri thức mới vì tờ đáp án đã có sẵn. Đây cũng là lý do giới học thuật chỉ ra rằng khi BACON tìm lại định luật Kepler, đó không phải là khám phá mà chỉ là một phép tính hồi quy.

Mặc dù vậy, lý do tái khám phá quan trọng là vì nó có thể chấm điểm được. Nếu đối chiếu công thức mà máy móc đưa ra với sách giáo khoa, chúng ta có thể biết ngay phương pháp này có hiệu quả hay không. Tái khám phá không phải là việc làm đầy kho tàng tri thức, mà là việc lấp đầy bàn thử nghiệm phương pháp. Chỉ những phương pháp vượt qua bàn thử nghiệm mới có đủ tư cách tiến tới những bài toán chưa có đáp án.

Thứ hai là dự đoán. Đó là việc đoán đúng cấu trúc hoặc tính chất mà chưa ai biết tới. AlphaFold đoán đúng cấu trúc 3 chiều của protein, hay Chemprop chỉ ra đặc tính kháng khuẩn của các hợp chất đều thuộc về nhóm này. Khác với tái khám phá, dự đoán không có sẵn tờ đáp án. Đúng hay sai chỉ có thể được phân định khi thực sự tạo ra chúng trong phòng thí nghiệm. Vì vậy, cỗ máy dự đoán luôn phải kết hợp với con người làm thí nghiệm để hoàn thiện tri thức.

Dự đoán có một giới hạn rõ ràng. Máy móc đưa ra câu trả lời đúng nhưng thường không thể nói rõ tại sao lại là câu trả lời đó. Đoán đúng và giải thích là hai việc hoàn toàn khác nhau. Nhìn nhận khoảng cách này như thế nào là vấn đề tranh luận xuyên suốt cho đến phần kết luận của cuốn sách này.

Thứ ba là thiết kế. Đó là việc tạo ra những thứ chưa từng tồn tại trong tự nhiên. RFdiffusion phác họa ra những protein mà quá trình tiến hóa chưa từng tạo ra suốt hàng trăm triệu năm là một ví dụ. Thiết kế có hướng đi ngược lại với dự đoán. Nếu dự đoán là việc đoán biết bản chất của thứ mà tự nhiên đã tạo ra, thì thiết kế là việc xác định trước tính chất mong muốn rồi từ đó phác họa ngược lại vật chất mang tính chất ấy.

Bảng điểm của thiết kế không phải là đúng hay sai, mà là có hoạt động hay không. Nếu phân tử được thiết kế hoạt động đúng như mong muốn trong ống nghiệm thì đó là thành công, còn nếu sụp đổ thì là thất bại. Vì vậy, thiết kế vừa mang gương mặt của khoa học, vừa mang gương mặt của kỹ thuật. Khoảnh khắc thiết kế trở nên khả thi, câu hỏi của khoa học cũng thay đổi. Bởi vì bên cạnh câu hỏi 'Tại sao tự nhiên lại có hình thù như thế này?', một câu hỏi mới được đặt ra: 'Trong số những thứ tự nhiên chưa từng tạo ra, chúng ta sẽ tạo ra cái gì?'.

Thứ tư là thực thi. Đó là việc lập quy trình thí nghiệm để kiểm chứng giả thuyết và vận hành trang thiết bị thực tế để tiến hành thí nghiệm đó. Nhà khoa học robot Adam đã tự mình tiến hành thí nghiệm nuôi cấy nấm men, và Coscientist (Gomes) đã điều khiển thiết bị thí nghiệm bằng mô hình ngôn ngữ. Đến bước thực thi, trí tuệ nhân tạo lần đầu tiên bước ra khỏi màn hình. Đây là giai đoạn mà trí thức có được đôi tay.

Lý do thực thi quan trọng là vì nó chạm tới điểm yếu lâu đời của khoa học: tính tái lập. Máy móc lưu lại các điều kiện thí nghiệm dưới dạng mã lệnh đến từng microliter. Khi quy trình vốn phụ thuộc vào cảm giác đầu ngón tay của con người trở thành mã lệnh, robot ở phòng thí nghiệm khác có thể lặp lại y nguyên thí nghiệm đó.

Thứ năm là tự chủ. Đó là việc tự mình quan sát kết quả thí nghiệm, sửa đổi giả thuyết và tự chọn thí nghiệm tiếp theo. Đó là trạng thái mà vòng tuần hoàn từ quan sát đến kiểm chứng quay tròn một vòng mà không cần con người, trạng thái mà cuốn sách này gọi là vòng lặp kín. Nếu bốn điều trước là việc bàn giao từng mắt xích của vòng tuần hoàn cho máy móc, thì tự chủ là việc nối các mắt xích đó lại để khép kín vòng lặp.

Tự chủ mang tính chất trí thức khác biệt so với bốn giai đoạn trước. Từ tái khám phá đến thực thi là cấu trúc mà máy móc giải bài toán do con người đưa ra. Nhưng khi đạt đến tự chủ, việc lựa chọn bài toán tiếp theo cần giải cũng trở thành phần việc của máy móc. Đây là giai đoạn máy móc lần đầu tiên chạm tay vào công việc mang tính con người nhất trong khoa học: tò mò về điều gì. Phần lớn các cuộc tranh luận được đề cập ở phần sau của cuốn sách này đều nảy sinh từ đây.

Năm giai đoạn này nối liền nhau như một chiếc thang. Càng bước lên nấc thang cao hơn, bàn tay con người càng giảm đi, và bản chất công việc còn lại cho con người cũng thay đổi. Cho đến nay, năm điều này thường bị gộp chung dưới một câu nói: máy móc tự mình làm khoa học. Thế nhưng, nói thành quả của tái khám phá như thể thành quả của dự đoán sẽ là phóng đại, và nói thành quả của thực thi như thể thành quả của tự chủ sẽ giống như nói rằng thời đại chưa tới đã đến rồi. Dù bạn bắt gặp thành tựu của hệ thống nào trong cuốn sách này, xin hãy tự hỏi trước rằng nó đang đứng ở đâu trong năm nấc thang ấy.

Mục lục của cuốn sách này bám sát chiếc thang đó. Các nhân vật ở Phần 1 là những người hướng tới dự đoán, Phần 2 là thiết kế, Phần 3 là tìm kiếm quy luật – phiên bản hiện đại của tái khám phá, Phần 4 là thực thi, và Phần 5 là tự chủ. Thứ tự các chương càng tiến tới thì vị trí của máy móc càng mở rộng. Ở cuối mỗi phần có một ghi chú của tác giả mang cùng tiêu đề 'Vai trò của con người đã chuyển dịch về đâu?',

nhằm tổng kết những gì đã rời khỏi bàn tay con người và những gì còn lại cho con người trong phần đó. Khi câu chuyện về 11 nhân vật khép lại, Chương 12 sẽ xem xét câu hỏi về pháp luật và thể chế: 'Ai là chủ nhân của khám phá?', và phần kết luận sẽ tìm kiếm con đường vượt qua sự trù phú của dự đoán mà không có sự thấu hiểu. Điều gì còn lại cho con người ở cuối chiếc thang ấy chính là câu hỏi cuối cùng của cuốn sách này.

Hai hệ thống trùng tên Trong cuốn sách này xuất hiện hai hệ thống có tên gọi giống nhau khi chuyển ngữ sang tiếng Hàn. Một là Coscientist do nhóm nghiên cứu của Gabriel Gomes tại Đại học Carnegie Mellon công bố trên tạp chí học thuật Nature vào năm 2023. Đây là hệ thống nơi mô hình ngôn ngữ lớn đọc các bài báo khoa học, viết mã thí nghiệm và vận hành thiết bị thí nghiệm hóa học thực tế. Hệ thống còn lại là Co-Scientist do Google DeepMind công bố trên Nature vào năm 2026. Đây là hệ thống nơi nhiều tác nhân trí tuệ nhân tạo lập giả thuyết, phản bác lẫn nhau và trau chuốt ý tưởng nghiên cứu, nhưng không vận hành thiết bị thí nghiệm. Nơi chế tạo, thời điểm lần công việc thực hiện đều khác nhau. Cuốn sách này phân biệt và ghi hệ thống trước là Coscientist (Gomes), hệ thống sau là Co-Scientist (DeepMind). Coscientist (Gomes) được đề cập ở Chương 9, còn Co-Scientist (DeepMind) được bàn đến ở phần kết luận.

Nội dung tường thuật trong chương này dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Phần 1

Cỗ máy dự đoán

Chương 1. Demis Hassabis: Từ trò chơi đến sự sống

Cậu bé bên bàn cờ vua

Vào mùa đông năm 1988, trong hội trường của một nhà thờ ở Liechtenstein, có một cậu bé đang ngồi. Đó là Demis Hassabis, mười hai tuổi. Trong hội trường, hàng trăm kỳ thủ tề tựu từ khắp nơi trên thế giới, mỗi người đều đang đối diện với bàn cờ của mình. Ngồi đối diện với Hassabis là một kỳ thủ cờ vua lão luyện ở độ tuổi giữa ba mươi, từng là cựu vô địch Đan Mạch. Ván cờ bắt đầu từ buổi sáng và kéo dài hơn mười tiếng đồng hồ mà vẫn chưa kết thúc. Vào thời đó, các giải đấu cờ vua chưa có luật giới hạn thời gian suy nghĩ. Cho đến khi mặt trời lặn ngoài cửa sổ và màu sắc của kính màu trong hội trường thay đổi, cả hai người vẫn không thể rời khỏi chỗ ngồi.

Khi bước vào tàn cuộc, thế cờ trên bàn vẫn còn một con đường rõ ràng để dẫn đến kết quả hòa. Bằng cách thí Hậu để khiến Vua của đối phương rơi vào thế bất động, cậu có thể chia điểm. Tuy nhiên, sau mười tiếng đồng hồ liên tục tính toán các nước đi, đầu óc của cậu bé đã kiệt sức. Các quân cờ trước mắt mờ đi và chồng chéo lên nhau. Tin rằng mình sắp rơi vào thế chiếu hết không thể tránh khỏi, Hassabis đã đứng dậy và tuyên bố xin hàng.

Đối thủ bật dậy, cười và bắt đầu sắp xếp lại các quân cờ. Ông chỉ ra từng nước cờ hòa mà Hassabis đã bỏ lỡ. Cậu bé đỏ mặt ngay tại chỗ. Nếu kiên trì thêm một chút nữa, cậu đã có thể chia điểm, nhưng chỉ vì đầu óc kiệt sức mà đành để mất trọn vẹn điểm số. Rời khỏi hội trường, cậu một mình bước đi trên con đường núi lạnh giá của Liechtenstein. Đón nhận cơn gió thổi từ chân dãy Alps, cậu đã ấp ủ một câu hỏi sẽ thay đổi cả cuộc đời mình.

Thời thơ ấu của Hassabis không bắt đầu từ cờ vua. Ông sinh ngày 27 tháng 7 năm 1976 tại London. Cha ông xuất thân từ một gia đình nhập cư người Síp gốc Hy Lạp, là một người ghét sự khuôn mẫu, từng mở tiệm bán đồ chơi rồi sáng tác và hát nhạc. Mẹ ông là người Singapore gốc Hoa, mồ côi cha mẹ từ nhỏ và phải bươn chải kiếm sống qua các công việc y tá và nhân viên bán hàng. Cả hai đều không ép buộc con trai phải theo một con đường sự nghiệp nhất định. Gia đình Hassabis liên tục chuyển nhà mỗi khi người cha đổi việc. Chỉ riêng thời thơ ấu, ông đã chuyển trường hơn mười lần. Tại mỗi ngôi trường mới, Hassabis nhanh chóng rèn luyện được trực giác nắm bắt các mối quan hệ giữa bạn bè. Mỗi khi mở cánh cửa lớp học xa lạ, việc nhận biết ai thân với ai và quy tắc nào mới là quy tắc thực sự đối với ông gần như là một trò chơi.

Cờ vua đến với ông năm lên bốn tuổi. Khi nhìn cha và chú chơi cờ trên bàn, cậu bé Hassabis đã vò vĩnh đòi dạy luật chơi. Cha mẹ chỉ nhẹ nhàng chỉ cho cậu cách di chuyển các quân cờ. Thế nhưng chỉ hai tuần sau khi học, cậu bé đã lần lượt đánh bại cả cha lẫn chú. Cha mẹ đưa cậu đến một câu lạc bộ cờ vua địa phương, và Hassabis thể hiện kỹ năng không hề lép vế ngay cả trước các kỳ thủ trưởng thành. Lên sáu tuổi, cậu giành chức vô địch lứa tuổi dưới 8 tại London, và sau đó đảm nhận vai trò đội trưởng đội tuyển trẻ quốc gia Anh trong suốt thời niên thiếu. Trong trận đấu giao hữu giữa các kỳ thủ xuất sắc của Anh và Mỹ diễn ra tại London năm 1986, cậu đảm nhận bàn 1, vị trí dẫn đầu của đội. Đó là vị trí mà một cậu bé mười tuổi đại diện cho đất nước chơi ván cờ đầu tiên.

Tháng 1 năm 1990, khi tròn mười ba tuổi, Hassabis đạt mức hệ số Elo 2300 của Liên đoàn Cờ vua Quốc tế và giành danh hiệu Kiện tướng Ứng viên. Đây là thành tích xếp thứ hai thế giới trong số các kỳ thủ dưới mười bốn tuổi, chỉ sau nữ kỳ thủ xuất sắc nhất thời bấy giờ là Judit Polgár. Tuy nhiên, vì hoàn cảnh gia đình khó khăn nên cha mẹ vẫn phải chắt chiu chi phí tham dự các giải đấu quốc tế, mỗi sai lầm trong từng ván đấu đều trở thành nỗi lo lắng cho kinh tế gia đình. Cờ vua đối với ông vừa là một trò chơi thú vị, vừa là một gánh nặng khó lòng trút bỏ.

Thế nhưng, khi bước đi trên con đường núi ở Liechtenstein, câu hỏi mà Hassabis đối mặt lại vượt ra ngoài thắng bại. So với các quy luật vật lý chi phối vũ trụ, bàn cờ 64 ô lưới chỉ là một thế giới khép kín. Dẫu vậy, những bộ óc hàng đầu thế giới lại dành cả cuộc đời chỉ để nghiên cứu cách đánh bại lẫn nhau. Liệu có thể dùng thời gian đó để tìm ra phương pháp chữa bệnh ung thư hay giải các phương trình giải thích lực hấp dẫn hay không, cậu bé đã tự hỏi như thế. Cậu cũng nhận ra rằng việc ghi nhớ các khai cuộc và tàn cuộc trong cờ vua là thứ tri thức hạn hẹp không thể chuyển giao sang các lĩnh vực khác. Câu hỏi của cậu bé mười hai tuổi không dừng lại ở đó. Nó còn tiến xa hơn đến việc: nếu dùng thời gian trau dồi cách chiến thắng trong cờ vua để làm sáng tỏ nguyên lý vận hành của chính tư duy con người thì sao? Vốn có thể sống như một kỳ thủ cờ vua chuyên nghiệp, Hassabis lại chọn con đường đào sâu vào chính phương thức tư duy. Đó là quyết tâm bóc tách nguyên lý tư duy nâng đỡ cờ vua và cấy ghép nó vào một thứ gì đó không phải con người. Gác lại bảng thành tích lẫy lừng của nhà vô địch trẻ, ông một lần nữa đứng trước bài toán lớn hơn và mong lung hơn nhiều: trí thông minh.

Quyết tâm này vẫn tiếp diễn cho đến tận ngày nay, gần 40 năm sau. Ngày 5 tháng 8 năm 2026, Hassabis từ chức Tổng giám đốc điều hành Google DeepMind để trở thành Chủ tịch Google DeepMind, đồng thời đảm nhận chức danh Nhà khoa học trưởng của Alphabet và tiếp tục giữ vai trò đại diện Isomorphic Labs. Việc quản lý điều hành hằng ngày do Phó Chủ tịch cấp cao Koray Kavukcuoglu phụ trách. Trong bức thư gửi nhân viên, Hassabis tuyên bố rằng trí tuệ nhân tạo

tổng quát đã ở ngay trước mắt, và bản thân ông quyết định dành thời gian cho sự an toàn cũng như chiến lược dài hạn của trí tuệ nhân tạo tổng quát. Trong một cuộc phỏng vấn, ông nhận định về cục diện cạnh tranh hiện nay khi nhiều quốc gia và doanh nghiệp cùng theo đuổi trí tuệ nhân tạo tổng quát rằng vốn dĩ mọi chuyện không nên diễn ra theo cách này. Vào tháng 7 cùng năm, ông cũng đề xuất rằng các quốc gia trên thế giới nên chung tay thành lập một cơ quan giám sát trước cuối năm nhằm đánh giá các mô hình trí tuệ nhân tạo nguy hiểm và có thể làm chậm tốc độ phát triển nếu cần thiết. Cuộc cải tổ này diễn ra giữa lúc việc ra mắt một số mô hình bên trong DeepMind bị trì hoãn và nhân tài rời đi. Cùng ngày hôm đó, Jeff Dean, người đặt nền móng cho trí tuệ nhân tạo của Google, đã rời Google sau 27 năm gắn bó để thành lập một công ty có tên là Discovery Loop. Vì đây là một công ty hướng tới việc tự động hóa nghiên cứu để tự cải thiện mà hầu như không cần sự can thiệp của con người, giấc mơ về vòng lặp khép kín mà cuốn sách này theo đuổi giờ đây đã xuất hiện dưới dạng tên của một công ty. Ngay cả trong khoảnh khắc tổ chức biến động, thay vì việc điều hành thường nhật của DeepMind, Hassabis dường như đã quay trở lại với câu hỏi mà cậu bé mười hai tuổi từng ấp ủ trên con đường núi: bộ óc kiệt xuất nên được sử dụng vào việc gì.

Đường vòng mang tên trò chơi

Năm 2003, trong một văn phòng ở London, khuôn mặt của những nhà phát triển ngồi trước màn hình máy tính đều đanh lại. Trên màn hình mở ra quốc gia hư cấu Novistrana ở Đông Âu, nơi một triệu cư dân trí tuệ nhân tạo phải vận động với niềm tin, lịch trình hằng ngày và cảm xúc riêng của từng người. Đó là một dự án đầy tham vọng với sự tham gia của các nhà phát hành lớn như Microsoft và Vivendi Universal. Thế nhưng vào thời điểm đó, một CPU Pentium dùng cho gia đình không thể gánh nổi khối lượng tính toán này. Rốt cuộc, số lượng cư dân từ một triệu người đã phải cắt giảm xuống quy mô vài trăm người. Tên của trò chơi này là Republic: The Revolution. Và người dẫn dắt dự án đó chính là Demis Hassabis.

Năm 1997, Hassabis hoàn thành chương trình cử nhân khoa học máy tính tại Queens' College thuộc Đại học Cambridge với kết quả 'Double First', tức hạng tối ưu xuất sắc nhất. Phía trường đại học mong muốn ông học thẳng lên chương trình tiến sĩ trí tuệ nhân tạo. Nhưng Hassabis đã chọn một con đường khác. Năm mười bảy tuổi, ông đã từng làm việc tại Bullfrog Productions với tư cách là lập trình viên trưởng kiêm đồng thiết kế của trò chơi Theme Park (1994). Theme Park không phải là một trò chơi

arcade tầm thường. Giữa thập niên 1990, đó đã là trò chơi mà hàng nghìn du khách ảo phản ứng với những cảm xúc khác nhau tùy thuộc vào giá vé trò chơi và vị mặn của đồ ăn vặt. Khái niệm mà ngày nay gọi là mô phỏng đa tác tử đã được ông hiện thực hóa trước tiên qua mã nguồn trò chơi vào thời kỳ đó.

Sau khi tốt nghiệp Cambridge, Hassabis gia nhập Lionhead Studios do Peter Molyneux thành lập với vai trò lập trình viên trí tuệ nhân tạo trưởng. Tại đây, ông đã tạo ra trò chơi Black & White (2001), trong đó người chơi đóng vai một vị thần nuôi dưỡng sinh vật ảo. Nếu người chơi vuốt ve lưng sinh vật thì đó là khen thưởng, còn nếu dùng tay đánh xuống thì là trừng phạt. Ông đã từ chối cách tiếp cận mà lập trình viên phải nhồi từng quy tắc một vào mã nguồn. Thay vào đó, ông thiết kế để sinh vật tự học hỏi từ phần thưởng và hình phạt theo hành vi của người chơi. Đó là dạng sơ khai của phương pháp mà ngày nay gọi là học tăng cường, nơi máy móc tự tìm ra câu trả lời thông qua thử và sai. Kết quả học tập của mỗi sinh vật là khác nhau, nên ngay cả những người mua cùng một trò chơi cũng sẽ nuôi dưỡng những sinh vật có tính cách hoàn toàn khác biệt.

Năm 1998, ở tuổi hai mươi hai, Hassabis thành lập Elixir Studios tại London và bắt tay vào phát triển Republic. Đội ngũ phát triển e ngại báo cáo các giới hạn kỹ thuật, và chỉ luôn đáp lại rằng không có vấn đề gì. Cuối cùng, vào tháng 4 năm 2005, Elixir Studios đã phải đóng cửa. Hassabis từ

thất bại này đã rút ra một bài học. Rằng tưởng tượng đi trước thế giới năm mươi năm chỉ là mộng tưởng, và chỉ có thiết kế đi trước đúng năm năm mới thực sự hoạt động được.

Sau khi Elixir sụp đổ, Hassabis không tiếp tục ở lại ngành công nghiệp trò chơi mà bước vào khám phá bộ não. Ông đăng ký học chương trình tiến sĩ khoa học thần kinh nhận thức tại University College London. Ông là người đã tự mình trải nghiệm trong thực tế phát triển trò chơi rằng trí tuệ nhân tạo được xây dựng từ các quy tắc định sẵn và cấu trúc cây theo tình huống rất dễ sụp đổ trước những tình huống bất ngờ. Ông nhận định rằng không thể tạo ra trí thông minh thực sự bằng cách để con người viết sẵn toàn bộ quy tắc. Vì vậy, ông quyết định tìm hiểu phương thức học tập của bộ não sinh học mà quá trình tiến hóa đã hoàn thiện qua hàng trăm triệu năm.

Giáo sư hướng dẫn của ông là Eleanor Maguire, một chuyên gia đầu ngành về khoa học thần kinh nhận thức. Thí nghiệm mà hai người thực hiện với các bệnh nhân mất trí nhớ đã chứng minh rằng hồi hải mã vừa là nơi ghi nhớ quá khứ vừa là nơi tưởng tượng ra tương lai. Quá trình phát hiện này được tái sinh thành thuật toán học tập của DeepMind sau này sẽ được tiếp tục ở phần tiếp theo.

Cậu thiếu niên mười bảy tuổi từng làm game, người nghiên cứu sinh từng chăm chú quan sát hồi hải mã, và nhà khoa học đang nói về trí tuệ nhân tạo tổng quát hiện nay rốt cuộc đều là một người. Hassabis đã khép lại công ty thất bại và chuyển hẳn sang một ngành học hoàn toàn khác. Hơn nữa, ông làm điều đó bằng cách quay lại làm sinh viên khi đã cận kề tuổi ba mươi. Nếu Republic thành công, có lẽ Hassabis giờ này vẫn là giám đốc của một công ty trò chơi ở London. Nhưng chính vì thất bại nên ông mới nghiên cứu bộ não, và nhờ nghiên cứu bộ não nên mới có DeepMind ngày hôm nay.

Bài học từ hồi hải mã

Năm 2007, tại Viện Thần kinh học thuộc UCL ở London, một người đàn ông đang ngồi trước một bệnh nhân nhắm nghiền mắt. Ông đề nghị bệnh nhân tưởng tượng về một bãi biển nhiệt đới yên tĩnh với bờ cát trắng trải dài dưới ánh nắng chói chang. Ông nhờ bệnh nhân nghĩ rằng mình đang nằm nghỉ trên một chiếc võng và kể lại từng thứ mình nhìn thấy. Bệnh nhân im lặng một lát rồi mở lời. Những từ ngữ như cát, biển xanh, hải âu xuất hiện, nhưng không có một khung cảnh nào được tạo dựng. Các từ ngữ trôi nổi rời rạc giữa không trung, và bên trong người bệnh nhân đó không có ai có thể gắn kết bầu trời, biển cả và bãi cát thành một bức tranh phong cảnh duy nhất. Người đàn ông đó, Demis Hassabis, đã ghi lại khoảnh khắc này vào biên bản thí nghiệm. Điều đó có nghĩa là người mất đi ký ức cũng không thể vẽ nên tương lai.

Như đã thấy ở phần trước, Hassabis vừa thanh lý công ty và bước vào chương trình tiến sĩ tại Viện Thần kinh học Queen Square thuộc UCL. Việc từ bỏ vị trí giám đốc công ty để trở lại làm một nghiên cứu sinh ngồi bệt trên sàn phòng thí nghiệm xử lý dữ liệu là một lựa chọn không thể giải thích bằng con đường phát triển sự nghiệp thông thường. Câu hỏi duy nhất mà ông ấp ủ là: bộ não con người làm thế nào để hình dung trong đầu một tương lai mà mình chưa từng trải qua? Ông gia nhập phòng thí nghiệm của Giáo sư Maguire và bắt đầu đi tìm câu trả lời.

Lý thuyết do Hassabis xây dựng có tên là cấu tạo cảnh quan. Cốt lõi của nó không hề phức tạp. Việc ghi nhớ quá khứ và tưởng tượng tương lai không nằm ở hai ngăn kéo khác nhau trong não bộ. Cả hai đều là một nhiệm vụ duy nhất do hồi hải mã đảm trách. Hồi hải mã không chỉ đơn thuần lấy các mảnh ký ức cũ ra khỏi kho để phát lại. Giống như một biên tập viên phim chọn lọc các đoạn phim rồi rọc để ghép thành một cảnh quay mới, hồi hải mã cũng kéo các mảnh màu sắc, hình khối và cảm xúc trong quá khứ từ vỏ não để ghép nối lại. Quá trình đó diễn ra theo khoảng bốn bước. Đầu tiên, nó trích xuất các mảnh ký ức phân tán. Tiếp theo, các tế bào lưới và tế bào vị trí xác định tọa độ không gian và thời gian nơi các mảnh ghép đó được đặt vào. Sau đó, hồi răng và vùng CA3 bên trong hồi hải mã loại bỏ sự sai lệch giữa các mảnh ghép và giữ chúng lại thành một cấu trúc thống nhất. Cuối cùng, trên khung cảnh đã hoàn thiện, bộ não bước thử vào những điều sắp xảy ra.

Trí tuệ nhân tạo dựa trên thống kê vốn là dòng chính thống thời bấy giờ còn cách rất xa năng lực này. Nó chỉ là một bộ nhận dạng mẫu tĩnh ghép nối đầu vào và kết quả dựa trên lượng dữ liệu khổng lồ, chứ không phải một cỗ máy tự xây dựng thế giới riêng trong tâm trí và trải qua quá trình thử và sai trong đó. Nguyên lý mà Hassabis tìm thấy từ hồi hải mã lại khác. Đó là nguyên lý mà tác tử tự mình xây dựng mô hình thế giới vượt qua giới hạn của dữ liệu quan sát được, tự mình trải qua thất bại trong thế giới ảo đó trước khi bước vào hành động thực tế. Nhận thức

sâu sắc này sau này đã trở thành bản phác thảo để DeepMind thiết kế phương pháp học từ thử và sai, tự mình phán đoán chỉ bằng dữ liệu và phần thưởng mà không cần bị áp đặt các quy tắc có sẵn.

Thí nghiệm xác thực lý thuyết này bằng dữ liệu đã được công bố trên tạp chí PNAS vào năm 2007. Tham gia thí nghiệm có 5 bệnh nhân bị tổn thương nặng cả hai bên hồi hải mã dẫn đến mất trí nhớ và 10 người khỏe mạnh; họ được yêu cầu tưởng tượng và mô tả bằng lời những tình huống mình chưa từng trải qua. Nhóm nghiên cứu đã chấm điểm các mô tả đó theo thang điểm 30 gọi là điểm nhất quán không gian. Đây là điểm số đánh giá khoảng cách giữa các vật thể, luật xa gần và cấu trúc của khung cảnh. Nhóm đối chứng khỏe mạnh đạt điểm trung bình 24,8. Nhóm bệnh nhân bị tổn thương hồi hải mã chỉ đạt điểm trung bình 9,3. Sự khác biệt mang tính thống kê rất rõ ràng: giá trị p nhỏ hơn 0,001. Nghiên cứu này đã được tạp chí Science bình chọn là một trong 10 thành tựu khoa học tiêu biểu của năm đó.

Năm 2010, Hassabis thành lập DeepMind và cấy ghép lý thuyết này vào máy móc. Năm 2013, DQN do DeepMind ra mắt đã tự học 57 trò chơi Atari chỉ bằng cách quan sát các điểm ảnh trên màn hình. Thiết bị giúp ổn định quá trình học tập này là bộ đệm phát lại trải nghiệm. Mô phỏng quá trình hồi hải mã phát lại ngẫu nhiên các ký ức ban ngày trong giấc ngủ của con người để khắc sâu vào vỏ não, máy móc cũng không học các trải nghiệm theo trình tự thời gian mà tích lũy chúng vào kho lưu trữ rồi trích xuất ngẫu nhiên để học. Nếu bị ràng buộc bởi trình tự thời gian, mạng nơ-ron sẽ bị cuốn theo những gì vừa mới học mà xóa sạch các quy tắc đã học trước đó. Bộ đệm phát lại đã giải quyết bài toán này.

Đến năm 2016, với AlphaGo, thiết kế này đã tiến thêm một bước. Mạng nơ-ron chính sách học các kỳ phổ của con người để thu hẹp các nước đi tiếp theo, còn mạng nơ-ron giá trị đọc thế cờ trên bàn để ước tính tỷ lệ thắng. Trên mô hình thế giới do hai mạng lưới này tạo ra, thuật toán tìm kiếm cây Monte Carlo sẽ chơi thử trước hàng chục triệu ván cờ trong tâm trí. AlphaZero ra đời sau đó thậm chí không cần đến kỳ phổ của con người mà lặp lại quá trình này chỉ thông qua việc tự đấu với chính mình. Nước đi thứ 37 trong ván đấu thứ hai với Lee Sedol là nước cờ mà con người chưa từng đi trong suốt 5.000 năm qua. Đó

là kết quả của việc đi thử trước trong trí tưởng tượng trước khi thực sự đặt quân cờ xuống. Bộ não ghi nhớ và cỗ máy tính toán đang sử dụng chung một bản thiết kế.

Công ty con phát triển thuốc mới Isomorphic Labs do ông thành lập sau này cũng đứng trên cùng một nguyên lý: thay vì chỉ ghi nhớ dữ liệu quan sát thế giới, nó xây dựng một thế giới phân tử trong tâm trí và tìm kiếm trước các câu trả lời bên trong đó.

Sự ra đời của DeepMind

Vào cuối mùa hè năm 2010, Demis Hassabis đứng trong hội trường của một dinh thự lớn ở California. Đó là Hội nghị Thượng đỉnh Singularity Summit do Peter Thiel tổ chức, và thời gian dành cho ông chỉ vỏn vẹn một phút. Hassabis không hề đề cập đến các con số đầu tư. Thay vào đó, biết rằng Thiel từng là một kỳ thủ cờ vua thời trẻ, ông đã hỏi: Sự căng thẳng đầy sáng tạo mà chúng ta cảm nhận được khi đối quân Tượng lấy quân Mã trên bàn cờ rốt cuộc là gì? Tại sao máy móc lại không thể tự mình tính toán giá trị đó? Thiel đã bị lay động bởi câu nói duy nhất đó và quyết định rót vốn ban đầu 1,4 triệu bảng Anh, tương đương khoảng 1,85 triệu USD ngay tại chỗ. Một câu hỏi hay đã tỏ ra có sức nặng hơn cả một tập kế hoạch kinh doanh dày cộp.

DeepMind đã khởi đầu trước đó vào mùa thu năm 2010 trong một văn phòng tồi tàn ở London. Shane Legg, một nhà lý thuyết đến từ New Zealand, là người định nghĩa giới hạn trí thông minh mà máy móc có thể đạt tới bằng toán học. Ông lập tức đồng tình với những ý tưởng khoa học não bộ mà Hassabis đưa ra. Tiếp đó, Mustafa Suleyman, người từng bỏ học Oxford năm mười chín tuổi và tự mình vận hành một đường dây nóng tư vấn, đã gia nhập. Suleyman không biết viết mã, nhưng lại có tài thuyết phục con người và tổ chức bộ máy hơn bất kỳ ai. Kế hoạch kinh doanh chỉ vỏn vẹn hai dòng: mục tiêu giải mã bản chất của trí thông minh, và mục tiêu dùng trí thông minh đó để giải quyết tất cả các vấn đề còn lại. Đó là một kế hoạch táo bạo đến mức khiến ngay cả các nhà đầu tư ở Thung lũng Silicon cũng phải bối rối.

Tiền bạc luôn trong tình trạng thiếu thốn. Vào mùa thu năm 2013, vòng gọi vốn trị giá 65 triệu USD đã sụp đổ do mâu thuẫn trong hội đồng quản trị, đến mức công ty không còn đủ tiền mặt để trả lương cho tháng tiếp theo. Suleyman đã thuyết phục được Horizons Ventures, công ty đầu tư của Lý Gia Thành tại Hồng Kông, giúp công ty tạm thời qua cơn nguy khốn. Hassabis muốn xây dựng công ty này giống như Bell Labs, một nơi mà các nhà nghiên cứu có thể đào sâu nghiên cứu những vấn đề mình trăn trở trong thời gian dài mà không phải chịu áp lực ra mắt sản phẩm. Khi Mark Zuckerberg của Facebook đưa ra lời đề nghị thuê, ông đã từ chối, nhưng đến tháng 1 năm 2014, ông đã nắm lấy bàn tay chìa ra của Larry Page từ Google. Số tiền mua lại là 650 triệu USD, và Suleyman đã đưa vào hợp đồng điều khoản không sử dụng công nghệ này cho vũ khí hoặc hệ thống giám sát, cùng điều khoản thành lập một hội đồng đạo đức độc lập.

Học tăng cường sâu do DeepMind tạo ra là kết quả của sự giao thoa giữa hai dòng nghiên cứu. Học sâu của Geoffrey Hinton trở thành đôi mắt tự đọc các điểm ảnh trên màn hình, còn học tăng cường của Richard Sutton trở thành quy tắc chấm điểm cho kết quả của hành động. Thí nghiệm trên não khi đã đưa ra câu trả lời sớm cho sự kết hợp này. Vào cuối thập niên 1990, khi các nhà khoa học thần kinh cấy điện cực vào não khỉ và quan sát, họ nhận thấy khi một lượng nước trái cây ngoài dự

đoán xuất hiện, các tế bào thần kinh dopamine phản ứng rất mạnh. Thế nhưng khi quá trình học kết thúc và khỉ biết trước thời điểm nước trái cây xuất hiện, phản ứng đó đã biến mất. Bộ não chỉ phản ứng với mức sai số - tức sự chênh lệch giữa kỳ vọng và thực tế - chứ không phải với giá trị phần thưởng thực tế. DeepMind đã chuyển thẳng phương pháp tính toán sai số này vào hàm mất mát của mạng nơ-ron.

Thứ được tạo ra theo nguyên lý này chính là Mạng Deep Q (DQN). Khi chọn hành động a từ trạng thái màn hình s , một mạng nơ-ron duy nhất sẽ ước tính toàn bộ giá trị Q , tức tổng phần thưởng sẽ nhận được trong tương lai. Phương trình Bellman làm nền tảng cho phép tính này, khi giải ra, tương đương với việc

cộng giá trị tối ưu của khoảnh khắc tiếp theo vào giá trị của ngày hôm nay. Vấn đề có hai điểm. Vấn đề các trải nghiệm tích lũy theo thứ tự thời gian đan xen vào nhau khiến việc học bị lệch về một phía đã được giải quyết bằng phương pháp phát lại trải nghiệm đã thấy ở phần trước. Cũng có vấn đề bản thân mục tiêu bị dao động do giá trị mục tiêu và giá trị dự đoán đều xuất phát từ cùng một mạng lưới. Lần này, họ đã ngăn chặn điều đó bằng mạng mục tiêu, tức giữ một bản sao riêng của mạng lưới tính toán mục tiêu và chỉ cập nhật theo chu kỳ nhất định.

Thành quả đã hiện rõ trong trò chơi bắn gạch. Trong 100 lần đầu tiên, nó di chuyển thanh trượt một cách ngẫu nhiên và đạt điểm trung bình 1,2. Vượt qua 300 lần, nó đã đọc trước vị trí bóng rơi và đặt sẵn thanh trượt chờ đợi. Đến 500 lần, nó tự tìm ra chiến lược đục một lỗ bên hông bức tường gạch rồi đưa bóng lên phía trên trần để quét sạch điểm số mà hầu như không cần di chuyển thanh trượt. Không có bất kỳ quy tắc nào do con người chỉ dạy. Trong kết quả được công bố trên tạp chí Nature năm 2015, DQN đã giải quyết toàn bộ 57 trò chơi chỉ bằng cùng một mạng nơ-ron và cùng một bộ tham số thiết lập. Trong trò bắn gạch, nó đạt điểm số gấp hơn 13 lần mức trung bình của con người, gấp 10 lần trong trò lặn biển Seaquest, và cũng đạt điểm cao hơn con người trong Space Invaders.

Có một quy luật kỳ lạ ẩn giấu trong quá trình học tập này. Nhìn vào dự án GT Sophy học trò chơi đua xe, chỉ cần 10% tổng thời gian học là đủ để bắt kịp trình độ của một người lái xe nghiệp dư bình thường. Thế nhưng, để thu hẹp khoảng cách còn lại nhằm bắt kịp cảm giác vào cua và phanh của các tay đua chuyên nghiệp hàng đầu, nó phải tốn đến 90% thời gian còn lại. Đường cong này giống như việc ôn thi. Điểm số có thể nhanh chóng đạt mức vừa đủ qua môn, nhưng càng tiến gần đến điểm tuyệt đối, thời gian cần thiết để làm đúng thêm một câu hỏi lại tăng theo cấp số nhân. Nguyên lý này sau đó cũng xuất hiện y hệt khi AlphaGo tiến hóa thành AlphaGo Zero.

DeepMind không dừng lại ở đây. Năm 2016, họ đã chinh phục 57 trò chơi nhanh hơn bằng cấu trúc A3C cho phép nhiều mạng nơ-ron học tập đồng thời, và đến năm 2018, cùng với ngôn ngữ trong không gian mê cung 3 chiều

họ đã giới thiệu IMPALA để tìm kiếm và khám phá. Đến năm 2022, tác tử Gato thậm chí đã xuất hiện, chuyển đổi chuyển động của cánh tay robot, điểm ảnh camera và câu tiếng Anh thành một mảng giá trị liên tục duy nhất để xử lý bằng một mạng nơ-ron. Điều này tương đương với bước chuyển từ một trí tuệ nhân tạo chỉ giỏi một trò chơi sang một trí tuệ nhân tạo vượt qua ranh giới của nhiều cơ thể và nhiều ngôn ngữ.

Năm 2021, các nhà nghiên cứu của DeepMind đã công bố bài báo có tựa đề 'Phần thưởng là tất cả'. Đó là lập luận rằng các năng lực dường như tách biệt như tri giác, ngôn ngữ và suy luận nhân quả thực chất là những kết quả cùng nảy nở trong quá trình theo đuổi một phần thưởng duy nhất. Họ giải thích rằng điều này cũng giống như nguyên lý bộ não trong quá trình tiến hóa đã theo đuổi một mục tiêu duy nhất là sinh tồn mà đồng thời thu được nhiều năng lực khác nhau. Ý tưởng này đã đi qua AlphaGo chiến thắng cờ vây và tiếp nối đến AlphaFold giải mã cấu trúc protein.

Đến năm 2026, công ty con Isomorphic Labs thậm chí đã cho ra mắt công cụ thiết kế thuốc mới vượt xa AlphaFold. Đội ngũ nhỏ ở London từng mơ ước về một Bell Labs sản sinh ra mười người đoạt giải Nobel giờ đây đã tiến tới điểm có thể thiết kế các chất ứng viên thuốc mới ngay bên trong máy tính.

Máy gia tốc hạt mang tên khoa học

Chàng trai mười chín tuổi Demis Hassabis đang cầm trên tay một cuốn sách trong căn phòng nhỏ ở London. Đó là cuốn sách 'Giấc mơ về một lý thuyết tối hậu' do Steven Weinberg, người đoạt giải Nobel Vật lý, viết. Ban đầu, ông muốn trở thành một nhà vật lý lý thuyết. Ông muốn giải thích mọi thứ, từ hạt nhỏ hơn móng tay cho đến toàn bộ thiên hà, chỉ bằng một công thức duy nhất. Thế nhưng, bàn tay đang lật từng trang sách của ông bỗng khựng lại. Đó là đoạn viết về việc các nhà vật lý nghiên cứu lý thuyết

siêu dây đang phải vật lộn suốt nhiều năm trước những phép tính phức tạp với hơn mười chiều không gian. Bộ não con người chỉ có thể hình dung rõ ràng tối đa ba chiều bằng giấy và bút chì. Tuy nhiên, các quy luật thực sự của vũ trụ lại vận hành ở số chiều nhiều hơn thế rất nhiều. Hassabis quyết định đổi hướng đi ngay tại chỗ. Thay vì ước mơ tự mình giải quyết vật lý, ông quyết tâm tạo ra một cỗ máy có thể giải quyết vật lý thay cho con người. Bản thân cậu thiếu niên ấy vào đêm hôm đó hẳn cũng không thể biết rằng quyết định đưa ra nơi góc thư viện sẽ dẫn đến giải Nobel Hóa học sau vài thập kỷ.

Bên trong quyết định này chứa đựng một thế giới quan. Hassabis coi vật liệu cơ bản cấu thành nên vũ trụ không phải là vật chất hay năng lượng, mà là thông tin. Giống như Einstein đã đảo lộn nền vật lý khi chứng minh rằng khối lượng và năng lượng có thể chuyển hóa lẫn nhau, ông cho rằng thông tin cũng là một thực thể có thể hoán đổi vị trí với vật chất và năng lượng. Điều đó có nghĩa là chuỗi protein hay tín hiệu thần kinh trong não rốt cuộc cũng chỉ là những hình thái mà thông tin thay hình đổi dạng khoác lên mình. Ông coi bản thân việc làm khoa học là một công việc gần như tôn giáo. Ông từng nói rằng việc làm sáng tỏ đến cùng các quy luật sâu sắc của tự nhiên là đức tin duy nhất của mình, và trí tuệ nhân tạo là công cụ để thực hành đức tin đó.

Vì vậy, ông ví trí tuệ nhân tạo với kính viễn vọng của Galileo. Khoảnh khắc Galileo ghép hai thấu kính lồi để lần đầu tiên nhìn thấy các vệ tinh của Sao Mộc, thuyết địa tâm kéo dài hơn một nghìn năm bắt đầu sụp đổ. Nhưng rốt cuộc kính viễn vọng cũng chỉ là công cụ mở rộng tầm mắt của con người. Việc giải thích ánh sáng đi qua thấu kính vẫn là phần việc của con người. Trí tuệ nhân tạo mà Hassabis muốn tạo ra còn tiến xa hơn một bước. Đó là cỗ máy không chỉ dừng lại ở việc hiển thị dữ liệu, mà còn tự mình tìm ra các quy luật ẩn giấu bên trong dữ liệu đó. Nó tự quyết định khu vực cần tìm kiếm và đảm nhận cả việc sàng lọc những giả thuyết không nhất quán. Có thể xem như nó đã chuyển vị thế từ một công cụ quan sát thành một người đồng nghiệp cùng nghiên cứu.

Nơi mà triết lý này có được hình hài thực tế chính là công ty con phát triển thuốc Isomorphic Labs, và câu chuyện đó sẽ được tiếp nối ở phần sau.

Các thành quả thực tế cũng đã tích lũy. DeepMind đã cùng với Trường Bách khoa Liên bang Lausanne (EPFL) của Thụy Sĩ tạo ra một trí tuệ nhân tạo kiểm soát plasma bên trong lò phản ứng nhiệt hạch Tokamak. Phương pháp tính toán truyền thống phải mất hai tuần để có được một kết quả, nhưng trí tuệ nhân tạo này đã điều chỉnh lực từ theo thời gian thực ở đơn vị một phần nghìn giây để giữ plasma nóng hơn tâm Mặt Trời mười lần theo hình dạng mong muốn. Nghiên cứu này đã được công bố trên tạp chí Nature vào năm 2022. Mô hình dự báo thời tiết GraphCast cũng tương tự. Dự báo thời tiết mười ngày mà siêu máy tính truyền thống phải chạy trong hai tuần mới có được thì GraphCast đã tạo ra chỉ trong vài phút. Khi theo dõi đường đi thực tế của cơn bão Melissa, độ chính xác của nó cũng cao hơn mô hình truyền thống của Trung tâm Dự báo Thời tiết Hạn vừa Châu Âu.

Hassabis cũng đề xuất một bài kiểm tra để đánh giá xem trí tuệ nhân tạo tổng quát thực sự đã đến hay chưa. Ông gọi bài kiểm tra này là Bài kiểm tra Einstein. Trí tuệ nhân tạo chỉ được cung cấp dữ liệu vật lý trước năm 1911 và dữ liệu quan sát cho thấy quỹ đạo của Sao Thủy lệch khỏi định luật Newton. Nếu chỉ với những dữ liệu này, cỗ máy có thể tự mình đi đến kết luận rằng không-thời gian bị uốn cong và suy ra lại phương trình trường hấp dẫn năm 1915 của Einstein, thì đó chính là thời điểm trí tuệ nhân tạo làm khoa học mà không cần sự trợ giúp của con người. Ông cũng nhắc đến bóng đen của Oppenheimer. Đó là quyết tâm không lặp lại lịch sử khi bị mê hoặc bởi việc có thể phân tách urani mà nhận ra quá muộn sức nặng của sức mạnh đó tác động lên thế giới. Vì vậy, ông liên tục lặp lại lời kêu gọi thiết lập các biện pháp an toàn trước tiên tại các giảng đường trên khắp thế giới.

Quyết tâm được xác lập tại thư viện năm mười chín tuổi vẫn còn là tiêu chuẩn định hình mỗi ngày của ông ngay cả sau ba mươi năm.

Hassabis cũng giải thích lý do tại sao khoa học lại được giải quyết nhờ trí tuệ nhân tạo bằng ngôn ngữ của lý thuyết tính toán. Người ta từng cho rằng để giải quyết các hiện tượng tự nhiên có vô số biến số đan xen như thời tiết, dòng chảy rối hay sự cuộn gấp protein bằng máy tính cổ điển thì cần lượng tính toán nhiều hơn số nguyên tử trong vũ trụ. Ông đưa ra quan sát rằng giới tự nhiên không bị phân tán ngẫu nhiên đến mức như vậy. Từ trình tự axit amin của protein cho đến sự hình thành của các thiên hà, tất cả đều đã được sàng lọc qua quy luật vật lý và áp lực chọn lọc của tiến hóa suốt hàng tỷ năm để đi vào một hành lang hẹp có số chiều thấp. Ông đặt tên cho loại hiện tượng tự nhiên đã được sàng lọc này là giới tự nhiên có thể học được. Khi mạng nơ-ron lần ngược lại các mẫu hình của hành lang hẹp này trong lượng dữ liệu khổng lồ, ngay cả những bài toán mà toán học từng khẳng định là không thể giải cũng sẽ cho ra giá trị xấp xỉ chỉ trong vài giờ. Triết lý ví việc khám phá khoa học như một máy gia tốc hạt của ông đứng vững trên lý thuyết này.

Isomorphic Labs

Trên màn hình chỉ hiện ra một chuỗi ký tự khoảng chừng hai mươi chữ cái Latinh. TIM3, tên của một loại protein bám trên bề mặt tế bào miễn dịch. Loại protein này đóng vai trò như một công tắc ra lệnh cho tế bào T trong cơ thể ngừng chiến đấu. Tế bào ung thư lên lút nhấn công tắc này để vô hiệu hóa tế bào miễn dịch và bảo vệ chính mình. Do đó, các công ty dược phẩm từ lâu đã tìm kiếm loại kháng thể có thể chặn công tắc này. Tuy nhiên, trên bề mặt TIM3 không thấy có rãnh rõ rệt nào để dùng làm điểm bám. Đó là lý do tại sao các nhà hóa học suốt nhiều thập kỷ đã phải quay đầu bỏ cuộc trước loại protein này.

Công cụ của Isomorphic Labs chỉ tiếp nhận chuỗi ký tự này. Không có cấu trúc 3 chiều, cũng không có ảnh tinh thể học. Thế nhưng chỉ trong vài phút, một bề mặt cong mới đã được vẽ ra trên màn hình. Đó là vị trí để kháng thể gắn vào, tức túi liên kết. Phép tính đã hoàn tất với sai số dưới 1 angstrom, tức là trong khoảng một phần mười đường kính của một nguyên tử. Đây là điều diễn ra chỉ sau 5 năm kể từ khi công ty này được thành lập.

Isomorphic Labs ra đời vào ngày 24 tháng 2 năm 2021 với tư cách là một pháp nhân độc lập thuộc Alphabet, công ty mẹ của Google. Vị trí CEO do Demis Hassabis, người đang lãnh đạo DeepMind, kiêm nhiệm. Vài tháng trước đó, DeepMind do Hassabis dẫn dắt đã giải quyết được một bài toán hóc búa tồn tại suốt năm mươi năm bằng AlphaFold 2. Đó là bài toán dự đoán chuỗi axit amin dài sẽ cuộn gấp thành hình dạng nào. Trong phòng thí nghiệm, việc tìm ra hình dạng của một protein duy nhất phải mất nhiều năm và hàng trăm nghìn đô la. AlphaFold 2 đã hoàn thành công việc này bên trong máy tính chỉ trong vài ngày, và DeepMind đã công khai miễn phí toàn bộ hơn 200 triệu cấu trúc protein. Nhờ đóng góp này, Hassabis đã nhận giải Nobel Hóa học năm 2024 cùng với John Jumper và David Baker.

Thế nhưng, ngay cả khi nhận giải thưởng, Hassabis và Jumper vẫn đang hướng tầm mắt về một nơi khác. Bởi vì biết được hình dạng tĩnh của một protein là một chuyện, còn việc protein đó kết hợp với thuốc để chữa bệnh lại là một vấn đề hoàn toàn khác. Trong cơ thể, protein không tồn tại đơn độc. Chúng di chuyển dọc theo DNA, bám lấy các protein khác và liên kết với các hợp chất nhỏ. Phải tính toán được lực liên kết này, việc liên kết này có gây hại cho những nơi khác hay không, và liệu nó có bị phân hủy quá nhanh ở gan hay không thì mới trở thành một loại thuốc thực sự. Tên công ty Isomorphic xuất phát từ niềm tin này. Nó có nghĩa là các quy luật của sự sống và các cấu trúc toán học mà trí tuệ nhân tạo học được có thể chồng khít lên nhau trong cùng một khuôn khổ. Đó là lời tuyên bố sẽ giải quyết sinh học như một bài toán xử lý thông tin.

Công ty này đã chứng minh được giá trị của mình chỉ sau 3 năm đi vào hoạt động. Vào tháng 1 năm 2024, Eli Lilly và Novartis mỗi bên đã trả một khoản tiền trả trước để bắt tay hợp tác với công ty. Tổng giá trị của hai hợp đồng lên tới 3 tỷ đô la, tương đương 4 nghìn tỷ won. Vào tháng 3 năm 2025, công ty

đầu tư Thrive Capital của Thung lũng Silicon đã dẫn đầu khoản đầu tư 600 triệu đô la. Google Ventures và Alphabet cũng góp mặt trong danh sách này. Vào tháng 5 năm 2026, một khoản đầu tư 2,1 tỷ đô la khác do Thrive Capital tiếp tục dẫn đầu đã được công bố, nâng tổng số tiền của hai đợt đầu tư lên 2,7 tỷ đô la.

Vào tháng 2 năm 2026, công ty này đã tung ra quân bài mà họ giấu kín bấy lâu. Đó là một hệ thống mới mang tên Công cụ Thiết kế Thuốc, gọi tắt là IsoDDE. Trong ngành, hệ thống này đôi khi còn được gọi là AlphaFold 4. Theo điểm chuẩn công bố trong sách trắng, độ chính xác của IsoDDE trong thử nghiệm dự đoán hình dạng liên kết giữa protein chưa từng thấy và hợp chất đã tăng gần gấp đôi so với AlphaFold 3.

IsoDDE đã chia nhỏ quy trình thiết kế thuốc mới thành bốn giai đoạn và tự động hóa chúng. Đầu tiên, nó chỉ nhìn vào trình tự axit amin để tìm vị trí liên kết. Tiếp theo, nó tự vẽ ra cấu trúc phân tử ứng viên khớp với vị trí đó. Sau đó, nó tính toán xem phân tử đó liên kết mạnh đến mức nào, và cuối cùng sàng lọc xem nó có xung đột với 20.000 loại protein trong cơ thể hay không, cũng như có tồn tại được trong gan và thận hay không. Phép tính ái lực liên kết mà phần mềm truyền thống phải chạy hàng nghìn siêu máy tính trong nhiều ngày mới xong thì IsoDDE đã rút ngắn xuống còn đơn vị mili giây. Sai số nhỏ đến mức gần như trùng khớp với giá trị đo được bằng thiết bị thí nghiệm chính xác.

Hiện tại, công ty đang đồng thời triển khai từ 18 đến 19 ứng viên thuốc mới nhắm vào ung thư, bệnh tự miễn và bệnh tim mạch. Kháng thể nhắm vào TIM3 chỉ là một trong số đó. Các chất ứng viên này đã vượt qua các thử nghiệm tế bào trên chuột và linh trưởng trong suốt năm 2025. Phân tử sinh ra trên màn hình đã bước sang cơ thể sống thực sự. Công ty thông báo sẽ bắt đầu thử nghiệm lâm sàng đầu tiên trên người vào cuối năm 2026.

So với phương pháp cũ, sự khác biệt càng trở nên rõ ràng hơn. Các công ty dược phẩm truyền thống từng vận hành quy trình sàng lọc thông lượng cao bằng cách dùng cánh tay robot nhỏ từng hợp chất trong số một triệu hợp chất vào ống nghiệm. Với phương pháp này, phải mất trung bình 4 đến 5 năm và ngân sách gần 1 tỷ đô la để tinh chỉnh cấu trúc của một ứng viên thuốc mới. Ngay cả những chất ứng viên được chọn lựa công phu như vậy cũng có tới 90% thất bại ngay trước ngưỡng cửa thử nghiệm lâm sàng. Isomorphic Labs đã đảo ngược phép tính này. Phòng thí nghiệm ướt chỉ cần tiếp nhận khoảng mười ứng viên đã được sàng lọc qua xác thực máy tính để thực hiện bước kiểm tra cuối cùng.

Khoản đầu tư 2,1 tỷ đô la được công bố vào ngày 13 tháng 5 năm 2026 không chỉ có Alphabet và Google Ventures mà còn có thêm các nhà đầu tư mới như MGX, Temasek, CapitalG và Quỹ đầu tư quốc gia Anh. Trước bức tường mười năm và một tỷ đô la để đưa một loại thuốc mới ra thế giới, công ty này đang đưa ra những con số chỉ tính bằng vài tuần và vài triệu đô la. Hassabis từng nói rằng khi chi phí phát triển thuốc mới giảm xuống, bàn tay phát triển có thể chạm tới cả những căn bệnh mà các công ty dược lớn chưa từng đụng tới.

Tế bào ảo

Con hẻm phía sau ga King's Cross ở London vẫn sáng rực ánh đèn từ các cửa sổ viện nghiên cứu ngay cả vào ban đêm. Hiệp sĩ Demis Hassabis đi dạo trên con hẻm này và cùng Hiệp sĩ Paul Nurse trao đổi về một câu hỏi suốt nhiều thập kỷ. Sơ đồ con đường chuyển hóa tế bào được vẽ trong sách giáo khoa chỉ gần giống như một bức tranh hoạt hình gồm các mũi tên và hình tròn. Hai người tự hỏi liệu việc đưa toàn bộ các phản ứng hóa học và chuyển động của protein diễn ra bên trong tế bào thực vào trong chip máy tính có khả thi về mặt vật lý hay không.

Hiệp sĩ Paul Nurse đã nhận giải Nobel Sinh lý học hoặc Y học năm 2001 và lãnh đạo Viện Francis Crick. Từ lâu, ông đã cho rằng sự sống phải được nhìn nhận như một hệ thống xử lý thông tin tinh vi chứ không phải là sự tiếp xúc cơ học của các chất hữu cơ. Người dẫn dắt Hassabis, người từng trải qua lĩnh

vực trò chơi và khoa học thần kinh, bước vào thế giới sinh học chính là Paul Nurse.

Mỗi khi câu hỏi này được gọi lại sau mỗi 5 năm, câu trả lời luôn là không. Bởi vì ngay cả việc tìm ra hình dạng cuộn gấp của một protein duy nhất cũng là bài toán mà các phòng thí nghiệm hàng đầu thế giới vẫn chưa thể giải quyết cho đến năm 2020.

Cuối năm 2020, AlphaFold 2 đã giải quyết bài toán này. Năm sau, vào năm 2021, Hassabis đã đến gặp Paul Nurse và nói rằng đã đến lúc bắt đầu dự án tế bào ảo. Đó là lời tuyên bố sẽ đưa toàn bộ tế bào sống vào máy tính trên nền tảng vững chắc mang tên AlphaFold.

Sinh vật được chọn làm mục tiêu đầu tiên là nấm men dùng để làm bánh mì và bia, *Saccharomyces cerevisiae*. Ở nấm men, một tế bào đơn lẻ chính là toàn bộ sinh vật. Nó không chịu sự can thiệp từ các cơ quan hay mô khác như tế bào người. Đồng thời, là một tế bào nhân thực có nhân, nó chia sẻ khung cấu trúc cơ bản với tế bào người và tế bào thực vật. Hơn nữa, trong suốt một trăm năm qua, các nhà vi sinh vật học trên toàn thế giới đã ghi chép tỉ mỉ về các gen và con đường chuyển hóa của nấm men. Có thể nói họ đã chọn một sinh vật mà dữ liệu đã được tích lũy sẵn.

Việc tạo ra một tế bào ảo không dừng lại ở một mạng nơ-ron đơn lẻ. Tăng tính toán lực ở cấp độ nguyên tử, tăng tái hiện hình dạng khi protein và DNA (axit deoxyribonucleic) gắn kết với nhau, tăng biểu diễn các phản ứng hóa học bên trong tế bào dưới dạng mạng lưới lần lượt được xếp chồng lên nhau. Vấn đề nằm ở thang đo thời gian. Thời gian để một phân tử dao động tính bằng một phần nghìn tỷ giây, còn thời gian để một tế bào phân chia tính bằng vài giờ. Nếu giải sự chênh lệch thang đo này từng chút một bằng phương trình Newton thì ngay cả siêu máy tính cũng mất nhiều tháng. Thay vì lần theo từng nguyên tử riêng lẻ, DeepMind đã chọn trạng thái của khối protein

và vượt qua bức tường này bằng cách tích hợp một mô hình thế giới dự đoán toàn bộ sự biến đổi.

Các phản ứng hóa học bên trong tế bào được liên kết lại thành một mạng lưới. Tại mỗi nút là một chất chuyển hóa, và tại mỗi đường nối các nút là một gen mã hóa enzyme. Ngay cả mô hình nấm men do robot Adam xử lý—sẽ được đề cập ở phần sau—cũng có hơn 1.000 gen và hơn 1.000 chất chuyển hóa đan xen lẫn nhau bên trong mạng lưới này. Một người nếu muốn tự tay vẽ và thấu hiểu mạng lưới này thì cả đời cũng không đủ.

Chỉ riêng trên cơ sở dữ liệu tài liệu y sinh PubMed, mỗi năm đã có hơn 1 triệu bài báo khoa học mới được tích lũy. Tri thức ngày càng nhiều nhưng chỉ nằm rải rác từng mảnh vụn trong đầu của mọi người, không có một bộ não nào trên Trái Đất có thể nhìn thấu toàn bộ tế bào trong một cái nhìn. Trí tuệ nhân tạo tế bào ảo đọc lượng tài liệu khổng lồ này trong chớp mắt và xâu chuỗi chúng lại thành một mạng lưới duy nhất.

Tính hiệu quả thực tế của phương pháp này đã được các nhà khoa học robot xác nhận. Quá trình mà các robot Adam và Eve của Ross King tự mình thực hiện từ việc đặt giả thuyết đến tiến hành thí nghiệm trong nghiên cứu về nấm men và bệnh sốt rét sẽ được trình bày chi tiết ở Chương 8.

Quy trình phát triển thuốc truyền thống ngay cả sau khi tìm ra một chất ứng viên vẫn phải huy động hàng nghìn con chuột và khỉ để trải qua các cuộc kiểm tra độc tính kéo dài nhiều năm. Khi tế bào ảo hoàn thiện, trước khi đưa thuốc vào cơ thể sống thực tế, người ta có thể đưa thuốc qua các tế bào gan và tế bào tim ảo trước để kiểm tra phản ứng. Công cụ thiết kế thuốc mới của Isomorphic Labs đã thấy ở phần trước đã giảm đáng kể số lượng chất cần phải trực tiếp tổng hợp và xác thực trong phòng thí nghiệm. Vào tháng 7 năm 2026, khi DeepMind và Isomorphic Labs công bố phương án chung nhằm ứng phó với các mối đe dọa sinh học, họ đã tuyên bố rằng các phòng thí nghiệm trong tương lai phải chuyển đổi từ nơi tạo ra dữ liệu thành nơi xác nhận những dự đoán do máy tính đưa ra. Phép tính bắt đầu từ một tế bào nấm men đang từng bước tiến tới vị trí ngăn ngừa bệnh tật cho con người.

Bức tranh cuối cùng mà Hassabis vẽ ra không dừng lại ở việc chữa trị bệnh tật. Ông nói rằng sẽ tái hiện lại thông qua mô phỏng máy tính về cách sự sống tự hình thành màng bao bọc và bắt đầu tự nhân bản từ món súp hóa học của Trái Đất nguyên thủy nơi các axit amin từng trôi nổi. Đó là việc tua lại những gì quá trình tiến hóa đã tạo ra qua hàng tỷ năm chỉ bằng vài phút tính toán của máy chủ. Việc thấu hiểu hoàn toàn một loài nấm men chính là bước đi đầu tiên.

Những vấn đề vẫn chưa được giải quyết

Không có bất đồng nào về việc AlphaFold dự đoán chính xác hình dạng gấp cuộn của chuỗi axit amin với sai số trong phạm vi đường kính của một nguyên tử. Tuy nhiên, giới sinh học cấu trúc liên tục chỉ ra rằng việc đoán đúng cấu trúc và việc thấu hiểu cấu trúc đó là hai phạm trù khác nhau. Các nhà nghiên cứu, bao gồm nhà sinh học cấu trúc Peter Moore của Đại học Yale, đã chỉ ra rằng AlphaFold không giải thích được lý do tại sao protein lại gấp lại theo đúng hình dạng đó, cũng như cách nó rung lắc, di chuyển và liên kết với các phân tử khác bên trong tế bào. Độ chính xác của dự đoán không đồng nghĩa với sự thấu hiểu về quan hệ nhân quả. Thứ mà AlphaFold đưa ra chỉ là một khung cảnh tĩnh, trong khi động học của protein sống liên tục thay đổi hình thái và vận hành theo từng khoảnh khắc lại không nằm trong khung cảnh đó.

Sự phân biệt này gắn liền với những phê phán về chủ nghĩa dự đoán đã được thảo luận từ lâu trong giới triết học khoa học. Việc một mô hình dự đoán chính xác kết quả không đồng nghĩa với việc mô hình đó giải thích đúng đắn cơ chế của hiện tượng. Mạng nơ-ron truy ngược các quy luật thống kê từ lượng dữ liệu khổng lồ có thể chỉ đưa ra câu trả lời mà không cung cấp lời giải thích vật lý về lý do tại sao quy luật đó lại hình thành. Giới học thuật chỉ ra rằng giữa việc có được câu trả lời và việc hiểu được nguyên do vẫn tồn tại một khoảng trống chưa thể lấp đầy.

Quan điểm thận trọng cũng được đặt ra đối với lộ trình tế bào ảo. Ý tưởng đưa toàn bộ gen và mạng lưới trao đổi chất của một loài nấm men vào máy tính chưa phải là thành tựu đã được kiểm chứng, mà gần như một lời hứa cần được thực hiện trong tương lai. Có ý kiến cho rằng phương pháp dùng mô hình thế giới để vượt qua khoảng cách thang thời gian giữa sự dao động của nguyên tử và sự phân chia của tế bào vẫn chưa được chứng minh đầy đủ về khả năng tái hiện hành vi thực tế của tế bào. Từ việc dự đoán hình dạng gấp cuộn đến việc tái hiện hoạt động sống của toàn bộ tế bào vẫn còn một khoảng cách dài ngăn cách giữa thành tựu và lời hứa.

Nội dung trong chương này được trình bày dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Chương 2. John Jumper: Người giải bài toán 50 năm

Nhà hóa học tình cờ

Đó là bữa trưa trước phòng thí nghiệm của khoa Vật lý, Đại học Vanderbilt. Thức ăn trên khay nguội dần trong khi cuộc trò chuyện vẫn chỉ xoay quanh cùng một chủ đề. Cậu sinh viên đại học John Jumper hỏi khi nào dự án máy dò hạt mà cậu đang theo đuổi — chính là thí nghiệm BABAR, một thí nghiệm máy gia tốc lớn của Mỹ — mới đưa ra dữ liệu thực sự. Nhà vật lý đàn anh ngồi cạnh trả lời với vẻ mặt dừng dưng. Anh ta bảo rằng có lẽ đến tận lúc mình nghỉ hưu cũng chẳng được thấy thiết bị đó phát ra tín hiệu, và phải đến thế hệ nghiên cứu sinh tiếp theo mới nhận được những con số ấy. Các đàn anh khác cũng gật đầu như thể đó là điều hiển nhiên. Hôm đó, trên bảng thông báo của phòng thí nghiệm vẫn dán lịch thu thập dữ liệu của tháng sau, nhưng chẳng ai bận tâm nhìn tới. Nghe vậy, Jumper đặt thìa xuống. Đó là một ngành học mà người ta phải chờ đợi trọn một thế hệ chỉ để thấy được một kết quả thí nghiệm.

Jumper sinh năm 1985 tại Little Rock, bang Arkansas, Mỹ. Từ nhỏ, cậu đã say mê cấu trúc đối xứng của toán học và các hiện tượng điện từ, luôn mong muốn trở thành một nhà vật lý tìm kiếm những quy luật cơ bản của vũ trụ. Sau khi tốt nghiệp Học viện Pulaski năm 2003, cậu theo học tại Đại học Vanderbilt ở Nashville, bang Tennessee, học song song cả vật lý và toán học. Giáo sư Mel Webster, người hướng dẫn cậu thời đại học, luôn truy vấn đến cùng nếu thấy một kẽ hở nhỏ trong lập luận, ngay cả khi đáp án đã đúng. Có lần Jumper đã tính ra chính xác giá trị thực nghiệm, nhưng Giáo sư Webster vẫn chỉ ra một điểm sai lệch nhỏ về xác suất trong quá trình tính toán và bắt cậu giải lại từ đầu. Đó là sự rèn luyện để coi trọng căn cứ hơn là đáp số. Thói quen này sau này đã dẫn đến công việc cấy các định luật vật lý vào mạng nơ-ron. Mỗi lần bước ra khỏi phòng thí nghiệm của Giáo sư Webster, đôi vai của Jumper trĩu nặng, nhưng đồng thời cậu cũng cảm thấy tư duy logic của mình ngày càng thêm vững chắc.

Sự thất vọng từ thí nghiệm BABAR đã làm thay đổi hướng đi của Jumper. Cậu tốt nghiệp thủ khoa Đại học Vanderbilt và nhận học bổng Marshall để sang Đại học Cambridge ở Anh. Học bổng Marshall chỉ tuyển chọn vài chục tài năng trẻ mỗi năm từ Anh và Mỹ, bản thân nó đã là minh chứng cho thành tích xuất sắc của cậu thời đại học. Cậu theo học tại Trường St Edmund's và học chương trình thạc sĩ vật lý chất ngưng tụ lý thuyết tại Phòng thí nghiệm Cavendish. Đây là lĩnh vực nghiên cứu các nguyên lý vận động tập thể của vô số electron bên trong chất rắn. Jumper miệt mài giải các trạng thái electron trong chất siêu dẫn bằng công thức toán học. Thế nhưng những công thức trên bảng đen vẫn còn quá xa vời với thực tế. Trở về Mỹ, cậu theo học chương trình tiến sĩ vật lý, nhưng việc từng chút một

cập nhật các lý thuyết cũ kỹ trong một cuộc sống như nhà máy sản xuất bài báo đã khiến cậu thất vọng. Cậu quyết định bỏ dở chương trình tiến sĩ chỉ sau một năm nhập học.

Thời điểm đó mùa nộp hồ sơ sau đại học đã kết thúc nên không có nhiều nơi để đi. Trong lúc cân nhắc tìm việc ở ngành tài chính, D. E. Shaw Research đã lọt vào tầm mắt của cậu. Viện nghiên cứu do nhà sáng lập quỹ phòng hộ David Shaw thành lập này đã tự chế tạo và vận hành Anton, một siêu máy tính chuyên dụng cho mô phỏng protein. Tại đây suốt 3 năm, Jumper đã quan sát trên màn hình quá trình protein cuộn gập theo phương trình chuyển động của Newton. Anton là cỗ máy đầu tiên mô tả liên mạch quỹ đạo cuộn gập của một protein lên tới khoảng thời gian 1 mili giây. Tuy nhiên, tốc độ tính toán lại chậm đến mức tàn nhẫn. Protein trong tế bào cuộn gập chỉ trong vài mili giây, nhưng máy tính phải chia nhỏ quá trình đó thành từng đơn vị 2 femto giây để tính toán. Anton phải chạy liên tục nhiều tháng mới có thể tái hiện được một lần cuộn gập. Hơn nữa, bản thân các công thức vật lý xác định độ lớn của lực lại là các giá trị xấp xỉ do con người tự điều chỉnh thủ công, và kết quả tạo ra thường không khớp với cấu trúc protein thực tế chụp bằng tia X hoặc kính hiển vi điện tử nghiệm lạnh. Một sai số nhỏ cũng làm biến dạng toàn bộ kết quả. Chính lúc này, Jumper đã quyết tâm phải tìm ra cách để học các định luật vật lý từ dữ liệu.

Bước ngoặt tiếp theo trong sự nghiệp lại đến từ cuộc sống hôn nhân. Khi vợ cậu, Carolyn, trúng tuyển chương trình tiến sĩ di truyền học tại Đại học Chicago, Jumper cũng phải tiếp tục việc học ở cùng thành phố. Cậu định nộp đơn vào khoa Vật lý của Đại học Chicago nhưng lại nhằm hạn chót là nửa đêm thay vì 5 giờ chiều, dẫn đến việc bỏ lỡ cơ hội nộp hồ sơ. Cậu đã kiên trì xin xem xét lại suốt một tháng nhưng bị khoa Vật lý từ chối. Đúng lúc đó, trong một bữa trưa, một đồng nghiệp ở bộ phận hóa học tính toán tại D. E. Shaw Research đã gợi ý khoa Hóa học của Đại học Chicago. Jumper chưa từng học một lớp hóa học chính quy nào kể từ thời trung học. Ký ức về những ngày học thuộc ký hiệu nguyên tố đã quá xa xôi. Dẫu vậy, cậu tin rằng hóa học rất cuộc cũng là một nhánh của vật lý nghiên cứu sự chuyển động của electron nên đã nộp đơn. Khoa Hóa học đã đặc cách mở lại cổng tiếp nhận hồ sơ trong một ngày cho cậu, sau khi xem xét hồ sơ đã quyết định nhận cậu vào học. Thế là cậu tự gọi mình là nhà hóa học tình cờ. Để được miễn học phí, cậu phải hướng dẫn thí nghiệm hóa học cho sinh viên năm nhất, và cậu đã vượt qua quãng thời gian làm trợ giảng bằng cách tự học trước sách giáo khoa đúng một tuần so với sinh viên mỗi tuần.

Trong chương trình tiến sĩ kéo dài từ năm 2011 đến năm 2017, cậu được Giáo sư Tobin Sosnick và Giáo sư Karl Freed hướng dẫn. Giáo sư Sosnick xử lý dữ liệu thực tế về quá trình protein cuộn gấp trong nước, còn Giáo sư Freed giảng dạy toán học chặt chẽ của vật lý thống kê. Jumper đã kết hợp hai mảng này để biểu diễn protein theo đơn vị các khối axit amin thay vì từng nguyên tử

riêng lẻ, rồi đúc kết thành luận văn. Điều này giống như thay vì xử lý từng khối Lego riêng lẻ, người ta xử lý toàn bộ các mảnh ghép đã được gom lại thành cụm. Các tham số lực vốn được con người căn chỉnh thủ công cũng được học từ dữ liệu. Ý tưởng kết hợp các ràng buộc vật lý với dữ liệu học máy này sau này đã trở thành bản thiết kế để xây dựng AlphaFold.

Vào tháng 6 năm 2026, Jumper rời Google DeepMind sau 9 năm gắn bó để chuyển sang Anthropic. Câu chuyện đằng sau sẽ được đề cập trong hai tiểu mục cuối của chương này. Giống như việc cậu bé say mê vật lý ở Arkansas trở thành nhà hóa học nhờ nhằm hạn chót, ở tuổi bốn mươi mốt, anh lại một lần nữa chọn cho mình một bàn thí nghiệm mới lạ. Hành trình sự nghiệp của anh cho thấy việc chuyển đổi chuyên ngành không phải là hồ sơ của sự thất bại, mà là dấu mốc của việc định hình lại hướng đi.

Cây gậy chỉ huy sau 6 tháng

Mùa thu năm 2017, tại văn phòng Google DeepMind ở King's Cross, London, một chàng trai 32 tuổi được nhận một bàn làm việc. Đó là John Jumper, người mới nhận bằng tiến sĩ được 6 tháng và vừa tốt nghiệp khoa Hóa học của Đại học Chicago. Demis Hassabis, Giám đốc điều hành của DeepMind, đã đích thân bước tới bên cạnh anh. Trên tay Hassabis là cuốn luận án tiến sĩ của Jumper. Đó là công trình nghiên cứu tính toán chuỗi protein bằng cách gộp theo đơn vị gốc cặn, đồng thời để dữ liệu tự tìm ra độ lớn của lực liên kết thay vì con người phải chỉnh tay. Trong đó, Hassabis đã nhìn thấy gương mặt của người sẽ giải quyết thách thức tiếp theo của DeepMind sau AlphaGo: bài toán cuộn gấp protein tồn tại suốt 50 năm. Một người mới tốt nghiệp như vậy đã bước lên vị trí trưởng nhóm sinh học tại DeepMind. Xét theo tiền lệ lâu đời của giới học thuật, đây là một quyết định bổ nhiệm không tưởng. Một sinh viên vật lý chưa từng theo học trọn vẹn một bài giảng hóa học chính quy nào lại đảm nhận vị trí lãnh đạo dự án sinh học của viện nghiên cứu hàng đầu thế giới. Thế nhưng Hassabis đã đưa ra phán đoán dựa trên lối tư duy thể hiện trong một bài báo khoa học thay vì tuổi tác và kinh nghiệm.

Vào thời điểm Jumper gia nhập DeepMind, AlphaFold 1 đã giành vị trí số một thế giới tại cuộc thi CASP13 năm 2018. CASP là tên cuộc thi nơi các phòng thí nghiệm trên toàn thế giới so tài dự đoán cấu trúc của các protein chưa được công bố. Điểm số tại đây được gọi là GDT, một thang điểm 100 đo mức độ trùng khớp giữa hình dạng do máy tính đưa ra và hình dạng thực tế được kiểm chứng bằng thực nghiệm. Dù đứng đầu thế giới, trong nội bộ nhóm nghiên cứu sự lo lắng lại lấn át niềm vui chúc mừng.

AlphaFold 1 có cấu trúc gồm hai giai đoạn: đầu tiên vẽ khoảng cách giữa các gốc axit amin thành bản đồ 2D, sau đó chuyển bản đồ đó sang một chương trình riêng biệt mang tên Rosetta để ghép thành hình dạng 3D. Nếu ví với khuôn mặt người, điều này giống như việc chỉ đánh dấu vị trí các đường nét ngũ quan trên giấy rồi giao việc tạc khuôn mặt thực tế cho một nhà điêu khắc khác. Với phương pháp này, điểm số bị chững lại ở mức quanh 70 điểm GDT. Để các công ty được phẩm có thể thiết kế thuốc mới, sai số của từng nguyên tử phải dưới 1 angstrom (tức một phần mười của một phần tỷ mét), tương đương mức vượt qua 90 điểm GDT, nhưng bức tường đó vẫn không hề suy chuyển.

Tròn sáu tháng sau khi gia nhập, Jumper tìm đến Hassabis và đưa ra một đề xuất táo bạo. Đó là xóa bỏ toàn bộ mã nguồn dự đoán khoảng cách đã xây dựng từ trước đến nay và viết lại từ đầu. Ý tưởng là hợp nhất người vẽ ngũ quan và người tạc khuôn mặt thành một, xử lý đồng thời thông tin tiến hóa và tọa độ 3D ngay trong một mạng nơ-ron duy nhất. Đề xuất hủy bỏ chỉ sau một đêm toàn bộ mã nguồn đã dày công xây dựng suốt nhiều năm ắt hẳn đã khiến bầu không khí phòng họp đóng băng. Nhưng Hassabis đã chấp thuận đề xuất đó. Dự án AlphaFold 2 đã bắt đầu như vậy. Thiết kế của mạng nơ-ron mới mà nhóm nghiên cứu

đã xây dựng trong suốt hai năm sẽ được tìm hiểu chi tiết ở tiểu mục tiếp theo.

Tháng 11 năm 2020, kết quả cuộc thi CASP14 được công bố. AlphaFold 2 đã đạt điểm GDT trung bình là 92,4 đối với các protein mà cấu trúc thậm chí còn chưa được làm sáng tỏ bằng thực nghiệm. Đội tham gia có thành tích tốt thứ hai chỉ đạt 90,8 điểm. Đây là con số tương đương với độ chính xác mà con người phải mất nhiều năm miệt mài với tia X hoặc kính hiển vi điện tử nghiệm lạnh mới có được. Sai số giữa tọa độ nguyên tử dự đoán và tọa độ thực tế có giá trị trung vị là 0,96 angstrom, một khoảng cách thậm chí còn nhỏ hơn đường kính của một nguyên tử. Tiến trình này là một bước nhảy vọt, bắt đầu từ 38 điểm tại cuộc thi CASP11 năm 2014, 40 điểm tại CASP12 năm 2016 và vươn tới 92,4 điểm tại CASP14 năm 2020. Đồ thị vốn chỉ nhích từng bước chậm chạp với những con số hai chữ số đã dựng đứng lên chỉ sau hai năm kể từ khi Jumper tiếp quản nhóm.

Thành quả này không dừng lại ở điểm số các cuộc thi. Vào năm 2025 và 2026, các nhà nghiên cứu của DeepMind đã sử dụng AlphaFold để giải mã cơ chế hoạt động của protein midnolin, vốn là một ẩn số trong hơn 40 năm. Tế bào người tự dọn dẹp các protein hết hạn sử dụng, nhưng cách thức "người dọn dẹp" này giữ lấy mục tiêu từng là một câu đố lâu năm. AlphaFold đã tìm ra hai điểm trên bề mặt midnolin giữ các chất khác theo hình chiếc kẹp. Các phòng thí nghiệm tại Harvard và MIT đã dùng kéo di truyền cắt bỏ chính xác hai điểm đó để kiểm chứng. Ngay lập tức, quá trình dọn dẹp các protein cũ bên trong tế bào đã dừng lại hoàn toàn. Đó là khoảnh khắc mà dự đoán trên màn hình máy tính đã khớp chính xác bên trong tế bào sống. Bước nhảy vọt 6 tháng của Jumper — từ một nghiên cứu sinh thức trắng đêm học thuộc sách giáo khoa hóa học trong thư viện Chicago trở thành nhà thiết kế cấy các định luật vật lý vào mạng nơ-ron — đã được hoàn tất như thế.

Sự ra đời của AlphaFold 2

Tháng 12 năm 2018, hội trường hội nghị tại Cancún, Mexico ngập tràn trong bầu không khí lễ hội. Đó là vì AlphaFold 1, trí tuệ nhân tạo đầu tiên do Google DeepMind đưa ra tranh tài tại Hội nghị Dự đoán Cấu trúc Protein lần thứ 13 (CASP13), đã giành chức vô địch. Giữa lúc các đồng nghiệp cùng nhau nắm tay reo hò, duy chỉ có Tiến sĩ John Jumper, nhà nghiên cứu mới gia nhập DeepMind, là không rời mắt khỏi màn hình máy tính xách tay. Anh đang bóc tách từng chi tiết bên trong mô hình vừa chiến thắng.

Anh là người tò mò về điểm dừng của kỷ lục hơn là bản thân kỷ lục mà đội của mình lập nên.

Nội tình bên trong AlphaFold 1 mà Jumper nhìn vào hoàn toàn khác với vẻ ngoài hào nhoáng của nó. Mô hình này gần giống như một sản phẩm lắp ghép bê nguyên mạng nơ-ron tích chập (CNN) vốn dùng trong nhận diện hình ảnh sang dữ liệu protein. Mạng nơ-ron vẽ bản đồ khoảng cách nhận đầu vào là

giống hàng đa trình tự (MSA) — một bảng so sánh đối chiếu các protein tương tự của nhiều loài sinh vật khác nhau đặt cạnh nhau. Hai giai đoạn được đề cập ở tiểu mục trước là hai chương trình riêng biệt, và giai đoạn 1 không có cách nào tiếp nhận lại sai số từ giai đoạn 2 để học hỏi. Từng bộ phận riêng lẻ có thể rất thông minh, nhưng nếu các kênh kết nối giữa chúng bị nghẽn thì hiệu năng tổng thể cuối cùng sẽ dừng lại ở một điểm. AlphaFold 1 cũng như vậy. Bức tường 70 điểm đã thấy ở tiểu mục trước chính là giới hạn bắt nguồn từ sự tắc nghẽn này.

Việc Jumper tiếp nhận tình trạng bế tắc này và đề xuất loại bỏ toàn bộ mã nguồn cũ đã diễn ra như đã trình bày ở tiểu mục trước. Anh chỉ đặt ra một nguyên tắc duy nhất: cùng học thông tin trình tự do tiến hóa để lại và cấu trúc hình học 3D nơi phân tử thực sự tồn tại ngay trong một mạng nơ-ron duy nhất, đảm bảo tính vi phân liên mạch từ đầu đến cuối.

Thành quả ra đời từ nguyên tắc này là Evoformer. Evoformer là một cơ chế tính toán được thiết kế để trao đổi qua lại 48 lần giữa hai không gian: biểu diễn MSA chứa thông tin trình tự và biểu diễn cặp chứa khoảng cách giữa các cặp gốc cận. Ở phía MSA, phép tính cơ chế chú ý theo trục luân phiên quét theo chiều hàng và cột để tìm ra các cặp gốc cận cùng biến đổi trong quá trình tiến hóa. Ở phía biểu diễn cặp, anh đã đưa vào phép tính cập nhật tích tam giác. Định luật vật lý quy định rằng khoảng cách giữa ba điểm không thể vi phạm quy tắc độ dài các cạnh của một tam giác đã được Jumper khắc trực tiếp vào mạng nơ-ron bằng các công thức toán học. Sau khi đưa cơ chế này vào, tỷ lệ bản đồ khoảng cách vi phạm quy tắc đó đã giảm xuống dưới một phần tư.

Phần tiếp nhận bản đồ do Evoformer vẽ ra là mô-đun cấu trúc. Tại đây, khung xương axit amin được coi như một khung tam giác cứng, và mỗi gốc cận được gắn một hệ quy chiếu vật rắn thể hiện phép quay và phép tịnh tiến. Jumper đã trực tiếp sáng chế ra phương thức chú ý mới để xử lý hệ quy chiếu này: cơ chế chú ý điểm bất biến (Invariant Point Attention). Đây là phương thức đưa trực tiếp vào tính toán đặc tính khoảng cách giữa các gốc cận không đổi bất kể toàn bộ cấu trúc quay theo hướng nào. Anh cũng tạo ra một hàm mất mát mới: Sai số điểm căn chỉnh theo hệ quy chiếu, viết tắt là FAPE. Phương pháp này đã khắc phục điểm yếu của các phép tính cũ, vốn khiến toàn bộ phần sau bị văng sang vị trí sai lệch hoàn toàn chỉ vì một góc ở giữa chuỗi hơi bị lệch. FAPE đặt một gốc cận làm gốc tọa độ tạm thời và đo lại vị trí của tất cả các gốc cận khác dựa trên gốc tọa độ đó. Ngay cả khi toàn bộ cấu trúc bị xoay hoặc dịch chuyển, khoảng cách tương đối này vẫn được giữ nguyên. Do đó, khi chỉ một đoạn chuỗi bị lệch, chỉ đoạn đó phải chịu điểm phạt, còn tín hiệu học tập của các phần còn lại không bị bóp méo.

Điểm xuất phát của mô-đun cấu trúc cũng ẩn chứa một quyết định táo bạo. Jumper cho phép tất cả các gốc cận tạo nên protein xếp chồng hoàn toàn lên nhau tại một điểm gốc tọa độ duy nhất để bắt đầu tính toán. Các đồng nghiệp gọi phương pháp này là khởi tạo lỗ đen. Dù bắt đầu bằng cách phớt lờ cả các định luật vật lý lẫn quy tắc chống va chạm, nhưng khi quá trình học diễn ra, các gốc cận tập trung tại gốc tọa độ bắt đầu bung tỏa ra bốn phía như một cuốn sách gấp tự mở ra và tìm về đúng vị trí của mình. Từng có đề xuất đưa vào các chốt an toàn để ép giữ khoảng cách nguyên tử, nhưng Jumper đã gạt bỏ chốt an toàn đó. Bởi lẽ nếu mạng nơ-ron có thể tự mình tìm lại các định luật vật lý, thì thiết bị đó vốn dĩ là một chiếc nạng không cần thiết ngay từ đầu.

Kết quả của việc cấy các định luật vật lý vào chính cấu trúc mạng nơ-ron đã được thể hiện qua các con số. AlphaFold 1 phải sử dụng toàn bộ khoảng 150.000 cấu trúc tích lũy trong cơ sở dữ liệu cấu trúc protein để đạt mức 70 điểm GDT. Khi phòng thí nghiệm của Giáo sư Mohammed AlQuraishi tại Đại học Columbia lấy nguyên cấu trúc của AlphaFold 2 để huấn luyện lại, họ đạt được cùng số điểm đó chỉ với 1.500 cấu trúc, tức 1% tổng số. Điều đó đồng nghĩa với việc chính thiết kế, chứ không phải dữ liệu, đã nâng cao hiệu suất học lên gấp 100 lần. Một sự đảo ngược tương tự cũng xảy ra khi nhóm của Jumper hoàn thiện bài báo, loại bỏ toàn bộ các lớp tích chập còn lại trong biểu diễn cặp và chỉ giữ lại cơ chế chú ý theo trục. Số lượng tham số giảm đi nhưng mất mát kiểm định lại được cải thiện. Theo thường thức học

được trong các lớp học máy, khi giảm dung lượng tính toán thì khả năng biểu diễn cũng phải giảm theo. Nhưng đối với protein, kết quả hoàn toàn ngược lại. Do đặc tính chỉ nhìn vào các vị trí lân cận, tích chập đã bỏ sót các cặp gốc cận ở xa nhau trên chuỗi nhưng lại gắn liền với nhau trong không gian 3D. Jumper đã chứng minh bằng số liệu thực tế rằng một công cụ rất phù hợp cho hình ảnh lại có thể gây cản trở cho các phân tử sinh học. Khi thiết kế đã được chứng minh, việc áp dụng mô hình này cho tất cả

protein trên Trái Đất đã trở thành nhiệm vụ tiếp theo.

200 triệu cấu trúc

Giáo sư Matthew Higgins đã vấp phải cùng một bức tường suốt hơn một thập kỷ trong phòng thí nghiệm tại Đại học Oxford. Ông cần kết tinh protein bề mặt Pfs48/45 của ký sinh trùng truyền bệnh sốt rét để quan sát hình dạng không gian 3D của nó, nhưng loại protein này nhất quyết không chịu đông đặc thành chất rắn. Mỗi lần kết tinh thất bại, manh mối để thiết kế kháng thể cũng biến mất theo. Thế nhưng sau khi AlphaFold 2 được công bố vào năm 2021, vấn đề này đã được giải quyết chỉ trong ba mươi phút. Nằm trong tay các tọa độ nguyên tử do máy tính phác họa, nhóm nghiên cứu đã xác định chính xác vị trí mà kháng thể sẽ bám vào. Kết quả này đã dẫn đến vắc-xin sốt rét R21/Matrix-M được Tổ chức Y tế Thế giới phê duyệt vào năm 2023.

Đằng sau phép màu này là một khoảng cách chênh lệch đã âm ỉ từ lâu. Khi công nghệ giải trình tự gen thế hệ mới trở nên rẻ đi nhanh chóng, các nhà khoa học đã liên tục tích lũy các trình tự gen của virus, thực vật và con người với chi phí gần như miễn phí. Kết quả là hàng trăm triệu trình tự protein đã được tích lũy trong cơ sở dữ liệu trình tự UniProt. Tuy nhiên, việc làm sáng tỏ xem các trình tự đó thực sự có hình dạng không gian 3D như thế nào lại diễn ra với một tốc độ hoàn toàn khác. Một nghiên cứu sinh phải mất bốn đến năm năm mới giải mã được cấu trúc của một protein bằng tinh thể học tia X hoặc kính hiển vi điện tử nghiệm lạnh. Thậm chí nếu không tạo được tinh thể, nghiên cứu sẽ bị đình trệ hoàn toàn. Dù đã tích lũy suốt 50 năm như vậy, số lượng cấu trúc mà nhân loại nắm giữ cũng không vượt quá 180.000 cấu trúc. Trình tự gen đã gia tăng nhanh gấp ba nghìn lần so với cấu trúc.

Đó chính là con số mà nhóm AlphaFold do Tiến sĩ John Jumper dẫn dắt phải đối mặt. Như đã thấy ở phần trước, thay vì động lực học phân tử vốn tính toán bằng cách di chuyển từng nguyên tử theo dòng thời gian, AlphaFold 2 đã giải quyết toàn bộ mối quan hệ về khoảng cách và góc giữa các axit amin chỉ trong một phép tính mạng nơ-ron duy nhất. Nhờ đó, thời gian cần thiết để thu được cấu trúc của một protein đã giảm từ vài tháng xuống nhiều nhất là vài phút. DeepMind đã dồn tốc độ này vào các trung tâm dữ liệu TPU của Google Cloud, nạp toàn bộ các trình tự đã biết trên Trái Đất vào hệ thống. Vào tháng 7 năm 2021, 350.000 cấu trúc đầu tiên đã được công bố, và đến tháng 7 năm 2022, 200 triệu cấu trúc đã được giải phóng cùng một lúc. Tính đến năm 2024, con số tích lũy đã tăng lên 214 triệu, và 98% protein do con người tạo ra đều nằm trong số này. Nếu con người xác định từng cấu trúc trong số 200 triệu này bằng thực nghiệm, tính theo cuộc đời của một nghiên cứu sinh, sẽ phải mất từ 800 triệu đến 1 tỷ năm. AlphaFold đã chuyển đổi khoảng thời gian này thành các phép tính toán trên máy tính chỉ tiêu tốn tiền điện và hoàn thành tất cả trong vòng một năm.

Tuy nhiên, việc chỉ có những con số tuôn ra không có nghĩa là các nhà khoa học có thể lập tức tin tưởng và sử dụng ngay. Nếu dự đoán sai nhưng lại hiển thị như một hình ảnh đầy chắc chắn thì còn nguy hiểm hơn. Để giải quyết vấn đề này, Tiến sĩ Jumper đã cùng thiết kế hai chỉ số để mô hình tự báo cáo độ tự tin của chính mình. Một là pLDDT. Khi giá trị này vượt quá 90, điều đó có nghĩa là cấu trúc dự đoán tương đương với cấu trúc được xác định qua thực nghiệm; còn nếu giảm xuống dưới 50, điều đó có nghĩa là vùng đó vốn dĩ không có hình dạng cố định và dao động linh hoạt trong nước. Chỉ số còn lại là PAE. Đối với các protein dài được chia thành nhiều vùng, chỉ số này cho biết mức độ gắn kết chặt chẽ và chuyển động cùng nhau của một vùng so với vùng khác theo đơn vị angstrom. Thông qua chỉ số này,

người ta có thể phân biệt xem đó là mối nối lỏng lẻo dao động như khớp nối, hay gắn kết vững chắc thành một khối thống nhất. AlphaFold thể hiện độ tin cậy này bằng màu sắc ngay trên hình ảnh cấu trúc. Màu đỏ dùng để tô các vùng có độ tin cậy thấp chính là màu sắc bộc lộ sự chưa biết chứ không hề che giấu.

Cơ sở dữ liệu được tích lũy như vậy đã mang lại kết quả ngay cả bên ngoài phòng thí nghiệm. Phức hợp lỗ nhân, cửa ngõ ra vào nhân tế bào, là một cỗ máy hình bánh donut khổng lồ với khối lượng phân tử hơn 100 triệu dalton và được cấu tạo từ hơn 1.000 mảnh protein, đến mức một bức ảnh kính hiển vi điện tử nghiệm lạnh không thể thu trọn toàn bộ các mối nối. Các nhóm nghiên cứu từ năm quốc gia đã ghép từng mô hình mảnh vỡ do AlphaFold tạo ra như những khối Lego để hoàn thiện bản đồ toàn thể. Enzyme phân hủy nhựa PETase do một loại vi khuẩn hiếm phát hiện dưới biển tạo ra cũng tương tự. Công việc cải tiến trước đây mất nhiều năm theo cách tạo đột biến ngẫu nhiên rồi chờ đợi may rủi, nay các nhà nghiên cứu chỉ mất một ngày sau khi kiểm tra cấu trúc AlphaFold và thiết kế lại vị trí liên kết. Enzyme tạo ra theo cách đó đã phân hủy nhựa nhanh hơn hàng chục lần so với trạng thái tự nhiên. Bài báo viết về cơ sở dữ liệu này đã được trích dẫn hơn 30.000 lần. Hơn 3 triệu nhà nghiên cứu tại 190 quốc gia đang sử dụng cơ sở dữ liệu này miễn phí. Ngay cả sinh viên đại học ở các nước đang phát triển không có tiền mua thiết bị đắt tiền, chỉ cần gõ tên protein vào thanh tìm kiếm là có thể nhận được bản đồ nguyên tử.

Xu hướng này vẫn tiếp tục trong năm 2026. EMBL-EBI và Google DeepMind, hai đơn vị quản lý Cơ sở dữ liệu AlphaFold, đã tái thiết kế giao diện trong năm nay, đưa cả các biến thể protein xuất phát từ cùng một gen vào phạm vi cấu trúc, đồng thời bổ sung thêm hơn 80.000 dự đoán về các cặp protein liên kết làm việc cùng nhau. Isomorphic Labs, mà chúng ta đã thấy ở Chương 1, đang chuẩn bị tiến hành thử nghiệm lâm sàng đầu tiên trên người trong năm nay với các ứng viên thuốc mới được thiết kế dựa trên AlphaFold 3. Tọa độ nguyên tử trên màn hình máy tính đã tiến sát đến thời điểm đi vào cơ thể người bệnh thực tế. Năng lực tính toán từng phá vỡ bức tường mười năm của phòng thí nghiệm Oxford chỉ trong ba mươi phút ấy, giờ đây đang bước vào vị trí sàng lọc các ứng viên thuốc mới

và đưa chúng vào cơ thể con người.

Người trẻ tuổi nhất sau 70 năm

Sáng ngày 9 tháng 10 năm 2024, một sự tĩnh lặng kỳ lạ bao trùm căn phòng ngủ ở London. Tiến sĩ John Jumper tắt báo thức rồi nằm xuống lại. Hôm đó là ngày công bố giải Nobel Hóa học, và ông là một trong những cái tên được nhắc đến như ứng viên sáng giá. Thế nhưng, ông tự tính toán xác suất đoạt giải của mình dưới 10%. Ông nghĩ rằng thà ngủ tiếp còn hơn phải trải qua 90% thất vọng khi kỳ vọng sụp đổ. Đến quá 10 giờ 30 phút sáng, điện thoại vẫn không reo. Ông nói với vợ: "Có lẽ năm nay chưa đến lượt chúng ta rồi". Vợ ông, Carolyn, nắm lấy tay ông và đáp: "Hãy chờ thêm một chút nữa xem sao". Lời nói vừa dứt, màn hình điện thoại ở đầu giường hiện lên thông báo cuộc gọi đến từ 'Thụy Điển'. Giọng nói bên kia đầu dây tuyên bố ông là người đồng đoạt giải Nobel Hóa học năm 2024. Năm đó ông 39 tuổi. Vào ngày đứng cạnh David Baker và Demis Hassabis để nhận giải, trên sân khấu lễ trao giải rất khó tìm thấy một gương mặt nào trẻ hơn ông.

Lý do ông nhận giải thưởng vô cùng rõ ràng. Bởi lẽ AlphaFold 2 do ông dẫn dắt đã giải quyết thực chất bài toán cuộn gấp protein tồn tại suốt 50 năm. Cuộn gấp protein là quá trình một chuỗi axit amin giống như sợi chỉ dài tự gấp lại để tạo thành một hình dạng ba chiều. Chỉ khi hình dạng này xuất hiện chuẩn xác thì chúng ta mới biết protein đó thực hiện chức năng gì trong cơ thể. Năm 1969, nhà khoa học người Mỹ Cyrus Levinthal đã tính toán rằng số lượng cấu hình khả dĩ mà một chuỗi axit amin có thể gấp lại nhiều đến mức dù có dùng thời gian dài hơn tuổi thọ vũ trụ cũng không thể đếm hết. Thế nhưng trong tế bào thực tế, protein chỉ mất vài giây để gấp thành hình dạng chính xác. Khoảng cách giữa hai điều này

được gọi là nghịch lý Levinthal, và các nhà khoa học đã nỗ lực giải quyết nghịch lý này chỉ bằng thực nghiệm suốt nửa thế kỷ. AlphaFold 2, tại cuộc thi CASP14 đã đề cập trước đó, đã giải quyết bài toán này trong máy tính chỉ sau vài phút với sai số nhỏ hơn bề rộng của một nguyên tử. Khi nghịch lý tồn tại nửa thế kỷ bị đảo ngược chỉ sau một ngày công bố tại cuộc thi, người ta kể lại rằng ngay cả các học giả tham gia ban đầu cũng không thể tin vào những con số ấy.

Ủy ban Nobel đánh giá thành tựu này là sự kiện phá vỡ bức tường ngăn cách giữa tính toán máy tính và khoa học sự sống. Như đã thấy ở phần đầu chương, cảm thấy bế tắc trước sự chờ đợi kéo dài cả một thế hệ của phương pháp thực nghiệm, ông đã chuyển từ vật lý sang hóa học và giải quyết bài toán sinh học nan giải lâu đời của nhân loại bằng tính toán. Do đó, khoảnh khắc Ủy ban Nobel xướng tên ông là người đoạt giải Hóa học cũng chính là khoảnh khắc vật lý và khoa học sự sống hội ngộ tại một điểm. Quá trình hoạt động của ông cho thấy một sự thật rằng: quan trọng không phải là công cụ bạn đã học, mà là bạn hướng công cụ đó vào giải quyết bài toán nào.

Đến Anthropic

Mùa hè năm 2026, Financial Times đưa tin rằng tổ chức chuyên trách phát triển AlphaFold đã dần bị giải tán trong suốt một năm qua. Các tác giả của bài báo gốc đã phân tán sang các tổ chức liên quan đến Gemini và Isomorphic Labs, đồng thời có thông tin tiếp nối rằng bản thân Jumper cũng đã chuyển sang một bộ phận nội bộ liên quan đến lập trình trước khi rời công ty. Sự việc này diễn ra không lâu sau khi ông nhận giải Nobel Hóa học năm 2024. Đôi bàn tay từng giải quyết bài toán sinh học nan giải 50 năm lại được bố trí vào vị trí hoàn thiện các công cụ lập trình. Tiến sĩ Jumper từ trước đến nay luôn nói rằng trí tuệ nhân tạo phải là công cụ chữa lành bệnh tật và làm sáng tỏ các quy luật vật lý của vũ trụ. Vị trí tiếp theo mà tài năng đó hướng tới đã sớm được hé lộ.

Từ năm 2015, Google DeepMind đã tự ví mình như Bell Labs của thế kỷ 20. Điều đó có nghĩa là một tổ chức cung cấp tài nguyên điện toán và thời gian cho các nhà khoa học xuất sắc mà không chịu áp lực thương mại. Tuy nhiên, vào tháng 11 năm 2022, khi OpenAI tung ChatGPT ra thế giới, tình thế đã thay đổi. Giám đốc điều hành Sundar Pichai lập tức ban bố tình trạng khẩn cấp Code Red, và đến năm 2023, DeepMind cùng Google Brain đã được sáp nhập thành một tổ chức duy nhất. Tổ chức từng nghiên cứu khoa học cơ bản đã chuyển thành tổ chức tập trung toàn lực vào phát triển sản phẩm Gemini.

Ngày 19 tháng 6 năm 2026, Tiến sĩ Jumper đã đăng một bài viết ngắn trên mạng xã hội cá nhân. Nội dung thông báo ông rời Google DeepMind sau 9 năm gắn bó để chuyển sang Anthropic. Một ngày trước đó, vào ngày 18 tháng 6, Noam Shazeer, người thiết kế Transformer, cũng tuyên bố rời Google để chuyển sang OpenAI. Chỉ trong hai ngày, Google đã mất đi cả bộ não lẫn trái tim của trí tuệ nhân tạo. Transformer hiện là khung sườn cho cách mọi mô hình trí tuệ nhân tạo đọc và viết văn bản, còn AlphaFold là đôi mắt đọc ra các mảnh ghép sự sống mang tên protein. Những người thiết kế hai công trình ấy đã lần lượt rời bỏ Google trong cùng một tuần. Ngày 22 tháng 6, giá cổ phiếu Alphabet có lúc giảm tới 7% trong phiên và đóng cửa với mức giảm khoảng 5%, khiến vốn hóa thị trường bốc hơi khoảng 225 tỷ USD. Con số này tương đương gần 300 nghìn tỷ won. Chỉ có hai người chuyển công ty, nhưng thị trường đã định giá sức nặng của sự kiện đó ngang bằng ngân sách của một quốc gia.

Ở đây có một điểm đáng chú ý. Alphabet, công ty mẹ của Google, là nhà đầu tư lớn nhất nắm giữ khoảng 14% cổ phần của Anthropic. Tuy nhiên, họ không có quyền biểu quyết và ghế trong hội đồng quản trị. Có thể nói, công ty do chính họ rót tiền đã đưa người đoạt giải Nobel của họ đi mất.

Ngày 30 tháng 6, Anthropic đã ra mắt Claude Science, một nền tảng dành riêng cho các nhà nghiên cứu. Cấu trúc của nền tảng này, nơi giúp cánh tay robot xác nhận ngay lập tức trong ống nghiệm những giả thuyết được xây dựng trong máy tính, sẽ được tìm hiểu chi tiết ở tiểu mục tiếp theo.

Năm 2007, Chủ tịch Intel Andy Grove từng đặt câu hỏi tại sao hiệu năng bán dẫn cứ 2 năm lại tăng gấp đôi, trong khi phát triển thuốc mới lại mất tới 10 năm. Nhà hóa dược Derek Lowe đã phản bác câu hỏi này là một sai lầm, cho rằng tế bào là sản phẩm được tạo nên từ 4 tỷ năm tiến hóa nên không thể thiết kế như chất bán dẫn. Tiến sĩ Jumper đang tìm cách lật ngược sự phản bác lâu đời này. Dario Amodei, Giám đốc điều hành của Anthropic, luôn khẳng định rằng trí tuệ nhân tạo có thể rút ngắn tiến bộ sinh học vốn mất từ 50 đến 100 năm của con người xuống còn 5 đến 10 năm. Cùng với việc ra mắt Claude Science, Anthropic cũng tuyên bố khởi động chương trình phát triển thuốc riêng nhắm vào các bệnh bị lãng quên và bệnh di truyền hiếm gặp mà các hãng dược lớn từng bỏ qua vì lợi nhuận thấp.

Hiện tại, vào tháng 8 năm 2026, Claude Science đang hoạt động dưới dạng dịch vụ beta. Anthropic thông báo sẽ hỗ trợ tài nguyên điện toán lên tới 30.000 USD cho mỗi dự án trong tối đa 50 dự án nghiên cứu để tiến hành thử nghiệm từ ngày 1 tháng 9 đến ngày 1 tháng 12. Bản thân Tiến sĩ Jumper vẫn chưa tiết lộ chức danh mới hay ngày bắt đầu chính thức, mà chỉ bày tỏ ý định tạm nghỉ ngơi lấy lại năng lượng trước. Hai tháng sau khi nhà khoa học đoạt giải Nobel từng làm việc với protein thay vì bán dẫn chuyển công ty, giờ đây ông đang chuẩn bị bắt đầu một cuộc thử nghiệm mới đối đầu với sự sống và bệnh tật.

Trí tuệ nhân tạo tiến đến bàn thí nghiệm

Ba giờ sáng, đèn trong một phòng thí nghiệm ở San Francisco đã tắt. Không có một bóng người. Thế nhưng trong lồng ấp ở góc phòng, các tế bào vẫn đang phát triển. Vài giờ trước, một tác tử trí tuệ nhân tạo mang tên Claude đã tự mình thiết kế cấu trúc kháng thể khớp chính xác với khe hở thụ thể đột biến trên bề mặt tế bào ung thư. Sau đó, nó chuyển bản thiết kế đó thành mã lệnh và truyền trực tiếp đến chemostat, một lò phản ứng hóa học của Coefficient Bio. Cổ máy suốt đêm điều chỉnh nhiệt độ và nồng độ oxy của dịch nuôi cấy, đồng thời tích lũy dữ liệu về cách tế bào phản ứng. Khi nhà nghiên cứu đi làm vào buổi sáng, thí nghiệm đã hoàn tất. Đó là khoảnh khắc trí tuệ nhân tạo từ việc đọc và tóm tắt bài báo chuyển sang tự tay vận hành bàn thí nghiệm. Ý nghĩa của không gian phòng thí nghiệm tự nó đã thay đổi.

Claude Science được đề cập ở tiểu mục trước chính là thành quả thương mại hóa chính thức theo định hướng này. Cấu trúc của nó gồm ba tầng lớp. Tầng trên cùng là mạng lưới suy luận nơi nhiều trí tuệ nhân tạo cùng hợp tác làm việc. Tác tử đảm nhận vai trò điều phối sẽ chia nhỏ các nhiệm vụ phức tạp rồi phân chia cho các tác tử cấp dưới, trong khi tác tử đảm nhận vai trò kiểm duyệt sẽ rà soát lại từng quy trình tính toán và nguồn trích dẫn. Tầng giữa là mạng lưới thu thập dữ liệu tất định. Đây là kênh truy xuất không sai sót từ 60 cơ sở dữ liệu sinh học công cộng, bao gồm NCBI và UniProt. Tầng dưới cùng là các mô hình nền tảng đã tiếp thu các quy luật vật lý và hình học protein, như Evo 2, Boltz-2, OpenFold3 được kết nối với NVIDIA BioNeMo. Tầng dưới cùng này được kết nối trực tiếp với bàn thí nghiệm robot của Coefficient Bio. Đó là một cấu trúc không chỉ dừng lại ở việc diễn giải dữ liệu mà còn đưa ra các mệnh lệnh điều khiển thiết bị thí nghiệm.

Các chỉ số độ tin cậy pLDDT và PAE do John Jumper thiết lập khi phát triển AlphaFold 2 cũng được tích hợp nguyên vẹn vào cấu trúc này. Cả hai chỉ số đều được tác tử sử dụng làm thước đo để tự phân biệt giữa sự chắc chắn và nghi ngờ. Để robot tiếp quản thí nghiệm, hệ thống phải vận hành mà không cần con người túc trực kiểm tra bên cạnh. Vì vậy, bảng điểm này thay thế đôi mắt con người để giữ vững tính chuẩn xác trong phán đoán của máy móc.

Một bài kiểm chuẩn đã chứng minh chính xác lý do tại sao kênh tất định này lại cần thiết. Nhóm nghiên cứu sinh học của Anthropic đã tạo ra VirBench, một bài kiểm chuẩn thử nghiệm trí tuệ nhân tạo trên 120 dữ liệu virus thực tế. Khi lặp lại ba lần cùng một câu hỏi yêu cầu tìm kiếm trình tự gen glycoprotein của virus Ebola, trí tuệ nhân tạo tự suy luận mà không có công cụ hỗ trợ đã đưa ra 106 kết quả ở lần thứ nhất, lần thứ

hai là 15 kết quả, và lần thứ ba là 5 kết quả. Đáp án chính xác là 266. Độ chính xác trung bình giảm xuống còn 16,9%. Thế nhưng chỉ cần gắn thêm một công cụ tắt định mang tên gget virus, độ chính xác của chính mô hình đó đã nhảy vọt lên 92,8%. GPT-5.5 khi gắn cùng công cụ này đã đạt mức 99,7%. Điều đó có nghĩa là việc tích hợp chuẩn xác một công cụ hiệu quả hơn nhiều so với việc chỉ tăng kích thước mô hình.

Cùng một nhóm nghiên cứu cũng đã tiến hành hai thử nghiệm tự động biên soạn nội dung. WikiCrow đã hoàn thành các mục từ bách khoa toàn thư cho 15.600 gen chưa được đưa đầy đủ lên Wikipedia trong tổng số 20.000 gen người chỉ trong 48 giờ. Đây là kết quả của tác tử trí tuệ nhân tạo PaperQA2 khi rà soát 80 triệu bài báo khoa học và thực hiện hơn 100.000 lượt tìm kiếm. Khi các nhà nghiên cứu cấp tiến sĩ từ Harvard, Viện Công nghệ Massachusetts và Viện Francis Crick chấm điểm theo phương pháp mù, độ chính xác trích dẫn của WikiCrow đạt 86,1%, vượt trội so với mức 71,2% của các bài viết Wikipedia do con người soạn thảo. Tỷ lệ tự thêm bớt các câu văn không có căn cứ cũng chỉ ở mức 13,5%, thấp hơn mức 24,9% của con người. Ngược lại, ContraCrow là tác tử tìm kiếm các tuyên bố mâu thuẫn lẫn nhau giữa các bài báo. Khi khảo sát 93 bài báo được chọn ngẫu nhiên, trung bình mỗi bài phát hiện ra 2,34 mâu thuẫn tiềm ẩn. Tỷ lệ trùng khớp với phán đoán của các chuyên gia vượt quá 70%, và điểm F1 đạt 0,82.

Hướng đi mà tất cả các thiết bị này chỉ ra là một. Khi Claude Science kết hợp các cơ sở dữ liệu này với các quy luật vật lý và nắm trong tay bản thí nghiệm robot của Coefficient Bio, tế bào đang trở thành đối tượng được thiết kế và kiểm chứng như chất bán dẫn, đúng như câu hỏi của Grove ở tiểu mục trước. Vào nửa đầu năm 2026, Anthropic đã mua lại Coefficient Bio, công ty khi đó chỉ có tám người, với giá khoảng 400 triệu USD. Đây là nơi các nhà sinh học tính toán xuất thân từ Genentech từng thực hiện việc thiết kế chu trình phát triển thuốc và phát hiện đích tác dụng.

Trong buổi công bố ngày 30 tháng 6 năm 2026, Anthropic đã trình diễn quy trình trí tuệ nhân tạo tìm kiếm các ứng viên điều trị bệnh di truyền hiếm gặp, và các công ty được phẩm cùng viện nghiên cứu trải nghiệm sớm đã chia sẻ rằng thời gian dành cho việc phân tích tài liệu và lựa chọn đích tác dụng đã giảm đi đáng kể.

Nghiên cứu trí tuệ nhân tạo giờ đây không còn dừng lại bên trong màn hình máy tính. Một dòng mã điều khiển bộ điều nhiệt của lồng ấp, và kết quả thu được lại tiếp tục dẫn tới dòng mã tiếp theo. Đôi bàn tay của John Jumper từng tính toán tọa độ của từng nguyên tử, nay đang kết nối thành những tín hiệu đóng mở van của lồng nuôi cấy tế bào.

Những bài toán vẫn chưa có lời giải

Tương xứng với quy mô của những lời tán dương dành cho AlphaFold, tiếng nói từ giới học thuật chỉ ra các hạn chế của nó cũng rất rõ ràng. Những gì mô hình này đưa ra là một hình ảnh tĩnh duy nhất của một protein. Tuy nhiên, protein trong thực tế liên tục thay đổi hình dạng khi hoạt động bên trong tế bào. Chúng gắn vào rồi tách ra khỏi các phân tử khác, và khi nhận tín hiệu, chúng mở ra rồi gấp lại các vùng cấu trúc. Ruth Nussinov, nhà sinh học cấu trúc tại Viện Ung thư Quốc gia Hoa Kỳ, đã chỉ ra rằng AlphaFold chỉ nắm bắt được một trạng thái duy nhất chứ không phải tập hợp các biến đổi hình dạng này, và những chức năng phát sinh từ chuyển động như điều hòa dị lập thể của protein không thể giải thích bằng một bức tranh tĩnh. Việc dự đoán cấu trúc và việc thấu hiểu cách cấu trúc đó vận động sống động là hai vấn đề hoàn toàn khác nhau.

Bức tường thứ hai lộ diện ở những vùng dữ liệu còn trống. Các protein vô định hình nội tại mềm dẻo không có hình dạng cố định nên không kết tinh thành tinh thể, khiến dữ liệu thực nghiệm vốn dĩ đã hiếm hoi; và nhiều bài báo tổng quan đã chỉ ra rằng AlphaFold cũng chỉ trả về độ tin cậy thấp hoặc dễ đoán sai ở những vùng như vậy. Đối với các protein chỉ định hình hoàn chỉnh khi liên kết với phối tử, ion kim loại hoặc các chuỗi khác, những dự đoán tĩnh khi tính toán đơn lẻ cũng rất dễ bỏ sót. Do đó, thái

độ tiếp nhận cấu trúc dự đoán chưa qua kiểm chứng thực nghiệm như một sự thật hiển nhiên là rất nguy hiểm. Các tọa độ nguyên tử trên màn hình chỉ là những giả thuyết hợp lý, và luôn đi kèm lời cảnh báo rằng chúng có thể sai sót cho đến khi được chứng minh bằng tia X, kính hiển vi điện tử nghiệm lạnh hoặc các thí nghiệm hóa sinh.

Vì vậy, các ý kiến phê bình cho rằng dự đoán cấu trúc không đồng nghĩa với việc thấu hiểu chức năng và cơ chế hoạt động vẫn tiếp tục được đưa ra. Việc biết protein trông như thế nào chỉ là điểm khởi đầu chứ không phải đích đến của việc hiểu protein đó làm gì và tại sao lại làm như vậy. Trọng tâm của lời phê bình này là bản đồ do AlphaFold vẽ ra không phải là công cụ xóa bỏ chuyển khảo sát thực nghiệm, mà gần giống như một chiếc la bàn chỉ hướng nơi cần khảo sát. Thành tựu vượt qua ngưỡng cửa bài toán nan giải 50 năm và những lĩnh vực về động lực học, tương tác cũng như sự vô định hình mà thành tựu đó chưa thể chạm tới, đang cùng tồn tại song hành như thế.

Nội dung trong chương này được trình bày dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Chương 3. Regina Barzilay: Con mắt tìm ra Halicin

Từ ngôn ngữ đến phân tử

Vào một đêm năm 2014, Giáo sư Regina Barzilay ngồi trên giường bệnh và bật màn hình máy tính xách tay. Trên màn hình hiện lên hình ảnh chụp nhũ ảnh của chính bà từ hai năm trước. Vài giờ trước đó, bà nhận được chẩn đoán mắc bệnh ung thư vú giai đoạn 3. Tế bào ung thư đã lan đến các hạch bạch huyết ở nách. Với bàn tay run rẩy, bà phóng to bức ảnh chụp hai năm trước. Trong góc của những điểm ảnh đen trắng, một vết mờ hiện ra. Khi ấy, bác sĩ điều trị đã xem bức ảnh này và kết luận là bình thường. Hai năm sau, vết mờ đó đã phát triển thành khối ung thư trong cơ thể bà.

Có lý do khiến cảnh tượng này có vẻ xa lạ. Giáo sư Barzilay vốn dĩ là người cách rất xa bệnh viện. Cuối thập niên 1990, bà theo học chương trình thạc sĩ tại Viện Công nghệ Technion ở Israel. Thời kỳ đó, lĩnh vực xử lý ngôn ngữ tự nhiên đang chuyển dịch từ việc nhập tay từng quy tắc ngữ pháp sang học ngôn ngữ bằng thống kê và xác suất. Giáo sư Barzilay đã đứng ở hàng đầu của sự chuyển đổi này. Sau đó, bà nhận bằng tiến sĩ khoa học máy tính tại Đại học Columbia và trở thành giáo sư tại Phòng thí nghiệm Khoa học Máy tính và Trí tuệ Nhân tạo (CSAIL) thuộc Viện Công nghệ Massachusetts (MIT). Bà đã tạo ra trí tuệ nhân tạo tóm tắt văn bản và trí tuệ nhân tạo dịch thuật qua lại giữa nhiều ngôn ngữ. Rồi bà đạt được thành tựu khiến giới học thuật kinh ngạc: giải mã các ký tự cổ đại đã biến mất bằng trí tuệ nhân tạo mà không cần từ điển hay sách dịch. Chữ hình nêm và chữ Linear B, những văn tự mà suốt hàng nghìn năm không ai đọc được. Nhóm của Giáo sư Barzilay đã dùng tính toán để lần lại cấu trúc của các ký tự này và khôi phục ý nghĩa của chúng. Việc hồi sinh một ngôn ngữ đã chết tự thân nó là một công việc đầy chất thơ.

Thế nhưng vào năm 2014, tại bệnh viện nơi bà đến khám sức khỏe định kỳ, cuộc đời bà bỗng dừng lại. Bà rời khỏi thư phòng mà mình đã dày công xây dựng cả đời để bước vào phòng bệnh truyền hóa chất và phòng xạ trị. Trong suốt quá trình điều trị, điều khiến bà dằn vặt nhất không phải là bản thân căn bệnh ung thư, mà là hồ sơ khám bệnh cũ của chính mình. Khi đó, các bác sĩ chỉ dựa vào cách ước lượng mật độ mô vú bằng mắt thường để đọc kết quả, và phương pháp này đã không thể nắm bắt được những tín hiệu vi mô. Hai năm lẽ ra đã rất khác nếu được phát hiện sớm cứ thế trôi qua mất.

Ở khu điều trị nội trú, bà phải chứng kiến từng bệnh nhân cùng tăng lần lượt qua đời mỗi ngày. Rồi một ngày nọ, bà nghĩ về giới học thuật nơi mình từng gắn bó. Hàng năm có hàng vạn bài báo khoa học được công bố, và các nhà nghiên cứu reo hò khi độ chính xác của dịch máy tăng thêm 0,5 phần trăm. Thế nhưng, phòng đọc ảnh y tế nơi sinh mạng con người bị đe dọa lại vẫn dựa vào mắt nhìn và trực giác của con người

như trước. Bà nằm trên giường bệnh và cảm nhận khoảng cách này bằng cả cơ thể mình. Rồi ngay trong phòng bệnh, bà đã hạ quyết tâm: một khi bình phục, bà sẽ gác lại nghiên cứu xử lý ngôn ngữ tự nhiên và dành trọn phần đời còn lại cho trí tuệ nhân tạo trong lĩnh vực y tế và phát triển thuốc mới.

Năm 2015, sau khi chiến thắng bệnh tật và trở lại MIT, Giáo sư Barzilay đã biến quyết tâm này thành hành động. Bà đảm nhận vị trí người phụ trách đầu tiên của Jameel Clinic, phòng thí nghiệm trí tuệ nhân tạo lâm sàng mới được thành lập tại MIT. Ngay sau đó, bà đã cho ra đời Mirai giúp nhận biết sớm ung thư vú và Sybil giúp định vị trước vị trí có nguy cơ phát triển ung thư phổi. Cả hai hệ thống trí tuệ nhân tạo này đều là kết quả sinh ra từ nỗi đau chẩn đoán sai mà chính bà đã trải qua.

Bước chân của Giáo sư Barzilay không dừng lại ở đó. Cùng với Giáo sư Jim Collins và Tiến sĩ Jonathan Stokes, bà đã dẫn thân vào việc phát triển trí tuệ nhân tạo tìm kiếm thuốc mới, và tìm ra Halicin – một phân tử lạ lẫm hầu như không trùng lặp cấu trúc với các loại kháng sinh hiện có. Nhờ thành tựu này, vào năm 2020, bà đã trở thành người đầu tiên nhận giải thưởng Squirrel do AAAI mới lập ra nhằm trao

cho các học giả có đóng góp cho phúc lợi nhân loại, nhận về giải thưởng trị giá 1 triệu đô la.

Bước chân của Giáo sư Barzilay vẫn đang tiếp tục cho đến hôm nay. Một bài báo trên tạp chí Nature vào tháng 6 năm 2026 cho biết lý do bà dẫn thân vào nghiên cứu kháng sinh không chỉ vì cuộc chiến chống ung thư của bản thân. Cha bà cũng từng bị nhiễm trùng cột sống và chịu nhiều đau đớn vì không tìm được loại kháng sinh phù hợp. Cũng trong bài báo đó, nhóm nghiên cứu của Tiến sĩ Jonathan Stokes đã sàng lọc 10.000 hợp chất để tìm ra một phân tử kháng khuẩn mới tên là Enterolin, và phòng thí nghiệm của Giáo sư Barzilay đã công bố công cụ mang tên DiffDock giúp dự đoán cách các phân tử nhỏ liên kết với protein. Cùng tháng đó, Boston Globe đã vinh danh bà trong danh sách 'Tech Power Players 50'. Điều này cho thấy vị học giả từng là nữ hoàng của lĩnh vực xử lý ngôn ngữ tự nhiên vẫn đang thực hiện lời hứa trên giường bệnh năm xưa, ngay cả khi hơn mười năm đã trôi qua.

Đẩy sớm thời điểm chẩn đoán

Qua khung cửa sổ phòng nghiên cứu của Giáo sư Barzilay, tòa nhà nghiên cứu và phát triển của Novartis hiện ra mỗi ngày. Đứng giữa khuôn viên Viện Công nghệ Massachusetts (MIT) nhìn sang tòa nhà của công ty dược phẩm đa quốc gia, bà không thể xua đi một suy nghĩ. Đó là công thức phân tử mà các nhà hóa học viết trên bảng đen, tức chuỗi ký tự SMILES, thực chất không khác gì một câu trong ngôn ngữ tự nhiên. Đó là sự giác ngộ rằng dữ liệu ngôn ngữ mà bà đã xử lý suốt gần 20 năm và dữ liệu phân tử mà các nhà hóa học làm việc cùng đều có chung một cấu trúc.

Cuộc gặp gỡ này bắt đầu từ một sự tình cờ. Khi Giáo sư Klavs Jensen và Giáo sư Connor Coley thuộc Khoa Kỹ thuật Hóa học của MIT chuẩn bị dự án dự đoán phản ứng tổng hợp ngược với sự tài trợ của DARPA, họ đã mời Giáo sư Barzilay với tư cách là một nhà khoa học máy tính chuyên xử lý cấu trúc đồ thị. Đó là khoảnh khắc mà vị chuyên gia hàng đầu về xử lý ngôn ngữ tự nhiên ngồi xuống trước các đồ thị phân tử. Con đường thúc đẩy tốc độ phát triển thuốc mới và kháng sinh lên gấp hàng trăm lần nằm chính tại nơi đó.

Trong lần hợp tác đầu tiên này, nhóm nghiên cứu đã cho ra mắt Chemprop – một mạng nơ-ron đọc phân tử dưới dạng đồ thị gồm các nguyên tử và liên kết thay vì chuỗi ký tự, cùng JT-VAE (Junction Tree Variational Autoencoder) – mô hình chỉ chọn lọc và tạo ra những phân tử có thể tồn tại về mặt hóa học. Cách thức hoạt động của các công cụ này sẽ được trình bày chi tiết trong phần tiếp theo.

Cùng thời điểm đó, thành quả mà Giáo sư Barzilay giới thiệu với thế giới ở mảng chẩn đoán là trí tuệ nhân tạo đọc trước nguy cơ ung thư vú và ung thư phổi. Cùng với Tiến sĩ Constance Lehman tại Bệnh viện Đa khoa Massachusetts, bà đã tạo ra Mirai. Hệ thống được huấn luyện trên hơn 128.000 bức ảnh chụp nhũ ảnh thu thập từ 7 bệnh viện thuộc 4 quốc gia gồm Mỹ, Thụy Điển và Đài Loan. Thay vì tiêu chuẩn mật độ tuyến vú mà các bác sĩ chẩn đoán hình ảnh đã dùng suốt nhiều thập kỷ, mô hình đọc các mẫu mô vi mô trong ảnh và dự đoán khả năng khởi phát ung thư vú trong vòng 5 năm với độ chính xác từ 75 đến 84 phần trăm. Kết quả này đã được công bố trên tạp chí học thuật Radiology. Hệ thống Sybil do bà phát triển cùng Tiến sĩ Lecia Sequist nhắm vào ung thư phổi. Sau khi được huấn luyện bằng dữ liệu Thử nghiệm Tầm soát Ung thư Phổi Quốc gia Hoa Kỳ (NLST), mô hình được áp dụng trực tiếp lên nhóm bệnh nhân Đài Loan hầu như không có tiền sử hút thuốc mà không cần huấn luyện lại. Mô hình đã dự đoán khả năng khởi phát ung thư phổi trong vòng 2 năm với độ chính xác AUC từ 0,80 đến 0,95, và trong vòng 6 năm với AUC từ 0,76 đến 0,80. Điều này có nghĩa là hệ thống đã đọc được trước dấu vết của khối u sẽ phát triển nhiều năm sau đó ngay trong những bức ảnh mà bác sĩ chẩn đoán hình ảnh từng kết luận là bình thường và cho bệnh nhân ra về.

Thành tựu giúp Giáo sư Barzilay được biết đến rộng rãi trên toàn thế giới là kháng sinh Halicin. Phương thức máy móc rà soát trước kho hợp chất ảo khổng lồ rồi con người kiểm chứng lại trên bàn thí nghiệm đã tìm ra loại kháng sinh có khung cấu trúc hoàn toàn mới như thế nào sẽ được tiếp tục trình

bày ở phần cuối của chương này.

Mirai không dừng lại ở mức bài báo trong phòng thí nghiệm. Tháng 4 năm nay, Hội đồng Đạo đức của Bệnh viện Santa Casa ở Porto Alegre, Brazil đã phê duyệt thử nghiệm lâm sàng áp dụng Mirai cho 500 bệnh nhân. Bệnh viện này có 741 giường bệnh trong khuôn khổ SUS – hệ thống y tế công cộng quy mô lớn nhất thế giới – và khám chữa bệnh cho 6 triệu lượt bệnh nhân mỗi năm. Sự hợp tác với MIT bắt đầu từ hai năm trước, sau khi trải qua nghiên cứu kiểm chứng hồi cứu được công bố trên tạp chí ung thư học Brazil, đã tiến sang thử nghiệm tiên cứu trên bệnh nhân thực tế. Tiến sĩ Carmela Farias da Silva, người dẫn dắt nghiên cứu, cho biết công nghệ này có thể tạo ra sự thay đổi thực sự khi bước vào thực tế của các bệnh viện địa phương.

Hơn mười năm kể từ khi đôi bàn tay từng đọc ngữ pháp trong văn bản chuyển sang đọc các quy luật vật lý ẩn giấu trong các bức ảnh phân tử và tế bào, đầu ngón tay ấy giờ đây đang vươn vào bên trong các phòng khám bệnh viện ở nửa kia bán cầu.

Học biểu diễn phân tử

Nhiều thập kỷ trước, trên bàn làm việc của các nhà hóa học luôn có một bảng kiểm. Khi một phân tử mới xuất hiện trong phòng thí nghiệm, nhà hóa học sẽ rà soát từng ô trong công thức cấu trúc của nó và đặt câu hỏi: Có vòng benzen không? Có nhóm amin gắn vào không? Có bao nhiêu nhóm methyl? Câu trả lời chỉ là 0 và 1. Có là 1, không có là 0. Tập hợp các số nhị phân được tạo ra như thế được gọi bằng cái tên 'dấu vân tay phân tử' (Molecular Fingerprints). Dấu vân tay của con người lưu giữ nguyên vẹn họa tiết độc nhất của riêng người đó. Thế nhưng, dấu vân tay phân tử chẳng qua chỉ là kết quả của một bảng khảo sát được điền câu trả lời theo bộ câu hỏi mà ai đó đã định sẵn từ trước.

Phương pháp bảng kiểm này bộc lộ những hạn chế rõ rệt. Tùy thuộc vào việc đặc tính cần dự đoán là độc tính, độ hòa tan trong nước hay lực liên kết với một thụ thể cụ thể, cấu trúc phụ nào đóng vai trò quyết định lại thay đổi mỗi lần. Danh sách câu hỏi cố định không thể dung nạp được sự khác biệt này. Nó cũng không thể thể hiện được gradient điện tích chạy giữa các nguyên tử, hay góc đối diện thực tế của các nguyên tử trong không gian 3 chiều. Nó giống như việc phán đoán một phân tử đứng ở dạng lập thể chỉ bằng cách nhìn vào cái bóng phẳng của nó.

Nhóm nghiên cứu tại Viện Công nghệ Massachusetts (MIT) do Giáo sư Regina Barzilay dẫn dắt đã quyết định gác lại toàn bộ bảng câu hỏi này. Họ thay đổi thiết kế để mạng nơ-ron tự đặt ra câu hỏi và tự tìm kiếm câu trả lời. Thành quả của việc đó là Chemprop. Khung xương của Chemprop là một thuật toán mang tên Mạng Nơ-ron Truyền Thông điệp (Message Passing Neural Network, MPNN). Nó chuyển đổi phân tử từ chuỗi ký tự SMILES thành một đồ thị $G(V, E)$ với các nguyên tử là các điểm và các liên kết là các đường thẳng. Từng nguyên tử trở thành các nút của đồ thị, và các liên kết hóa học trở thành các cạnh nối các nút đó.

Bên trong đồ thị này, các nguyên tử trao đổi mẩu giấy nhắn cho nhau. Một nguyên tử nhận thông tin từ các nguyên tử lân cận liên kết với nó và chỉnh sửa dần trạng thái của chính mình. Lặp lại quá trình này qua vài bước, tình trạng của không chỉ nguyên tử kế bên mà cả những nguyên tử cách đó hai, ba bước cũng thấm sâu vào trạng thái của nó. Công việc mà nhà hóa học từng làm trên bàn thí nghiệm – đối chiếu bằng mắt vòng benzen và nhóm carboxyl – giờ đây được mạng nơ-ron thay thế bằng việc trao đổi các vectơ số. Khi thu thập và nén toàn bộ thông tin mà tất cả các nguyên tử đã trao đổi lại với nhau, một phân tử được nén thành một vectơ duy nhất. Nhóm nghiên cứu gọi vectơ này là vectơ không gian ẩn. Không phải từ bảng câu hỏi do con người định sẵn, mà trong quá trình huấn luyện bằng việc quan sát dữ liệu,

các đặc trưng do chính mạng nơ-ron tự chắt lọc đã được chứa đựng bên trong vectơ đó.

Quá trình nén này cũng ẩn chứa cạm bẫy. Các mạng nơ-ron đồ thị thời kỳ đầu tạo ra một vectơ phân tử duy nhất bằng cách cộng dồn tất cả các vectơ nguyên tử lại với nhau. Tuy nhiên, ngay cả với hai phân tử có khung xương hoàn toàn khác nhau, nếu loại và số lượng nguyên tử tương tự nhau thì tổng giá trị cộng lại có thể trùng khớp một cách ngẫu nhiên. Điều này giống như việc ghi nhận hai người có khuôn mặt khác nhau là cùng một người chỉ vì họ có tổng trọng lượng bằng nhau. Để giải quyết vấn đề này, nhóm nghiên cứu của Barzilay đã đưa vào các mẫu nguyên mẫu Wasserstein (Wasserstein Prototypes) và kỹ thuật pooling chú ý đa chiều. Thay vì chỉ cộng đơn thuần, họ đã sửa đổi phương pháp tính toán để nếu hình dạng của phân tử có chút khác biệt, sự khác biệt vì mô đó cũng được khắc ghi khác đi trong vectơ. Phép tính này tương tự như việc tóm tắt một cuốn tiểu thuyết. Dù hai cuốn tiểu thuyết có cùng số lượng từ, nhưng nếu cốt truyện khác nhau thì ấn tượng để lại cho người đọc hoàn toàn khác nhau. Một bản tóm tắt tốt sẽ không xóa nhòa sự khác biệt đó mà giữ cho nó sống động. Việc mà bước pooling của Chemprop hướng tới cũng không hề khác biệt.

Khi đã nắm trong tay cách biểu diễn này, nhóm nghiên cứu tiến thêm một bước nữa. Wengong Jin, khi đó là nghiên cứu sinh sau tiến sĩ, đã chỉ ra điểm yếu của các mô hình sinh hiện có vốn tạo ra phân tử mới bằng cách nối từng nguyên tử lại với nhau. Khi gắn từng nguyên tử một, hơn chín phần mười kết quả tạo ra không thể tồn tại về mặt hóa học do vi phạm quy tắc hóa trị hoặc không thể đóng kín vòng. Mô hình Junction Tree Variational Autoencoder (JT-VAE) do Wengong Jin tạo ra đã thay đổi tư duy. Thay vì gắn từng nguyên tử riêng lẻ, mô hình chọn cách chuẩn bị sẵn các khối đã ổn định về mặt hóa học như vòng benzen hay nhóm carboxyl giống như các khối Lego, rồi lắp ráp các khối đó lại. Quy trình thực hiện là dựng cấu trúc cây ở cấp độ khối tương ứng với bức tranh tổng thể trước, sau đó mới điền tiếp việc sắp xếp các nguyên tử bên trong. Vì bản thân các khối đã tuân thủ các định luật hóa học, nên phân tử được tạo ra theo cách này ngay từ đầu đã không tạo ra những cấu trúc bất khả thi về mặt hóa học.

Cách biểu diễn phân tử được trau chuốt như vậy sau này đã trở thành nền tảng cho công việc rà soát toàn bộ thư viện chứa hàng trăm triệu hợp chất ảo. Câu chuyện về loại kháng sinh ra đời từ bên trong thư viện đó sẽ được tiếp tục ở phần sau.

Bước tiến của phòng thí nghiệm này vẫn không dừng lại trong năm 2026. Đài phát thanh công cộng WBUR của Hoa Kỳ đưa tin vào tháng 5 năm 2026 rằng, dù mô hình dự đoán nguy cơ ung thư vú Mirai do nhóm của Giáo sư Barzilay phát triển có thể đọc trước các tín hiệu nguy hiểm sớm hơn con người,

nhưng mô hình lại đang lan tỏa một cách chậm chạp vào thực tế bệnh viện. Bản tin chỉ ra rằng dù toán học chuyển đổi dữ liệu thành vectơ đã hoàn tất khâu kiểm chứng, nhưng các thể chế và tập quán của bệnh viện trong việc tiếp nhận kết quả đó vẫn chưa theo kịp. Vào tháng 6 cùng năm, Quý Abdul Latif Jameel đã cùng với MIT Jameel Clinic tổ chức một buổi hội thảo ăn trưa với chủ đề y học dự phòng, và hãng tin Ynetnews của Israel vào tháng 1 cũng đăng tải phát biểu của Giáo sư Barzilay rằng Israel có thể đi đầu trong việc hợp tác quốc tế về trí tuệ nhân tạo y tế.

Khoảnh khắc bảng câu hỏi được buông khỏi tay và việc biểu diễn được giao lại cho mạng nơ-ron, khoảng cách giữa hóa học và trí tuệ nhân tạo đã được thu hẹp một cách rõ rệt.

Phát hiện ra Halicin

Một đêm năm 2019, trên bàn thí nghiệm của Viện Broad gần Viện Công nghệ Massachusetts (MIT), chín mươi chín đĩa nuôi cấy nhỏ được đặt ngay ngắn. Tiến sĩ Jonathan Stokes nhỏ một giọt hợp chất vào từng đĩa đang nuôi cấy vi khuẩn E. coli. Một trong số đó là chất từng được phát triển làm ứng viên thuốc điều trị tiểu đường từ vài năm trước nhưng bị xếp xó trong ngăn kéo vì không có hiệu quả. Đó là hợp chất thậm chí không có tên gọi mà chỉ được gọi bằng mã SU-3327. Ngày hôm sau, khi kiểm tra đĩa nuôi cấy, Tiến sĩ Stokes không dám tin vào mắt mình. Vi khuẩn E. coli hoàn toàn không phát triển. Hợp chất này sau đó đã được đặt tên là Halicin. Đó là loại kháng sinh do máy tính chứ không phải con người chọn

ra.

Người dẫn dắt thí nghiệm này là Giáo sư Regina Barzilay thuộc MIT, một học giả về xử lý ngôn ngữ tự nhiên. Kể từ penicillin vào thập niên 1940, thuốc kháng sinh đã phát triển bằng cách khai thác từng chất một từ vi sinh vật trong đất. Tuy nhiên, từ cuối thập niên 1960, các chất có cùng cấu trúc cứ lặp đi lặp lại. Các công ty dược phẩm chuyển sang phương thức dùng máy móc để thử nghiệm từng chất trong số hàng triệu hợp chất, nhưng chi phí và thời gian lại quá lớn. Suốt 30 năm, hầu như không có loại kháng sinh nào mang cấu trúc mới xuất hiện. Hơn nữa, kháng sinh là loại thuốc chỉ uống vài ngày là khỏi bệnh nên không mang lại nhiều lợi nhuận cho các hãng dược. Trong khi các công ty dược phẩm lớn từ bỏ nghiên cứu kháng sinh, vi khuẩn đa kháng thuốc lại gia tăng nhanh chóng. Đặt máy móc vào vị trí mà con người đã buông tay chính là lựa chọn của Giáo sư Barzilay.

Thật khó để con người có thể đoán biết hiệu quả kháng khuẩn chỉ bằng cách nhìn vào công thức hóa học. Thế nhưng Chemprop, mạng nơ-ron do nhóm của Giáo sư Barzilay xây dựng, đã chỉ ra được hiệu quả của những phân tử chưa từng thấy bao giờ dù chỉ mới học từ 2.335 hợp chất đã được xác nhận khả năng ức chế vi khuẩn E. coli. Nếu theo cách cũ, con người phải mất nhiều năm để tổng hợp từng hợp chất trong phòng thí nghiệm rồi thử nghiệm lên vi khuẩn E. coli, thì mạng nơ-ron này đã hoàn thành cùng một phán đoán đó bên trong máy tính chỉ trong vài giờ.

Tiến sĩ Stokes đã đưa 6.000 hợp chất được lưu trữ tại Viện Broad vào mạng nơ-ron này. Khi tiến hành thử nghiệm thực tế trên 99 hợp chất hàng đầu do mạng nơ-ron lựa chọn, có tới 51 hợp chất tiêu diệt được vi khuẩn E. coli. Tỷ lệ này là hơn một nửa. Tỷ lệ thành công theo phương pháp truyền thống thường dưới 0,1 phần trăm. Giáo sư Barzilay đã tiến thêm một bước nữa từ điểm này. Bà lọc bỏ tất cả các phân tử có cấu trúc dù chỉ hơi giống với các loại kháng sinh hiện có. Vứt bỏ những phân tử quen thuộc và chỉ giữ lại những phân tử lạ lẫm. Chất còn lại sau quá trình đó chính là Halicin. Chỉ số thể hiện mức độ tương đồng với các kháng sinh hiện có chỉ dừng lại ở mức 0,18. Đó là một bộ khung cấu trúc hoàn toàn mới. Nhóm nghiên cứu đã lấy tên máy tính trí tuệ nhân tạo HAL trong phim để

đặt tên cho chất này là Halicin. Họ đã khắc ghi ngay trong tên gọi sự thật rằng đây là loại kháng sinh đầu tiên do máy móc chứ không phải con người chọn ra. Mạng nơ-ron cũng chỉ ra phần nào của phân tử tạo ra hoạt tính kháng khuẩn. Điều này giúp các nhà hóa học có thể tận mắt xác nhận lý do vì sao họ nên tin tưởng vào vật chất này.

Thử thách thực sự nằm ở bước tiếp theo. Nhóm nghiên cứu đã để vi khuẩn E. coli tiếp xúc liên tục với Halicin trong suốt 30 ngày. Mục đích là để xem liệu vi khuẩn có thích nghi với thuốc và phát triển khả năng kháng thuốc hay không. Ciprofloxacin, loại thuốc dùng làm nhóm đối chứng, đã xuất hiện vi khuẩn kháng thuốc chỉ sau một ngày, và đến ngày thứ 30, lượng thuốc cần thiết đã tăng lên gấp 250 lần. Ngược lại, với Halicin, không có một vi khuẩn kháng thuốc nào xuất hiện trong suốt 30 ngày. Lý do chỉ được sáng tỏ sau khi nhóm nghiên cứu quan sát điện thế trên màng vi khuẩn. Khi đo điện thế màng bằng thuốc nhuộm đặc biệt, Halicin đã phá vỡ ngay lập tức điện thế màng giúp vi khuẩn tạo ra năng lượng, hay còn gọi là lực đẩy proton. Giống như việc máy phát điện bị tắt ngúm trước khi vi khuẩn kịp có cơ hội đột biến.

Halicin cũng thể hiện sức mạnh của mình khi bước ra khỏi phòng thí nghiệm. Nó đã tiêu diệt ba mươi sáu chủng vi khuẩn đa kháng thuốc nguy hiểm do Tổ chức Y tế Thế giới chỉ định chỉ trong vòng một ngày. Khi thoa thuốc mỡ Halicin chỉ một lần duy nhất lên vết thương bị nhiễm vi khuẩn đa kháng thuốc trên lưng chuột, vi khuẩn đã biến mất không dấu vết chỉ sau một ngày. Khi cho chuột bị viêm ruột uống thuốc qua đường miệng, toàn bộ vi khuẩn xâm nhập sâu vào trong ruột đều biến mất sau bốn ngày. Tuy nhiên, nhóm nghiên cứu cũng không giấu giếm mà ghi rõ sự thật rằng thuốc không có tác dụng đối với trực khuẩn mủ xanh (*Pseudomonas aeruginosa*) có màng ngoài dày. Đó là chi tiết thể hiện sự trung thực khi ghi lại cả những loại vi khuẩn thất bại thay vì chỉ chọn lọc khoe khoang những kết quả thành công.

Nhóm nghiên cứu không dừng lại ở đó mà tiếp tục rà soát lại 107 triệu hợp chất đáp ứng điều kiện kháng sinh từ cơ sở dữ liệu ZINC15 chứa 1,5 tỷ hợp chất. Trong số 23 ứng viên cuối cùng, 8 hợp chất đã thể hiện hoạt tính kháng khuẩn thực tế. Giáo sư Barzilay đã không đăng ký bằng sáng chế cho kết quả này mà công khai toàn bộ. Đó là với mục đích để ngay cả những quốc gia nghèo bị bỏ quên vì các căn bệnh không mang lại lợi nhuận cũng có thể sử dụng nguyên vẹn danh mục này.

Phương thức này hiện vẫn đang mở rộng sang các mầm bệnh khác. Vào tháng 7 năm 2026, ấn phẩm học thuật BioTechniques đã đưa tin về nghiên cứu sử dụng trí tuệ nhân tạo để tìm ra ứng viên kháng sinh mới nhắm vào vi khuẩn lặn đa kháng thuốc. Điều này có nghĩa là cách tiếp cận có cùng khung sườn như cách tạo ra Halicin—tức máy móc rà soát trên diện rộng trước rồi con người kiểm chứng trên bàn thí nghiệm—đang lan sang các loại vi khuẩn khác. Cùng tháng đó, CNBC cũng đưa tin một nhà nghiên cứu khác tại MIT đang thử nghiệm hướng đi dùng virus thay cho vi khuẩn để tiêu diệt vi khuẩn kháng thuốc. Tuy nhiên, dù máy móc có sàng lọc hàng trăm triệu ứng viên, việc xác nhận vi khuẩn trên đĩa cấy sống hay chết cuối cùng vẫn nhờ vào mắt người. Máy móc nhìn rộng, con người kiểm chứng cự ly gần. Nếu chỉ có một trong hai thì đã không thể đánh bại được vi khuẩn đa kháng thuốc.

Những vấn đề vẫn chưa có lời giải

Dù hoan nghênh thành tựu này, giới học thuật vẫn cho rằng cần phải xem xét lại từ chính dữ liệu huấn luyện nằm bên dưới. 2.335 hợp chất mà Chemprop học được cũng như kho lưu trữ ZINC15 được rà soát sau đó đều thiên về các phân tử mà con người đã từng tổng hợp hoặc đưa lên cơ sở dữ liệu. Có ý kiến chỉ ra rằng mạng nơ-ron rất khó đưa ra phán đoán vượt ra ngoài không gian hóa học mà nó chưa từng thấy. Mối lo ngại là nếu dữ liệu huấn luyện thiên lệch về một thư viện cụ thể, các ứng viên mà mô hình tìm ra cũng chỉ luẩn quẩn trong bóng râm đó. Các tạp chí học thuật, bao gồm cả Nature, cũng đã nhiều lần đặt ra câu hỏi về giới hạn của việc tìm kiếm thuốc mới bằng trí tuệ nhân tạo khi bị giam hãm trong phân phối dữ liệu.

Kỳ vọng dành cho Halicin cũng đi kèm những dấu hỏi. Việc tiêu diệt vi khuẩn trong đĩa nuôi cấy và trên chuột là một chuyện hoàn toàn khác với việc sử dụng an toàn trên cơ thể người. Giới phát triển thuốc mới từ lâu đã xem đây là lẽ thường: từ một ứng viên thuốc mới vượt qua thử nghiệm trên động vật và kiểm tra độc tính, đi qua mọi cửa ải từ thử nghiệm lâm sàng giai đoạn 1 đến giai đoạn 3 để trở thành thuốc kê đơn, thường phải mất hơn mười năm cùng tỷ lệ thất bại rất cao. Nhiều năm trôi qua kể từ khi được công bố, Halicin vẫn chưa vượt qua được ngưỡng cửa thử nghiệm lâm sàng trên người, và xuất hiện quan điểm thận trọng cho rằng liệu hiệu quả kháng khuẩn trong phòng thí nghiệm có chuyển hóa thành hiệu quả điều trị thực tế cho bệnh nhân hay không vẫn là phần việc cần phải kiểm chứng trong các thử nghiệm lâm sàng tương lai.

Các mô hình chẩn đoán như Mirai và Sybil cũng phải đối mặt với câu hỏi tương tự. Hiện tượng một mô hình được huấn luyện bằng hình ảnh thu thập từ một bệnh viện, một loại thiết bị hay một nhóm dân số nhất định bị giảm độ chính xác khi đứng trước các khu vực khác và thiết bị khác—gọi là vấn đề dịch chuyển phân phối—là điểm yếu liên tục được ghi nhận trong toàn bộ lĩnh vực trí tuệ nhân tạo y tế. Người ta chỉ ra rằng nếu nhóm dân số trong dữ liệu huấn luyện bị thiên lệch về một phía, dự đoán đối với những bệnh nhân nằm ngoài nhóm đó có thể kém tin cậy hơn. Khi nhiều trường hợp mô hình xây dựng từ dữ liệu của một bệnh viện không duy trì được hiệu năng ở bệnh viện khác được biết đến, giới học thuật ngày càng nhấn mạnh rằng việc kiểm chứng sâu rộng trên nhiều khu vực và nhiều thiết bị phải được đặt lên hàng đầu trước khi áp dụng vào lâm sàng.

Nội dung chương này dựa trên các tư liệu được công bố tính đến tháng 8 năm 2026.

Ghi chú của tác giả. Vai trò của con người đã chuyển dịch về đâu

Ba nhân vật ở phần 1 đã giao phó việc dự đoán cho máy móc. Hassabis và Jumper để máy móc đoán trước protein sẽ cuộn gấp thành hình dạng nào, còn Barzilay để máy móc đoán phân tử nào sẽ tiêu diệt mầm bệnh. Cả ba trường hợp đều là những ví dụ mà máy móc đã hoàn toàn vượt lên trước công việc mà trước đây con người phải kiểm chứng từng bước một bằng thực nghiệm.

Với tư cách là một luật sư, điều khiến tôi băn khoăn khi đọc ba chương này chính là khoảng cách giữa việc đoán trúng và sự thấu hiểu. AlphaFold đã đưa ra cấu trúc của 200 triệu protein nhưng không giải thích được vì sao chúng lại cuộn gấp như thế, còn Chemprop đã chỉ ra Halicin nhưng cơ chế tác dụng của chất đó mãi về sau mới được làm sáng tỏ. Máy móc trao câu trả lời trước, và để lại lý do như một bài tập về nhà cho con người.

Vì vậy, bên cạnh cỗ máy dự đoán, vị trí của con người không hề biến mất mà đã dịch chuyển. Đó là vị trí kiểm chứng dự đoán của máy móc trên bàn thí nghiệm, truy vấn lý do vì sao nó đúng, và ước lượng khi nào nó sẽ sai. Phần 1 là nơi câu trả lời của cuốn sách này lần đầu tiên lộ diện: dự đoán chỉ là nguyên liệu của khoa học, chứ không phải thành phẩm. Khi máy móc tích lũy nguyên liệu càng nhanh, bàn tay con người biến nguyên liệu đó thành khoa học lại càng trở nên quan trọng hơn.

Phần 2

Cỗ máy thiết kế những thứ chưa từng tồn tại

Chương 4. David Baker: Tạo ra những protein chưa từng tồn tại

Sự tiến hóa không chờ đợi

Đã từng có thời điểm hàng trăm thùng máy tính tràn ra tận hành lang trong một góc tòa nhà Khoa Hóa sinh thuộc Đại học Washington ở Seattle. Vào giữa những năm 1990, trong phòng thí nghiệm của Giáo sư David Baker, những cỗ máy đó ngày đêm miệt mài giải quyết duy nhất một bài toán. Đó là việc làm thế nào một chuỗi axit amin nối liền nhau thành một dải lại có thể tự gấp lại một cách chính xác thành một hình dạng ba chiều duy nhất, tựa như một sợi chỉ dài tự gấp mình thành hình con hạc giấy.

Baker vốn không phải là một nhà sinh học ngay từ đầu. Sinh năm 1962 tại bang Washington, khi vào Đại học Harvard vào đầu những năm 1980, ông vẫn còn đắm chìm trong triết học và khoa học xã hội. Thế nhưng, một lớp học sinh học phát triển mà ông tình cờ theo học vào năm cuối đại học đã thay đổi ông. Trong khi các lớp triết học cứ luẩn quẩn quanh những câu từ của hàng trăm năm trước, thì trong lớp sinh học, những phát hiện cách đó vài năm đã lập tức viết lại sách giáo khoa. Baker rời bỏ mê cung ngôn từ để chuyển đến UC Berkeley, rồi nhận bằng tiến sĩ sinh học phân tử dưới sự hướng dẫn của Giáo sư Randy Schekman, người sau này đoạt giải Nobel Sinh lý học hoặc Y học.

Năm 1993, ông trở về quê nhà Seattle với tư cách là trợ lý giáo sư tại Đại học Washington. Khi lần giở từng con đường trao đổi chất và tín hiệu miễn dịch bên trong tế bào, ông sớm đứng trước bức tường cổ xưa của ngành hóa sinh. Dù trình tự axit amin được xác định thì cấu trúc 3 chiều gấp lại cũng tự khắc được định hình, nhưng để thực sự xác nhận hình dạng đó, người ta phải kết tinh protein rồi chiếu tia X, và nhiều protein kiên quyết từ chối kết tinh. Không hiếm nhà nghiên cứu cống hiến cả đời mà khi nghỉ hưu vẫn chưa từng tận mắt thấy được một cấu trúc nào. Đến việc xác định hình dạng của một chuỗi đơn lẻ còn khó khăn như vậy, nên việc tạo ra protein mới với hình dạng mong muốn là điều không ai dám tưởng tượng vào thời kỳ đó.

Tại đây, Baker đã đảo ngược câu hỏi. Thay vì đi theo chiều tự nhiên là đưa trình tự vào để nhận về cấu trúc, ông tự hỏi liệu có thể vẽ ra cấu trúc mong muốn trước rồi dùng máy tính tính ngược lại trình tự phù hợp hay không. Tự nhiên đã mất 4 tỷ năm trải qua các đột biến ngẫu nhiên và cuộc cạnh tranh sinh tồn để chọn lọc ra các protein ngày nay. Ý tưởng không cần chờ đợi khoảng thời gian đằng đẳng đó mà tạo ngay phân tử cần thiết theo đúng hình dạng mong muốn chính là nguồn gốc khai sinh ra chương trình Rosetta.

Vấn đề nằm ở khối lượng tính toán. Số lượng trường hợp mà một chuỗi có thể gấp lại lớn đến mức một phòng máy chủ đại học không thể nào gánh nổi, và Baker đã giải quyết vấn đề này bằng cách mượn máy tính nhàn rỗi của các công dân trên khắp thế giới thông qua

các dự án và trò chơi mang tên Rosetta@home và Foldit. Dựa trên những phép tính đó, vào năm 2003, Top7 – một loại protein chưa từng tồn tại trong tự nhiên – đã được tạo ra chính xác theo bản thiết kế, và sau khi AlphaFold xuất hiện vào năm 2020, hàng loạt công cụ dùng học sâu để xây dựng các protein chưa từng có đã liên tiếp ra đời.

Triết lý vẽ nên phân tử cần thiết ngay tại chỗ mà không cần nhờ đến bàn tay tiến hóa nay đã vượt ra ngoài khuôn viên phòng thí nghiệm đại học, mở rộng từ vắc-xin và thuốc chống ung thư sang đến các enzym phân hủy nhựa và chất ô nhiễm. Ý tưởng không còn phải chờ đợi thời gian biểu 4 tỷ năm của tự nhiên đã bắt đầu từ một phòng làm việc nhỏ của vị trợ lý giáo sư ở Seattle và tiến xa đến tận ngày hôm nay.

Viện nghiên cứu Baker

Ba giờ chiều thứ Tư, trên bàn phòng phân tích của Viện Thiết kế Protein (IPD) thuộc Đại học Washington, những hộp sô-cô-la và bánh donut được xếp chồng lên nhau. Đó là món ăn nhẹ do Giáo sư David Baker tự bỏ tiền túi ra mua. Những người trẻ chuyên ngành máy tính vừa mới đây còn đang vật lộn với sai số tọa độ nguyên tử trước màn hình liên đứng dậy. Các nhà nghiên cứu phòng thí nghiệm ướt suốt đêm qua cầm pipet bên bàn thí nghiệm cũng bước ra ngoài. Hai nhóm hòa vào nhau, vừa cắn một miếng bánh donut vừa trò chuyện. "Cấu trúc liên kết này do RFDiffusion tạo ra trông thế nào?" phía máy tính hỏi. "Tốt đấy, ngày mai hãy đưa vào khuẩn E. coli xem sao," phía thực nghiệm đáp lại. Buổi gặp gỡ này chính là 'Thứ Tư Sô-cô-la'. Đó là thói quen mà IPD duy trì đều đặn hàng tuần không hề gián đoạn.

Chỉ riêng thói quen này không thể giải thích hết về IPD. IPD không phải là một phòng thí nghiệm chỉ giỏi tạo ra một bộ thuật toán, mà là một hệ sinh thái gắn kết giữa tính toán và thực nghiệm. Vào giữa những năm 2000, phòng thí nghiệm Baker tràn ngập nhân lực nhưng lại thiếu máy tính. Dù các thùng máy tỏa nhiệt xếp đầy dọc các hành lang, chúng vẫn quá ít ỏi để thiết kế hàng loạt vắc-xin hay thuốc điều trị các căn bệnh nan y. Để giải quyết vấn đề này, Giáo sư Baker đã chính thức thành lập IPD vào năm 2012. Đây là nơi các nhà khoa học máy tính, nhà vật lý và nhà hóa sinh cùng làm việc dưới một mái nhà để tạo ra protein từ con số không.

Năm 2019, Quỹ Bill & Melinda Gates và Dự án Audacious của TED đã đầu tư hàng trăm triệu đô la vào thử nghiệm này. Số tiền đó không được dùng để xây dựng các tòa nhà. Nó được rót vào việc thiết kế các enzym tiêu diệt nấm gây hại mùa màng ở các nước nghèo, và các loại vắc-xin vô hiệu hóa các virus gây đại dịch.

Một cơ chế khác do Giáo sư Baker thiết lập là RosettaCommons. Hơn 100 trường đại học và viện nghiên cứu trên khắp thế giới, bao gồm Đại học Washington, Đại học Johns Hopkins, UCSF và Đại học Stanford, cùng nhau chỉnh sửa mã nguồn Rosetta trên một kho lưu trữ GitHub duy nhất. Khi một nghiên cứu sinh tạo ra một hàm năng lượng mới, mã nguồn đó sẽ lập tức được mở cho toàn bộ RosettaCommons. Cách làm này khác hẳn với việc đăng ký bằng sáng chế rồi khóa lại và xem nhau như đối thủ cạnh tranh. Hơn 90 học trò từng qua đào tạo tại IPD đã trở thành giáo sư tại các trường đại học trên toàn thế giới và nhân rộng phương thức này vào các phòng thí nghiệm của riêng họ.

Công nghệ cũng phát triển theo con đường này. Bắt đầu từ RoseTTAFold do Giáo sư Baek Min-kyung thuộc Khoa Khoa học Sinh học Đại học Quốc gia Seoul dẫn dắt hoàn thiện khi còn làm việc tại IPD, cho đến RFDiffusion vẽ khung protein mới và ProteinMPNN điền trình tự tương ứng với khung đó, tất cả đều lần lượt ra đời tại đây. Với sự kết hợp này, những cấu trúc protein mà tự nhiên trong suốt 4 tỷ

năm chưa từng thử nghiệm đã được máy tính vẽ ra chỉ trong vài chục giây.

Công cụ đưa năng lực tính toán này ra ngoài phòng thí nghiệm là ColabFold. Trước đây, để chạy các phép tính như vậy cần phải có máy chủ card đồ họa trị giá hàng chục triệu won mỗi chiếc. ColabFold đã mở ra khả năng tính toán này hoàn toàn miễn phí chỉ qua một cửa sổ trình duyệt Google Colab. Một học giả trẻ nghiên cứu nấm gây bệnh héo rũ trên lúa mì ở châu Phi có thể dùng một chiếc Chromebook cấu hình thấp kết nối với máy chủ ở Seattle để tạo ra cấu trúc peptide nhân tạo tiêu diệt nấm. Nhóm nghiên cứu IPD cùng với Đại học Oxford và Viện Bộ gen Sáng tạo thuộc UC Berkeley đã phun protein ức chế được tạo ra bằng phương pháp này lên nấm thực tế. Kết quả là sự phát triển của nấm đã giảm 99,4% và khả năng chịu nhiệt của cây trồng tăng hơn 4 độ C.

Tốc độ kiểm chứng bản vẽ máy tính bằng vật thể thực tế cũng là vũ khí của IPD. Nhóm nghiên cứu gọi phương pháp này là 'Hóa sinh Cao bồi' (Cowboy Biochemistry). Sáng thứ Hai, họ vẽ hàng ngàn bản thiết kế protein bằng RFDiffusion và ProteinMPNN; thứ Ba tổng hợp DNA; thứ Tư đưa vào vi khuẩn E. coli để lên men; và chiều thứ Năm đo lực liên kết. Đây là cách bỏ qua quy trình tinh chế và đưa thẳng dịch nuôi cấy tế bào vào thiết bị phân tích. Sai số nằm trong khoảng 0,4 angstrom. Chỉ trong ba ngày, giả thuyết trên máy tính đã trở thành dữ liệu thực tế, và dữ liệu đó lại được phản ánh vào thiết kế của tuần tiếp

theo.

Vắc-xin hạt nano tự lắp ráp do nhóm của Giáo sư Neil King phát triển cũng ra đời từ hệ sinh thái này, trải qua các thử nghiệm lâm sàng và được đưa vào tiêm chủng thực tế.

Giáo sư Baker cũng xử lý lợi nhuận từ bằng sáng chế sinh ra từ hệ sinh thái này theo một cách khác biệt. Trong hơn 100 hợp đồng cấp phép sáng chế và 21 doanh nghiệp công nghệ sinh học tách ra từ IPD, ông đều đưa vào điều khoản 'Global Health Carve-out' (Ngoại lệ Y tế Toàn cầu). Đây là điều khoản quy định khi phát triển thuốc điều trị các căn bệnh chỉ lây lan ở các nước nghèo như sốt rét hay bệnh Chagas, họ sẽ không thu phí bản quyền sáng chế mà chuyển giao công nghệ miễn phí cho các tổ chức phi lợi nhuận như Tổ chức Y tế Thế giới. Đó là quy tắc: thị trường sinh lời thuộc về công ty, còn các căn bệnh không sinh lời được trao trả lại cho thế giới. Hiếm có viện nghiên cứu nào đưa một điều khoản như vậy vào từng dòng hợp đồng bằng sáng chế.

Hệ sinh thái này đang mở rộng sang cả các thử nghiệm gắn kết protein với chip silicon. Các nhà nghiên cứu IPD đang thử nghiệm cảm biến 'mũi nhân tạo', đặt các thụ thể protein được thiết kế lên bề mặt bán dẫn để chuyển thành tín hiệu điện tử ngay khi các phân tử mầm bệnh hay thành phần ma túy lơ lửng trong không khí chạm vào thụ thể. Đó là bức tranh nơi các protein được vẽ ra từ tính toán không chỉ dừng lại trên bản vẽ giấy mà bước ra tận bàn xét nghiệm bệnh viện và cổng

kiểm soát an ninh sân bay.

Những phép tính bắt đầu từ một hành lang ở Seattle giờ đây đã kết nối tới các phòng thí nghiệm, bệnh viện trên toàn thế giới và cả những cánh đồng lúa mì ở châu Phi.

Khởi đầu của Rosetta

Đã có một đêm chuông báo cháy vang lên trong tòa nhà Khoa Hóa sinh của Đại học Washington tại Seattle. Xe cứu hỏa lao tới nhưng không hề có đám cháy nào. Nguyên nhân là do hàng trăm chiếc máy tính cũ kỹ mà phòng thí nghiệm của Giáo sư Baker xếp thành từng tầng dọc hành lang, dưới gầm bàn và cả trong góc phòng trà. Hệ thống làm mát không thể chịu nổi nhiệt lượng từ các thùng máy chạy ngày đêm, và cầu dao của tòa nhà sập xuống như cơm bữa. Lính cứu hỏa Seattle đã phải đến kiểm tra nhiều lần rồi quay về sau khi xác nhận đó chỉ là máy tính chứ không phải hỏa hoạn. Đó là một phòng thí nghiệm mà tiếng quạt tản nhiệt của thùng máy thay thế ống nghiệm và kính hiển vi để lấp đầy màn đêm.

Câu chuyện Giáo sư Baker từ một sinh viên triết học chuyển hướng sang sinh học đã được đề cập ở phần trước. Điểm khởi đầu của ông là theo học dưới sự hướng dẫn của Randy Schekman tại UC Berkeley, nơi ông tìm hiểu các quy luật vật lý chi phối sự di chuyển của protein bên trong tế bào.

Được bổ nhiệm làm trợ lý giáo sư tại Đại học Washington vào năm 1993, Giáo sư Baker đã đối mặt với bức tường mà ngành sinh học cấu trúc đã vấp phải suốt nửa thế kỷ. Khi nhận giải Nobel năm 1972, nhà hóa học Christian Anfinsen đã đưa ra giả thuyết rằng chuỗi axit amin trong nước sẽ tự gấp lại thành hình dạng ba chiều ổn định có năng lượng tự do thấp nhất. Vấn đề là việc xác nhận hình dạng đó trên thực tế. Để làm sáng tỏ cấu trúc của một protein bằng phương pháp tinh thể học tia X, một nghiên cứu sinh phải mất từ 4 đến 5 năm, và chi phí cho mỗi lần vượt quá 130 triệu won. Trong khi đó, thông tin trình tự gen đã tích lũy tới hàng trăm triệu mục, nhưng số lượng protein thực sự được giải mã cấu trúc chỉ dừng lại ở con số 180.000. Khoảng cách giữa trình tự và cấu trúc ngày càng bị nới rộng.

Giáo sư Baker quyết định lấp đầy khoảng trống này bằng tính toán chứ không phải bằng thực nghiệm. Ý tưởng của ông là nếu có thể chuyển đổi chính xác chuyển động phân tử của nước và lực giữa các nguyên tử thành các công thức toán học, thì có thể tìm ra hình dạng gấp bên trong máy tính mà không cần đến máy gia tốc khổng lồ. Chương trình được khởi xướng theo cách này vào giữa những năm 1990

chính là Rosetta. Thế nhưng, nó sớm đụng phải một bức tường khác. Số lượng trường hợp mà một chuỗi axit amin có thể gấp lại lên tới 10 lũy thừa 300. Con số này còn lớn hơn tổng số nguyên tử tồn tại trong vũ trụ. Để đưa từng con số này vào tính toán thì dù gộp tất cả máy tính hiện có lại cũng phải mất khoảng thời gian bằng cả tuổi của vũ trụ. Số lượng máy tính để tính toán thiếu thốn trầm trọng. Chuông báo cháy được nhắc đến ở trên chính là dấu tích của cơn khát tính toán đó.

Đầu những năm 2000, Giáo sư Baker đã tìm kiếm câu trả lời từ bên ngoài trường đại học. Ông khởi động dự án điện toán phân tán có tên Rosetta@home, cho phép máy tính của người dân bình thường trên khắp thế giới chạy các phép tính thay cho phòng lab trong thời gian nhàn rỗi.

Hơn 350.000 người từ 190 quốc gia đã cài đặt chương trình, và khi máy tính chuyển sang trạng thái nhàn rỗi, hình ảnh chuỗi axit amin đang gấp lại sẽ hiện lên trên màn hình như một trình bảo vệ màn hình. Máy tính trong phòng khách của các gia đình ở Mỹ, phòng máy tính trường tiểu học ở Đức và các thư viện công cộng ở Hàn Quốc đều tham gia vào phép tính này. Quy mô năng lực tính toán này đã vượt xa trung tâm siêu máy tính của hầu hết các trường đại học.

Năm 2008, họ tiến thêm một bước nữa. Những người tham gia đã đăng đề xuất lên diễn đàn của Rosetta@home. Họ bày tỏ rằng thay vì chỉ ngồi nhìn máy tính quay cuồng tính toán, họ muốn tự tay kéo thử các nút thắt. Trò chơi Foldit đã ra đời như vậy. Thành quả đã được chứng minh vào năm 2011. Cấu trúc protease retrovirus của virus khỉ Mason-Pfizer (M-PMV) vốn khiến các phòng thí nghiệm tầm cỡ đoạt giải Nobel bế tắc suốt 15 năm, nhưng 50.000 người chơi Foldit đã giải mã được cấu trúc này chỉ trong 3 tuần với sai số không vượt quá 0,5 angstrom. Kết quả hoàn toàn trùng khớp khi đối chiếu với ảnh chụp tia X thực tế. Các game thủ năm đó đã được ghi tên với tư cách là đồng tác giả của bài báo trên tạp chí học thuật Nature Structural & Molecular Biology.

Nguyên lý giúp Rosetta đạt được những kết quả này bất ngờ thay lại khá mộc mạc. Nó cắt trình tự protein thành các đoạn ngắn gồm 9 hoặc 3 đơn vị, rồi lấy nguyên gốc độ gấp thực tế của các đoạn tương tự từ cơ sở dữ liệu cấu trúc protein đã biết. Thay vì tính toán chuyển động của từng nguyên tử ngay từ đầu, phương pháp này ghép nối các mảnh ghép mà tự nhiên đã kiểm chứng sẵn như những khối xếp hình. Mỗi khi lắp thêm một mảnh mới, Rosetta sẽ chấm điểm xem kết quả có ổn định hơn trước hay không. Nếu điểm số tốt hơn, nó sẽ chấp nhận ngay; còn nếu điểm số kém đi, nó vẫn chấp nhận với một xác suất thấp. Chỉ khi làm như vậy, quá trình tính toán mới không bị mắc kẹt ở một hình dạng lơ lửng lừng chừng mà có thể tìm đến hình dạng thực sự ổn định. Phương pháp này đã được chứng minh tại CASP, cuộc thi dự đoán diễn ra hai năm một lần. Vào năm 1994 khi chưa có Rosetta, điểm số độ chính xác dự đoán tại CASP1 chỉ dừng ở mức 20 đến 30 điểm. Đến năm 1998 khi Rosetta lần đầu xuất hiện tại CASP3, con số đã nhảy vọt lên 50 đến 60 điểm, và tại CASP5 năm 2002 đã tăng lên mức 70 đến 75 điểm. Năm 2003, Top7 – một loại protein không có trong tự nhiên – được tạo ra chỉ từ bản thiết kế và đã được tổng hợp trên thực tế, với sai số giữa thiết kế máy tính và cấu trúc tinh thể thực tế chỉ là 1,2 angstrom. Đó là mức sai số tương đương với đường kính của một nguyên tử carbon. Top7 là hình thái mà tự nhiên chưa từng tạo ra trong quá trình tiến hóa. Đây là trường hợp đầu tiên con người vẽ ra hình dạng mong muốn trước rồi hiện thực hóa nó thành vật thể thực, thay vì ngồi chờ sự lựa chọn ngẫu nhiên của tự nhiên.

Dòng chảy này vẫn đang tiếp diễn vào năm 2026. Tháng 5 năm đó, tại lễ hội khoa học quy mô lớn nhất Brazil do Trung tâm Nghiên cứu Năng lượng và Vật liệu Brazil tổ chức, 30.000 khách tham quan đã trực tiếp

trải nghiệm việc gấp protein bằng Foldit. Khung cảnh người dân chạm tay vào protein trước màn hình không khác mấy so với năm 2008 khi ấy. Khởi đầu từ bên trong máy tính, Rosetta giờ đây đã vượt ra ngoài bức tường phòng thí nghiệm đại học để mở rộng sân khấu tính toán tới phòng khách của người dân, các thư viện và các lễ hội.

Sự tiến hóa của công cụ thiết kế

Ngày 30 tháng 11 năm 2020, Giáo sư David Baker và các nhà nghiên cứu đã tập trung trước màn hình trong văn phòng của Viện Thiết kế Protein thuộc Đại học Washington. Trên màn hình hiển thị bảng điểm của cuộc thi mang tên CASP14. Đây là cuộc thi dự đoán bằng máy tính xem protein sẽ gấp lại thành hình dạng nào. AlphaFold 2 của Google DeepMind đã đạt điểm trung vị GDT_TS là 92,4 điểm. Độ chính xác này tương đương với cấu trúc tinh thể chụp mất vài tuần trong phòng thí nghiệm. Bài toán gấp protein làm đau đầu các nhà sinh học suốt 50 năm đã được giải quyết chỉ trong chớp mắt. Về mặt của những người có mặt tại đó chia làm hai thái cực. Có người cảm thấy công việc của mình sắp biến mất. Nhưng Giáo sư Baker thì khác. Ông nhận định rằng nếu AlphaFold 2 đã khai thông con đường đi từ trình tự đến cấu trúc, thì con đường theo chiều ngược lại cũng có thể mở ra. Đó là con đường vẽ ra hình dạng mong muốn trước, rồi để máy móc tìm ngược lại trình tự có thể gấp thành hình dạng đó.

Ngay từ ngày hôm sau, Giáo sư Baker đã tập hợp những nhân sự tinh hoa của viện nghiên cứu. Bài báo về AlphaFold 2 khi đó vẫn chưa được công bố. DeepMind đã giấu kín mã nguồn để bảo vệ bằng sáng chế và quyền ưu tiên. Đội ngũ của Baker chỉ dựa vào các tài liệu thuyết trình tại hội nghị và các bài báo toán học để xây dựng lại mọi thứ từ đầu. Tốc độ nghiên cứu được đẩy nhanh khi Tiến sĩ Baek Min-kyung, một nhà sinh học tính toán người Hàn Quốc, đứng ra dẫn dắt dự án. RoseTTAFold được tạo ra như thế đã được đăng tải trên tạp chí Science vào tháng 7 năm 2021, đúng vào tuần mà bài báo về AlphaFold 2 xuất hiện. Ngay khi bài báo được xuất bản, nhóm của Baker đã mở mã nguồn miễn phí trên GitHub. Đó là vì họ muốn ngăn chặn nguy cơ một tập đoàn lớn độc quyền công nghệ.

Sự khác biệt giữa phương pháp cũ và phương pháp mới thể hiện rõ qua các con số. Trước đó, nhóm của Baker thiết kế protein mới bằng mô phỏng vật lý, tính toán từng lực tĩnh điện và lực van der Waals tác động giữa các nguyên tử. Chỉ cần số lượng nguyên tử cần tính toán tăng lên một chút, thời gian thực hiện đã tăng theo cấp số nhân. Tỷ lệ thành công khi thực sự tổng hợp các bản thiết kế tạo ra theo cách đó chưa tới một phần mười. Sau khi chuyển sang thiết kế dựa trên mạng nơ-ron, tỷ lệ thành công đã tăng vọt lên mức dao động quanh 90%. Đó là kết quả của việc mạng nơ-ron đã học trọn vẹn hình thái và quy luật để thay thế cho việc tính toán các định luật vật lý ở cấp độ từng nguyên tử.

RoseTTAFold chia thông tin thành ba nhánh và xử lý đồng thời: tầng đọc thứ tự các chữ cái axit amin, tầng tính toán khoảng cách giữa hai axit amin, và tầng vẽ tọa độ 3 chiều thực tế. Tầng thứ hai lập bảng khoảng cách giữa từng gốc axit amin và liên tục kiểm tra quy tắc rằng khoảng cách giữa ba điểm phải

nhất quán logic giống như ba cạnh của một hình tam giác. Cả ba tầng liên tục trao đổi kết quả và hiệu chỉnh lẫn nhau ở từng bước tính toán. Hình ảnh này tương tự như ba người cùng dệt vải và liên tục kiểm tra xem các sợi chỉ có bị rối vào nhau hay không. Phiên bản mới nhất, RoseTTAFold All-Atom, không chỉ xử lý protein mà còn đưa cả DNA, RNA và các hợp chất nhỏ vào cùng một phép tính.

Việc vẽ khung sườn do RFdiffusion đảm nhận. Nguyên lý của nó tương tự như AI tạo ảnh. Giống như Midjourney tạo ra bức tranh từ lớp nhiễu mờ昧, RFdiffusion rút ra khung protein từ đám mây nguyên tử phân tán ngẫu nhiên. Quá trình học diễn ra theo chiều ngược lại. Máy tính được huấn luyện lặp đi lặp lại phép tính phủ dần nhiễu lên cấu trúc protein có thực để biến nó thành trạng thái nhiễu hoàn toàn, rồi từng bước loại bỏ lớp nhiễu đó để phục hồi về hình dạng ban đầu. Khi đã thuần thục cách khử nhiễu, máy tính có thể tự tạo ra một khung protein hoàn toàn mới và hợp lý ngay cả từ trạng thái ban đầu chỉ toàn là nhiễu. Phép tính này được thực hiện dựa trên quy tắc toán học đối xứng SE(3), nơi kết quả không bị thay đổi dù các nguyên tử xoay hay dịch chuyển. Nếu đưa hình dạng mục tiêu mong muốn vào làm điều kiện, nó sẽ vẽ ra khung sườn mới phù hợp với điều kiện đó chỉ trong vài giây.

Chỉ riêng khung sườn thì chưa thể hoàn thiện một protein. Cần phải xác định xem loại nào trong số hai mươi loại axit amin sẽ được đặt vào vị trí nào. Công việc này do ProteinMPNN đảm nhận. Bằng cách truyền thông tin mà mỗi điểm trên khung trao đổi với các điểm lân cận qua mạng nơ-ron đồ thị, nó lần

lượt chọn và điền loại axit amin hoàn toàn phù hợp vào từng vị trí. Việc này tương tự như cắt may vải theo mẫu rập quần áo. Trình tự thu được sau đó sẽ được đưa trở lại vào AlphaFold hoặc RoseTTAFold để dự đoán cấu trúc. Nếu hình dạng dự đoán không khớp với khung sườn ban đầu trong phạm vi sai số cấp nguyên tử 1,0 angstrom, bản thiết kế đó sẽ bị loại bỏ ngay lập tức. Đây là phương thức sàng lọc trước trên máy tính và chỉ chuyển những ứng viên đạt yêu cầu sang phòng thí nghiệm.

Kể từ khi ba công cụ này được liên kết như một chuỗi xích, cách thức làm việc của viện nghiên cứu đã hoàn toàn thay đổi. Trước đây, việc thiết kế một protein mới bằng tính toán vật lý, sau đó tổng hợp và xác minh phải mất từ vài tuần đến vài tháng; nhưng giờ đây, hàng ngàn ứng viên có thể được phác thảo và sàng lọc trong máy tính chỉ trong một ngày, và chỉ một số ít vượt qua mới được chuyển đến phòng thí nghiệm. Khi chu kỳ máy tính tưởng tượng và phòng thí nghiệm kiểm chứng được đẩy nhanh như vậy, số lượng giả thuyết mà một nhà nghiên cứu có thể thử nghiệm trong suốt cuộc đời đã thay đổi hoàn toàn.

Xu hướng này vẫn không dừng lại trong năm 2026. Phòng thí nghiệm Baker đã phát hành mã nguồn mở RFdiffusion3 vào tháng 12 năm 2025. Bằng cách mở rộng các loại phân tử có thể xử lý so với phiên bản trước, họ đã đưa cả các protein liên kết với DNA

và những enzyme tinh vi vào phạm vi thiết kế. Năm 2026, thành quả cũng đã xuất hiện trong việc thiết kế kháng thể. Sau khi huấn luyện lại RFdiffusion2 bằng dữ liệu kháng thể và kết hợp với kỹ thuật sàng lọc bề mặt nấm men, một bài báo đăng trên tạp chí Nature đã công bố việc tạo ra thành công từ đầu các kháng thể có thể bám chặt vào vị trí đích mong muốn với độ chính xác ở cấp độ nguyên tử. Kháng thể vốn là lĩnh vực từ trước đến nay phải dựa vào việc kích thích phản ứng miễn dịch ở động vật sống, hoặc sàng lọc từng ứng viên một trong số hàng chục ngàn ứng viên. Nỗ lực thay thế phương pháp lâu đời đó bằng quy trình thiết kế trên máy tính giờ đây đã bắt đầu những bước đi đầu tiên.

Đã sáu năm trôi qua kể từ khi Giáo sư Baker đổi hướng trước bảng điểm CASP14. Trong khoảng thời gian đó, phòng thí nghiệm của ông đã chuyển dịch hoàn toàn từ nơi dự đoán protein sang nơi kiến tạo protein. Quyết định duy nhất đó đã sinh ra RoseTTAFold, và RoseTTAFold lại sinh ra RFdiffusion cùng ProteinMPNN. Ngay trong khoảnh khắc này, tại một phòng thí nghiệm nào đó, người ta vẫn đang đặt ba công cụ này cạnh nhau để vẽ nên một loại protein hoàn toàn mới chưa từng tồn tại trên thế giới.

Những phân tử tạo ra từ hư không

Quy trình phát triển thuốc truyền thống bắt đầu bằng việc tiêm protein đích vào cơ thể chuột hoặc lạc đà không bướu (llama) sống. Đó là cách chờ đợi hệ miễn dịch của động vật tình cờ tạo ra kháng thể trong vài tháng, rồi lấy tế bào lách ra để vớt vát lấy một ứng viên trong số hàng chục triệu ứng viên. Đưa protein có nguồn gốc từ động vật vào cơ thể người đôi khi còn gây ra phản ứng đào thải miễn dịch. Giáo sư Baker đã đảo ngược hoàn toàn phương pháp này. Chỉ cần nhập hình dạng 3D của protein đích vào máy tính, trí tuệ nhân tạo sẽ tạo ra từ đầu một protein mới có thể bám chặt vào hình dạng đó.

Nơi đầu tiên mà các protein do máy tính thiết kế gây kinh ngạc cho thế giới chính là vắc-xin. Giáo sư Neil King tại Đại học Washington đã thiết kế các hạt nano nhân tạo, trong đó sáu mươi mảnh protein tự lắp ráp thành hình quả bóng đá. Khi gắn dày đặc các vị trí liên kết của protein gai SARS-CoV-2 lên bề mặt đó, nồng độ kháng thể trung hòa đã tăng hơn một trăm lần so với vắc-xin hiện có. Loại vắc-xin này đã vượt qua thử nghiệm lâm sàng giai đoạn 3 dưới tên gọi SKYCovione và nhận được phê duyệt chính thức từ Bộ An toàn Thực phẩm và Dược phẩm Hàn Quốc vào năm 2022. Đây là trường hợp đầu tiên một protein được máy tính thiết kế từ đầu đến cuối được đưa vào cơ thể người như một loại dược phẩm chính thức.

Công nghệ tương tự cũng lan rộng sang các phòng thí nghiệm nghiên cứu ung thư và sa sút trí tuệ. Nhóm nghiên cứu đã thiết kế một protein nhân tạo làm vô hiệu hóa protein lá chắn PD-L1 mà tế bào ung thư đưa ra để tránh sự tấn công của tế bào miễn dịch, cùng với thụ thể tế bào T (TCR) nhân tạo chỉ nhận

diện tế bào ung thư. Trên mô hình động vật mắc ung thư tuyến tụy vốn nổi tiếng là khó điều trị, kích thước khối u đã giảm đi rõ rệt. Trong nghiên cứu bệnh Alzheimer, họ đã tạo ra một protein bẫy dạng bánh kẹp (sandwich) để bao bọc và giam giữ các khối protein amyloid beta vốn cắt đứt kết nối giữa các tế bào não. Khi đưa chiếc bẫy này vào môi trường nuôi cấy tế bào thần kinh, tốc độ tích tụ amyloid đã chậm lại đáng kể.

Bất kỳ ai từng viết mã đều sẽ không thấy xa lạ với quy trình này, bởi đó là việc xác định trước kết quả mong muốn rồi viết thuật toán ngược lại. Tuy nhiên, đối tượng mà Giáo sư Baker xử lý không phải là những con số trên màn hình, mà là sinh mạng con người và nhiệt độ của Trái Đất. Tiến hóa tự nhiên đã tạo ra 2 triệu loài protein, nhưng mục đích duy nhất của chúng chỉ là sinh tồn và sinh sản. Khả năng phân hủy nhựa hay thu giữ carbon dioxide vốn dĩ không nằm trong kế hoạch của tiến hóa. Nhóm nghiên cứu của Baker giờ đây đã bắt tay vào việc thiết kế từ đầu các protein lấp đầy khoảng trống đó, những cỗ máy phân tử thu giữ carbon mà tự nhiên có lẽ phải mất hàng chục triệu năm mới tình cờ tạo ra được. Khả năng vốn không hề tồn tại ở vi sinh vật tự nhiên

nay đang được viết mới bằng những dòng mã.

Vào tháng 5 năm 2026, viện nghiên cứu thông báo rằng họ đã thiết kế thành công một vỏ protein hình virus lớn gấp hai đến ba lần bình thường bằng cách kết hợp các mảnh ngũ giác và lục giác. Điều này tương đương với việc tạo ra một hộp vận chuyển lớn hơn để đưa liệu pháp gen vào sâu bên trong tế bào. Cũng trong tháng đó, viện đã mở một phòng nghiên cứu chuyên biệt về ung thư học và tuyên bố sẽ mở rộng thiết kế protein tập trung vào ung thư và các bệnh miễn dịch. Khởi đầu vào năm 2003 từ một protein nhân tạo từng bị coi là không thể gấp cuộn, giờ đây công việc này đã vượt qua vắc-xin và thuốc chống ung thư để vươn tới công nghệ vận chuyển liệu pháp gen.

Giải Nobel năm 2024

2 giây. Đó là thời gian cần thiết để thiết kế một protein chưa từng tồn tại trên Trái Đất. Năm 2021, vài nhà nghiên cứu đã tập hợp trước màn hình máy tính tại Viện Thiết kế Protein thuộc Đại học Washington. Trên màn hình hiển thị cấu trúc bề mặt của thụ thể yếu tố hoại tử khối u (TNF). Đó chính là loại protein gây bùng phát viêm nhiễm trong cơ thể bệnh nhân viêm khớp dạng thấp. Nhóm nghiên cứu đã nạp thông tin bề mặt của thụ thể này vào chương trình trí tuệ nhân tạo và ra lệnh. Thứ mà máy tính vẽ ra là một bộ khung protein nhân tạo chưa từng tồn tại trên Trái Đất. Nhóm nghiên cứu đã tạo ra gen theo bản thiết kế này rồi đưa vào vi khuẩn E. coli. Ba ngày sau, khi chiết xuất protein từ đĩa nuôi cấy và đo ái lực liên kết, con số thu được là 9 picomole. Tức là 10 lũy thừa âm 12 mol, hay nói cách khác là 9 phần 10 nghìn tỷ. Đây là ái lực liên kết khó lòng đạt tới theo phương pháp tiêm thụ thể vào chuột và lọc đã không bước rồi sàng lọc kháng thể suốt nhiều năm của các tập đoàn dược phẩm khổng lồ như Sanofi hay Roche.

Để hiểu tại sao cảnh tượng này lại quan trọng, chúng ta phải quay ngược lại nửa thế kỷ trước. Năm 1972, Christian Anfinsen khẳng định rằng chỉ cần biết trình tự là có thể xác định được hình dạng gấp cuộn, nhưng Cyrus Levinthal đã chỉ ra rằng số lượng cấu hình gấp cuộn có thể xảy ra là quá lớn, đến mức dù có chạy siêu máy tính trong suốt thời gian bằng tuổi của vũ trụ cũng không thể tìm ra câu trả lời. Đó chính là nghịch lý Levinthal. Baker đã đảo ngược vấn đề này và vào năm 2003, ông đã tạo ra Top7 – một protein nhân tạo gồm 93 axit amin đúng y như bản thiết kế. Cho đến lúc đó, giới học thuật vẫn tin rằng con người không thể bắt chước những protein mà 4 tỷ năm tiến hóa đã chọn lọc, và gọi thách thức của Baker là sự điên rồ; nhưng sự chế giễu đó đã phải im bật chỉ sau bài báo duy nhất này. Vài năm sau, Baker nhận Giải thưởng Feynman vinh danh những thành tựu trong kỹ thuật phân tử vi mô, và Eric Drexler, người được mệnh danh là nhà tiên tri của công nghệ nano, đã phát biểu tại lễ trao giải rằng thời đại chế tạo máy móc bằng cách xếp từng nguyên tử một đang đến. Phòng thí nghiệm của Baker đã biến

lời nói đó thành hiện thực cụ thể chứ không còn là lý thuyết.

Công nghệ không dừng lại ở đây. Năm 2023, phòng thí nghiệm của Baker cho ra mắt RFdiffusion để phác thảo khung xương và ProteinMPNN để lấp đầy các axit amin vào đó; sự kết hợp của hai chương trình này đã rút ngắn thời gian thiết kế từ vài tháng xuống còn vài giờ. Vắc-xin hạt nano COVID-19 do nhóm của Giáo sư Neil King phát triển bằng dòng công nghệ này đã vượt qua các thử nghiệm lâm sàng và được phê duyệt, trở thành trường hợp đầu tiên một protein được thiết kế từ đầu bằng máy tính được đưa vào cơ thể người một cách chính thức. Việc thao tác với protein như xử lý các hình học tô pô rất cuộc đã mang lại những vắc-xin và phương pháp điều trị thực sự cứu sống con người.

Dấu ấn từ các công cụ của Baker không chỉ giới hạn trong lĩnh vực dược phẩm. Bên cạnh các enzyme nhân tạo nhằm vào vi nhựa và khí nhà kính, người ta còn tạo ra các vòng protein nhân tạo giữ hai phân tử diệp lục ở góc độ chính xác để hấp thụ ánh sáng mặt trời bước sóng đỏ mà thực vật vốn bỏ lỡ, và cả những protein giúp sắp xếp tinh thể oxit kẽm – một vật liệu bán dẫn – theo cấu trúc mạng tinh thể mong muốn. Thứ mà Baker tạo ra không phải là thuốc điều trị cho một căn bệnh đơn lẻ, mà là công cụ dùng chung cho cả hóa học, sinh học lẫn khoa học vật liệu.

Vào ngày công bố giải Nobel Hóa học năm 2024, Viện Hàn lâm Khoa học Hoàng gia Thụy Điển đã trao một nửa giải thưởng cho David Baker, và nửa còn lại chia đều cho John Jumper và Demis Hassabis. Jumper và Hassabis nhận giải nhờ việc giải mã cấu trúc protein, còn Baker nhận giải nhờ việc tạo ra các protein chưa từng tồn tại. Jumper sau này đã đánh giá công trình của Baker là một thành tựu ghi dấu ấn trong lịch sử khoa học tự nhiên. Người giải quyết vấn đề và người tạo ra những bài toán mới đã cùng chia sẻ một giải thưởng trong cùng một năm. Ngay cả sau khi nhận giải, Baker vẫn không đóng cửa phòng thí nghiệm. Những công cụ do ông tạo ra không bị giấu sau bằng sáng chế mà được phát hành dưới dạng mã nguồn mở để bất kỳ nhà nghiên cứu nào trên thế giới cũng có thể tải về sử dụng, và từ nền tảng mã nguồn này đã có hơn hai mươi một công ty khởi nghiệp ra đời. Giải thưởng này đã trở thành một điển hình cho việc một phòng thí nghiệm đại học công khai chính phương pháp để thúc đẩy cả ngành công nghiệp phát triển, thay vì để các công ty lớn giấu kín kết quả như bí mật kinh doanh.

Năm 2026, phòng thí nghiệm này vẫn không ngừng lại. Công trình nghiên cứu dùng máy tính để phác thảo các kháng thể khớp với vị trí đích mong muốn ở cấp độ nguyên tử đã được đăng trên tạp chí Nature, và vào tháng 3 cùng năm, Washington Research Foundation đã đầu tư 7 triệu USD vào phòng thí nghiệm này để khởi động một dự án 4 năm nhằm đưa các bản thiết kế trong máy tính thành thuốc điều trị và vật liệu công nghiệp thực tế. Khoản đầu tư này không chỉ dùng cho chi phí nghiên cứu mà còn dùng để đào tạo các nhà khoa học mới và hỗ trợ khởi nghiệp. Giải Nobel là tấm huân chương ghi nhận chặng đường đã qua, nhưng ánh đèn tại tòa nhà nghiên cứu của Đại học Washington vẫn sáng đến tận đêm khuya. Họ không dừng lại ở vị trí đã nhận giải thưởng, mà đôi tay vẫn tiếp tục chuyển động để vẽ nên loại protein tiếp theo.

Những vấn đề vẫn chưa có lời giải

Cái nhìn của giới học thuật khi theo dõi những thành tựu này không hoàn toàn chỉ có lạc quan. Protein do máy tính phác thảo có thể gấp cuộn đúng theo bản thiết kế trong dung dịch đậm tinh khiết của ống nghiệm, nhưng tình hình bên trong tế bào sống lại rất khác. Bên trong tế bào là một không gian chật chội chen chúc đủ loại phân tử, nhiệt độ liên tục biến động và các nguyên tử luôn chịu sự dao động nhiệt động lực học không ngừng. Giới sinh học cấu trúc liên tục chỉ ra rằng các protein nhân tạo vốn hoàn hảo trong ống nghiệm khi đặt vào môi trường này có thể mất đi hình dạng gấp cuộn hoặc bám vào các phân tử không mong muốn, khiến chúng không thể hoạt động như ý định. Cũng có những trường hợp được báo cáo cho thấy ái lực liên kết dự đoán từ thiết kế chênh lệch rất lớn so với ái lực liên kết thực tế đo được bên trong tế bào. Tính toán cho thấy ổn định trong máy tính không đồng nghĩa với sự bảo

đảm rằng nó sẽ hoạt động bên trong tế bào.

Vấn đề tỷ lệ thành công cũng vẫn còn đó. Dù tỷ lệ sàng lọc ứng viên đã tăng vọt sau khi chuyển sang thiết kế bằng mạng nơ-ron, nhưng ngay cả những bản thiết kế đã vượt qua vòng này khi thực sự tổng hợp thì phần lớn vẫn không gấp cuộn được hoặc không thực hiện được chức năng mong muốn. Nhiều ý kiến cho biết việc phải kiểm chứng từng ứng viên một trong số hàng trăm đến hàng ngàn ứng viên tại phòng thí nghiệm để có được một thể liên kết bám chính xác vào mục tiêu là điều không hiếm gặp. Tính toán không thay thế được thực nghiệm mà chỉ sắp xếp lại hàng đầu của thực nghiệm, và việc kiểm chứng trên bàn thí nghiệm ướt (wet lab) vẫn là một cánh cổng không thể bỏ qua — đó là nhận định chung của các nhà nghiên cứu trong lĩnh vực này. Khác với dự đoán cấu trúc vốn có độ chính xác dự đoán cao, việc thiết kế xây dựng chức năng hoàn toàn mới từ đầu vẫn chưa có cả những tiêu chuẩn hoàn chỉnh để phân định trước thành công và thất bại.

Cũng có những mối lo ngại thuộc loại khác. Đó là vấn đề lưỡng dụng (dual-use), khi công nghệ phác thảo phân tử theo hình dạng mong muốn từ đầu này có thể không được dùng làm thuốc điều trị mà bị lợi dụng để tăng cường độc tố hoặc mầm bệnh. Năm 2024, khoảng một trăm nhà khoa học nghiên cứu thiết kế protein đã cùng ký tên vào cam kết 'Thiết kế sinh học AI có trách nhiệm' (Responsible AI x Biodesign), đồng lòng nhất trí sàng lọc các trình tự nguy hiểm ở giai đoạn tổng hợp gen và giám sát việc lạm dụng các công cụ thiết kế. Bản thân Giáo sư Baker cũng là một trong những người ký tên tham gia cam kết này. Sự cởi mở khi phát hành công nghệ dưới dạng mã nguồn mở từng là động lực thúc đẩy lĩnh vực này phát triển nhanh chóng, nhưng sự cởi mở đó cũng có thể đồng thời mở ra cánh cửa cho sự lạm dụng — sự căng thẳng này vẫn là một bài toán mà công nghệ này phải gánh vác.

Nội dung trong chương này dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Chương 5. Alán Aspuru-Guzik: Phòng thí nghiệm tự hành

Động lực mang tên khủng hoảng khí hậu

Vào đầu những năm 2000, tại hàng ghế đầu của một giảng đường ở Đại học Harvard, có một trợ lý giáo sư khoa hóa học trẻ tuổi đang ngồi. Trên màn hình hiển thị đồ thị nhiệt độ bề mặt Trái Đất và mực nước biển. Diễn giả là nhà khí hậu học, Ngài David King, cựu cố vấn khoa học trưởng của Thủ tướng Anh Tony Blair. Ngài King nói rằng nếu carbon dioxide tiếp tục tích tụ với tốc độ như hiện nay, trong vòng vài thập kỷ tới, một nửa bờ biển nước Anh sẽ chìm trong nước và bang Florida của Mỹ có thể biến mất khỏi bản đồ. Đường đồ thị bắt đầu thoải rồi dốc đứng dần về phía cuối. Người trợ lý giáo sư ngồi ở hàng ghế đầu tên là Alán Aspuru-Guzik. Khoảnh khắc ấy, ông cảm thấy sống lưng lạnh toát. Ông thấy bản thân thật thảm hại khi viết các bài báo nhằm cải thiện hiệu suất pin điện thoại di động trong tay thêm 0,1%. Rời khỏi giảng đường, đường đồ thị dốc đứng ấy vẫn không hề phai mờ trước mắt ông.

Aspuru-Guzik lớn lên ở Mexico. Từ nhỏ, ông là một cậu bé tò mò đến mức đọc trọn vẹn cuốn bách khoa toàn thư hai lần. Khi còn là học sinh trung học, ông đại diện cho Mexico tham dự Kỳ thi Olympic Hóa học Quốc tế tại Na Uy. Tại đó, ông đã gặp gỡ bạn bè đồng trang lứa từ khắp nơi trên thế giới và bị cuốn hút bởi việc hóa học diễn giải các quy luật tự nhiên bằng công thức toán học. Ông chia sẻ rằng mình yêu thích việc từng nguyên tử vô hình chuyển động một cách trung thực bên trong các phương trình. Sau khi học hóa học tại Đại học Tự trị Quốc gia Mexico, ông theo học cao học tại UC Berkeley (Mỹ) và học hóa học lượng tử dưới sự hướng dẫn của giáo sư hóa học tính toán Martin Head-Gordon. Năm 2004, ông nhận bằng tiến sĩ và trong nghiên cứu tiếp theo, ông là người đầu tiên trên thế giới nghĩ ra thuật toán mô phỏng phân tử bằng máy tính lượng tử, công bố trên tạp chí Science. Đó là bài báo tính toán năng lượng trạng thái cơ bản của phân tử nước và phân tử liti hydrua thông qua mô phỏng trên máy tính lượng tử. Nhờ thành tựu đó, ông được bổ nhiệm làm trợ lý giáo sư khoa hóa học tại Đại học Harvard khi mới ở độ tuổi 20. Trong giới hóa học lý thuyết, ông đã là một gương mặt trẻ đầy triển vọng.

Sau khi nghe bài giảng của Ngài King, Aspuru-Guzik đã thay đổi hướng nghiên cứu. Năm 2006, ông khởi động Dự án Năng lượng Sạch Harvard. Bằng cách mượn sức mạnh tính toán nhàn rỗi trong khi hàng triệu tình nguyện viên trên khắp thế giới chạy trình bảo vệ màn hình máy tính, ông giao cho hệ thống thực hiện các phép tính tìm kiếm phân tử hữu cơ dùng cho pin mặt trời. Trong đêm khi mọi người đang ngủ, hàng triệu chiếc máy tính âm thầm tính toán cấu trúc electron của từng phân tử. Trong vòng 3 năm, mạng lưới này đã tính toán tính chất vật lý của hơn 3 triệu phân tử hữu cơ. Trong máy tính, đó là một thành công rực rỡ.

Thế nhưng bức tường thực sự lại xuất hiện trong phòng thí nghiệm. Khi bản thiết kế các phân tử triển vọng do máy tính chọn lọc được chuyển cho các nhà hóa học, một nhà hóa học phải mất trung bình hai đến ba tuần để tổng hợp thực tế một phân tử và đo lường các tính chất vật lý của nó. Những thất bại cũng xảy ra thường xuyên. Có những ngày thuốc thử trong lọ thủy tinh phản ứng khác với dự kiến, hoặc hoàn toàn không tạo thành tinh thể. Trong số 3 triệu ứng viên, số vật liệu được chế tạo thành linh kiện thực tế chỉ vỏn vẹn có ba. Máy tính không hề biết liệu phân tử đó có thể thực sự được tạo ra trong phòng thí nghiệm hay không. Aspuru-Guzik gọi trải nghiệm này là một thất bại thành công. Tính toán thì hoàn hảo nhưng tổng hợp thì không. Ông đã tận mắt chứng kiến con sông rộng lớn ngăn cách giữa tính toán và tổng hợp.

Cùng thời điểm đó, ông cũng nghiên cứu về pin dòng chảy cùng các đồng nghiệp là giáo sư Michael Aziz và Roy Gordon. Pin dòng chảy là thiết bị lưu trữ điện bằng cách dùng bơm tuần hoàn chất điện phân lỏng chứa trong các bồn lớn. Đây là phương thức phù hợp để lưu trữ ngày đêm nguồn điện trời sput

theo thời tiết như năng lượng mặt trời hay gió. Pin dòng chảy vanadi hiện có tuy có hiệu suất tốt, nhưng lượng vanadi trên Trái Đất quá ít nên không thể đáp ứng cho toàn bộ lưới điện thế giới. Nhóm nghiên cứu đã đưa ra ý tưởng chế tạo chất điện phân bằng phân tử quinone hữu cơ, vốn có thể chiết xuất từ cặn dầu thô hoặc gỗ thay vì vanadi. Tuy nhiên, để tìm ra một loại quinone dùng được, họ đã phải tự tay tổng hợp thủ công hơn 150 chất ứng viên, tốn mất 5 năm và 5 triệu USD. Ông nhận định rằng nếu mất 10 hay 20 năm để đưa một vật liệu pin mới ra thị trường, Trái Đất sẽ phải đối mặt với thảm họa khí hậu trước khi điều đó kịp diễn ra.

Mùa thu năm 2016, Donald Trump đắc cử trong cuộc bầu cử tổng thống Mỹ. Chính quyền mới đã rút khỏi Thỏa thuận Khí hậu Paris và cắt giảm ngân sách liên bang dành cho Mission Innovation, một liên minh nghiên cứu năng lượng sạch quốc tế mà Aspuru-Guzik làm cố vấn. Đối với một người đã quyết tâm chiến đấu chống lại khủng hoảng khí hậu như ông, quyết định này thật khó chấp nhận. Ông đã từ bỏ vị trí giáo sư chính thức ổn định tại Harvard và chuyển đến Đại học Toronto của Canada vào năm 2018. Từ một cậu bé Mexico qua Mỹ rồi định cư tại Canada, điều dẫn dắt ông không phải là quốc tịch mà là chiếc đồng hồ khí hậu. Trở thành giáo sư Đại học Toronto kiêm nhà nghiên cứu tại Viện Vector, ngay sau đó ông đã lãnh đạo Acceleration Consortium quy mô 200 triệu dollar Canada (khoảng 200 tỷ won) do chính phủ Canada và trường đại học cùng thành lập. Sân khấu để hiện thực hóa ý tưởng ấp ủ từ thời ở Harvard – lấp đầy khoảng cách giữa tính toán và tổng hợp bằng robot và trí tuệ nhân tạo – cuối cùng đã được chuẩn bị.

Mục tiêu mà Acceleration Consortium đặt ra rất rõ ràng. Nếu việc tìm kiếm một vật liệu mới thường tốn hơn 10 triệu USD và hơn 10 năm, thì họ hướng tới việc giảm con số đó xuống còn 1 triệu USD và 1 năm.

chính là như vậy. Điều đó có nghĩa là nén chi phí xuống còn một phần mười và thời gian cũng xuống một phần mười. Ông tin rằng điều đó hoàn toàn khả thi nếu tạo ra một vòng lặp khép kín: cánh tay robot thực hiện thí nghiệm, trí tuệ nhân tạo thiết kế thí nghiệm tiếp theo, và kết quả lại được phản hồi cho trí tuệ nhân tạo. Thất bại thành công trải qua ở Harvard chính là bài toán mà thực tế đã để lại cho ông.

Nếu viết sai một dòng mã, chương trình sẽ dừng lại ngay lập tức và báo lỗi. Nhưng hóa học thì khác. Máy tính không thể cho biết tại sao một phân tử lại không tổng hợp được, dù nhiều tuần đã trôi qua. Điều mà Aspuru-Guzik đã trải qua suốt gần 20 năm chính là sự im lặng này. Ông đã vượt Đại Tây Dương và băng qua biên giới để gắn thêm mắt và tay cho phòng thí nghiệm đang im lặng. Để sự cấp bách của một con người phát triển thành mạng lưới hợp tác xuyên lục địa, khởi đầu chính là người trợ lý giáo sư trẻ tuổi ở độ tuổi 20 từng ngồi ở hàng ghế đầu của giảng đường năm ấy.

Từ lý thuyết đến robot

Thất bại thành công được thấy ở phần trước là kết quả trái ngược hoàn toàn với lý lịch của ông. Bởi lẽ, một người vốn đi trước bất kỳ ai trong lĩnh vực lý thuyết như ông lại liên tục bị thời gian đánh bại ở giai đoạn biến các phân tử do tính toán chọn lọc thành vật chất thực tế. Giữa câu trả lời trong đầu và kết quả tạo ra bằng tay luôn có dòng thời gian trôi khác nhau. Ông quyết định tìm kiếm con đường rút ngắn khoảng thời gian này không phải ở lý thuyết, mà là ở robot.

Chuyển đến Đại học Toronto vào năm 2018, ông đã từng bước xây dựng các công cụ để lấp đầy khoảng cách giữa tính toán và tổng hợp. Công cụ đầu tiên là thay đổi chính cách biểu diễn phân tử bằng ký tự. Ký hiệu SMILES truyền thống chỉ đơn thuần liệt kê các liên kết nguyên tử dưới dạng chuỗi ký tự, nên nếu trí tuệ nhân tạo sắp xếp sai dù chỉ một ký tự, 9 trên 10 phân tử tạo ra sẽ bị lỗi thành các phân tử không thể tồn tại. Nhà nghiên cứu sau tiến sĩ Mario Krenn đã tạo ra một hệ thống ký hiệu mới mang tên SELFIES bằng cách nhúng chính các quy tắc hóa học vào ngữ pháp – chẳng hạn như carbon chỉ có tối đa

4 liên kết, oxy có 2 và hydro có 1. Ví von như nghệ thuật gấp giấy origami, điều này giống như việc kẻ sẵn các nếp gấp trên giấy để dù gấp thế nào đi nữa thì hình dạng con hạc vẫn luôn xuất hiện. Kết quả là, ngay cả khi trí tuệ nhân tạo xáo trộn ngẫu nhiên các ký tự, các phân tử sinh ra vẫn 100% là các cấu trúc có thể thực sự tồn tại.

Công cụ thứ hai là các thuật toán Gryffin và Chimera do học trò của ông, Florian Häse, tạo ra. Đôi khi người ta phải giải quyết cùng lúc nhiều mục tiêu xung đột nhau, chẳng hạn như tăng độ dẫn điện, giảm giá thành và duy trì độ ổn định. Nếu người làm thí nghiệm phải tự tay thử từng tổ hợp vật liệu không dứt khoát như mười loại dung môi và năm loại chất xúc tác, thời gian tiêu tốn sẽ tương ứng với số lượng tổ hợp. Gryffin tìm ra những điểm tương đồng tiềm ẩn giữa các vật liệu này để chọn ra tổ hợp tiếp theo cần thử nghiệm bằng thống kê, trong khi Chimera tính toán điểm cân bằng giữa các mục tiêu xung đột như độ dẫn điện, độ bền và chi phí để xác định thứ tự thí nghiệm. Cánh tay robot chỉ cần thực sự thử nghiệm một số lượng rất ít các tổ hợp mà tính toán này chỉ định. Công cụ thứ ba là mắt của robot. Cốc thủy tinh trong suốt phản xạ ánh sáng khiến camera không thể định vị chính xác. Cùng với các giáo sư Animesh Garg và Florian Shkurti, ông đã thu thập hình ảnh phòng thí nghiệm thực tế đăng trên Instagram và Twitter để huấn luyện mạng nơ-ron nhận dạng mang tên TransparentNet. Giờ đây, cánh tay robot có thể gấp dụng cụ thủy tinh mà không bị trượt ngay cả trên bàn thí nghiệm bừa bộn.

Khi ba công cụ này kết hợp lại, các con số đã thay đổi. Trong dự án tìm kiếm laser hữu cơ của DARPA (Cơ quan Chỉ đạo các Dự án Nghiên cứu Tiên tiến Quốc phòng Mỹ), hệ thống đã rà soát 170.000 ứng viên phân tử và chỉ trong vài tuần đã tìm ra phân tử có ngưỡng kích thích thấp hơn 50% và độ lợi cao gấp 2 lần so với vật liệu tốt nhất trước đó

của Giáo sư Chihaya Adachi. Năm 2024, họ đã hoàn thành sự hợp tác kết nối các phòng thí nghiệm của bốn quốc gia gồm Ba Lan, Canada, Mỹ và Nhật Bản thành một mạng lưới không dây duy nhất. Khi viện nghiên cứu Warsaw ở Ba Lan dự đoán phân tử ứng viên, robot ở Toronto (Canada) sẽ phân chia một lượng nhỏ, nhóm của Giáo sư Marty Burke tại Đại học Illinois (Mỹ) ghép nối chúng bằng thiết bị tổng hợp liên tục, rồi gửi sản phẩm đó qua FedEx đến Đại học Kyushu (Nhật Bản) để hoàn thiện thành linh kiện bằng thiết bị lắng đọng chân không của Giáo sư Chihaya Adachi. Phần mềm điều phối hậu cần và chênh lệch múi giờ mang tên 'OrganizeIt'. Hơn 500 loại bột bán dẫn hữu cơ – điều mà 100 nhà nghiên cứu con người dốc sức cũng không thể làm được – đã được tổng hợp chỉ trong 3 tháng bằng phương pháp này. Một người khởi đầu từ lý thuyết cuối cùng đã bước tới vị trí tự tay lắp ráp những cánh tay robot và camera.

Thành quả này đã trở thành cơ sở để chính phủ Canada tiếp thêm sức mạnh cho Acceleration Consortium. Vào tháng 3 năm 2026, Đại học Toronto đã làm lễ động thổ xây dựng tòa nhà nghiên cứu mới cho liên minh sử dụng bằng cách mở rộng Tòa nhà Lash Miller ngay giữa khuôn viên trường. Tòa nhà nơi 10.000 lò phản ứng hoạt động đồng thời dự kiến sẽ dựng xong khung xương trong năm 2026. Vào tháng 6 cùng năm, một hội nghị quốc tế do liên minh của ông phối hợp tổ chức cùng Đại học Quốc gia Singapore đã diễn ra tại Singapore, cho thấy ý tưởng về phòng thí nghiệm tự hành đang vượt ra ngoài Toronto để bước lên bàn thí nghiệm của các trường đại học châu Á. Điều này diễn ra chưa đầy một năm sau khi ông cùng đồng nghiệp là Giáo sư Eugenia Kumacheva công bố bài báo về một phòng thí nghiệm tự hành riêng biệt để tổng hợp hạt nano. Nhà lý thuyết từng giải phương trình Schrödinger ở độ tuổi 20 nay đã trở thành người chia sẻ những bức ảnh công trường xây dựng với bê tông đang đông kết như thể đó là thành quả nghiên cứu của mình. Thời kỳ đánh mất 20 năm giữa tính toán và tổng hợp đang dần khép lại trên nền tảng những công cụ mà ông đã dày công xây dựng.

Phòng thí nghiệm tự hành

Khoảng cách chênh lệch trong Dự án Năng lượng Sạch Harvard được thấy ở phần mở đầu đã để lại cho Giáo sư Aspuru-Guzik một kết luận rõ ràng. Năng lực bên trong máy tính đủ sức quét qua cả vũ trụ nhưng lại không thể nhúc nhích nổi một đầu ngón tay trước bàn thí nghiệm, đồng nghĩa với việc nếu không thay đổi tốc độ của chính thí nghiệm thì tính toán có tinh vi đến đâu cũng vô dụng. 5 năm lao động thủ công trong nghiên cứu pin dòng chảy quinone cũng củng cố thêm kết luận tương tự. Trải qua những năm tháng đó, ông đã để lại câu nói: "Không còn thời gian cho cách làm khoa học như thông thường nữa (There is no time for science as usual)". Ý của ông là khủng hoảng khí hậu sẽ ập đến trước cả khi một viên pin mới kịp lắp ráp hoàn chỉnh. Sân khấu đưa kết luận này vào thực tiễn chính là Acceleration Consortium tại Toronto. Khoản tài trợ do chính phủ liên bang Canada cấp đến từ CFREF, và việc chuyển nơi ở của một cá nhân đã làm thay đổi cả diện mạo nghiên cứu của một quốc gia.

Cấu trúc mà liên minh muốn đưa vào phòng thí nghiệm chỉ có một. Đó là cấy ghép toàn bộ cấu trúc tuần hoàn vào phòng thí nghiệm: trí tuệ nhân tạo thiết kế phân tử, robot trực tiếp tổng hợp thiết kế đó, và máy móc đo lường kết quả rồi gửi ngược lại cho trí tuệ nhân tạo. Bằng cách thay thế cánh tay robot và cảm biến vào các công đoạn mà con người đóng vai trò cầu nối làm tiêu tốn hàng tuần lễ, tốc độ máy tính phác họa phân tử và tốc độ phòng thí nghiệm thực sự tạo ra phân tử đó cuối cùng cũng có thể đồng tốc. Một phòng thí nghiệm nơi tính toán, tổng hợp và đo lường hoàn thành một vòng tuần hoàn mà không qua tay con người như vậy được gọi là 'phòng thí nghiệm tự hành (Self-Driving Laboratory)', và nền tảng nghiên cứu không người do nhiều phòng thí nghiệm như vậy hợp lại được gọi là 'Nền tảng Tăng tốc Vật liệu (Materials Acceleration Platforms, MAPs)'. Giống như ô tô tự tìm đường mà không cần người lái, phòng thí nghiệm cũng tự đặt giả thuyết và hoàn thành việc kiểm chứng mà không cần nhà hóa học. Khi thiết bị đầu tiên bắt đầu vận hành tại MATER Lab của Đại học Toronto, sự thất vọng mà Giáo sư Aspuru-Guzik từng cảm nhận trước màn hình máy tính năm 2006 cuối cùng đã được chuyển giao sang những đầu ngón tay của máy móc.

Trong những năm đầu sau khi Acceleration Consortium được thành lập, Giáo sư Aspuru-Guzik đã từ chối gọi phòng thí nghiệm này chỉ đơn thuần là cơ sở tự động hóa có gắn cánh tay robot. Điều ông muốn xác lập không phải là một cánh tay máy, mà là chính trình tự của khoa học: đặt giả thuyết, làm thí nghiệm, xem xét kết quả và điều chỉnh giả thuyết tiếp theo. Thay cho vị trí con người phải mất nhiều tháng tra cứu tài liệu và thiết kế thí nghiệm, ông đã đặt vào đó một cấu trúc nơi trí tuệ nhân tạo và robot lặp lại chu trình này nhiều lần mỗi ngày. Việc xác lập này thực sự bước vào thử thách sau khi phòng thí nghiệm mở cửa. Khoản tài trợ của chính phủ không phải là tiền để mua robot, mà là tiền được dùng để chuyển giao công việc lặp đi lặp lại mà con người không kham nổi cho máy móc và đưa nhà hóa học sang vị trí phán đoán

được sử dụng vào mục đích đó.

Bảng thành tích của chu trình tuần hoàn này đã được đăng trên tạp chí khoa học Science vào tháng 5 năm 2024. Năm phòng thí nghiệm tại Toronto, Vancouver, Glasgow, Illinois và Fukuoka đã được liên kết thành một vòng lặp khép kín, tổng hợp và thử nghiệm hơn 1.000 loại vật liệu độ lợi dùng cho laser trạng thái rắn hữu cơ. Phương thức hoạt động là mô hình thống kê dự đoán trước hiệu suất của các chất chưa từng được tổng hợp, chọn ra các ứng viên có xác suất thành công cao nhưng chưa từng được các nhà hóa học con người chạm tới, rồi phân chia cho robot ở năm thành phố đảm nhiệm. Cuộc tìm kiếm mà một nhóm nghiên cứu con người phải mất nhiều năm đã kết thúc chỉ trong vài tháng, và 21 ứng viên hàng đầu đã được chọn lọc.

Phòng thí nghiệm này vẫn đang tiếp tục phát triển. Tạp chí học thuật quốc tế Nature trong một bài viết đặc biệt vào tháng 3 năm 2026 đã đưa tin rằng đội ngũ nghiên cứu của Giáo sư Aspuru-Guzik đang vận hành đồng thời 50 robot trải rộng khắp nhiều trường đại học và viện nghiên cứu. Cuộc tìm kiếm từng do một chiếc máy tính và vài nhà hóa học gánh vác, nay được 50 robot vận hành ngày đêm quanh các bàn

thí nghiệm đảm trách. Tạp chí Chemical & Engineering News của Hội Hóa học Hoa Kỳ trong bài báo tháng 6 năm 2026 cũng chỉ ra rằng phòng thí nghiệm tự hành đang thay đổi chính cách thức làm việc của các nhà hóa học. Nỗi thất vọng bắt đầu trước màn hình máy tính tại Harvard 20 năm trước, nay đang dần định hình thành một tòa nhà sừng sững trong khuôn viên Toronto. Từ ngữ 'phòng thí nghiệm' giờ đây không còn chỉ một căn phòng được bao quanh bởi bốn bức tường và cửa sổ nữa, mà đã trở thành từ chỉ một chu trình tự vận hành suốt đêm mà không cần đến con người.

Vòng lặp khép kín

Ban đêm ở MATER Lab tại Toronto không hề yên tĩnh. Tiếng cánh tay robot bắn từng giọt chất lỏng lên đĩa vang lên đều đặn, và các đồ thị trên màn hình tự động cập nhật ngay cả lúc rạng sáng. Mười năm trước, phòng thí nghiệm khoa hóa học của Harvard cách rất xa những đêm như thế này. Trong nghiên cứu pin dòng chảy quinone đã đề cập trước đó, các sinh viên phải tự tay dùng ống hút pipette để đun nóng, làm nguội và rửa sạch từng chất ứng viên một. Đó là thời kỳ mà hóa học được thúc đẩy bằng những đầu ngón tay. Như kết cục của Dự án Năng lượng Sạch đã cho thấy, máy tính dự đoán chính xác tính chất của phân tử nhưng liệu phân tử đó có thể thực sự được lắp ráp bằng thuốc thử hay không lại không nằm trong các phép tính. Giữa bit và nguyên tử dường như không có cây cầu nối. Cũng giống như việc máy tính dù tưởng tượng chính xác hương vị đến đâu, nếu không có gian bếp để thực sự nấu món ăn đó thì cũng chẳng thể dọn lên bàn ăn. Khi định cư tại Toronto, ông đã hạ quyết tâm nối liền các dự đoán trong máy tính với việc tổng hợp thực tế trong phòng thí nghiệm thành một mạch duy nhất.

Phòng thí nghiệm tự hành do Acceleration Consortium tạo ra vận hành qua bốn giai đoạn. Ở giai đoạn đầu tiên, các phân tử được thiết kế bằng ký hiệu SELFIES đã nêu trước đó. Ở giai đoạn thứ hai, tính chất của phân tử được xem xét trước thông qua các phép tính cơ học lượng tử. Nếu tính toán một phân tử bằng Lý thuyết phiếm hàm mật độ (DFT) mất 13 giờ, thì sử dụng phương pháp liên kết chặt bán thực nghiệm chỉ mất 15 phút. Giai đoạn thứ ba là tổng hợp. Hệ điều hành chuyên dụng cho hóa học mang tên ChemOS do nhóm nghiên cứu phát triển sẽ chuyển đổi quy trình thí nghiệm viết bằng ngôn ngữ tự nhiên thành các câu lệnh mà cánh tay robot có thể hiểu được, rồi robot Chemspeed của Thụy Sĩ và robot Medusa tự phát triển sẽ nhận các lệnh này để thực sự trộn thuốc thử. Giai đoạn thứ tư là phân tích. Sóng siêu âm bắn chính xác các giọt chất lỏng kích thước 2,5 nanolit, và sắc ký lỏng cùng khối phổ kế sẽ xác định danh tính của chất. Các giá trị đo được bằng tia cực tím, huỳnh quang và phương pháp quét thể vòng được đưa vào ba mô hình học Bayes mang tên Gryffin, Phoebus và Chimera để định hình lại hướng đi cho thí nghiệm tiếp theo. Nếu dự đoán và thực nghiệm sai lệch, máy móc sẽ tự động sửa đổi giả thuyết.

Chu trình này đã dẫn đến những thành quả thực tế. Kết quả hợp tác quốc tế về laser hữu cơ nói trên đã được đăng trên tạp chí Science vào năm 2024. Sáu nhóm nghiên cứu từ ba châu lục đã hợp tác để tìm ra 21 vật liệu laser mới. Đó là thành quả đạt được trong một lĩnh vực mà con người phải mất nhiều thập kỷ mới tìm ra được khoảng mười loại. Người học trò Andrés đã xây dựng một phòng thí nghiệm tự hành trị giá 300 USD tại Guadalajara (Mexico) chỉ bằng cánh tay robot in 3D, máy bơm Arduino giá 100 USD và một webcam. Giáo sư Aspuru-Guzik đã phát hành miễn phí toàn bộ phần mềm trên GitHub. Ông cho rằng không có lý do gì chỉ những trường đại học giàu có mới được sở hữu phòng thí nghiệm tự hành. Bức tranh mà ông vẽ ra không hề bị giam hãm sau những bức tường của các trường đại học danh tiếng.

Ông nói rằng chỉ cần có kết nối Internet, sinh viên đại học ở vùng nông thôn Nam Mỹ hay nghiên cứu sinh ở châu Phi cũng phải có thể tự thiết kế vật liệu lọc để làm sạch nước giếng bị ô nhiễm và kiểm chứng bằng robot.

Nhìn lại lịch sử khoa học, xu hướng này là một làn sóng mới. Từng có thời kỳ con người quan sát tự nhiên bằng mắt thường. Tiếp nối là thời kỳ thiết lập các định luật bằng công thức toán học như Newton.

Tiếp theo là thời kỳ mô phỏng hiện tượng bằng siêu máy tính, rồi đến thời kỳ tìm kiếm mối tương quan từ lượng dữ liệu khổng lồ. Giáo sư Aspuru-Guzik gọi xu hướng hiện nay, nơi máy móc tự đặt giả thuyết và hoàn thành thí nghiệm bằng robot, là làn sóng thứ năm. Ông dự đoán rằng khi các công ty như Zapata Computing do ông đồng sáng lập hoàn thiện máy tính lượng tử tự sửa lỗi, người ta sẽ có thể dự đoán không sai số ngay cả hành vi phức tạp của phân tử mà siêu máy tính hiện nay khó lòng đảm đương. Đó là điều sẽ diễn ra trong vòng 10 đến 15 năm tới. Nhà triết học khoa học Thomas Kuhn từng viết rằng khi một hệ hình thay đổi, dù thế giới vật lý vẫn giữ nguyên thì thế giới mà các nhà khoa học đang sống đã hoàn toàn khác biệt. Thế giới đun sôi thử công các ứng viên quinone và phòng thí nghiệm Toronto ngày nay nơi máy móc tự sửa đổi giả thuyết suốt đêm tuy cùng là hóa học, nhưng lại gần như thuộc về hai thế giới khác nhau.

Vào tháng 6 năm 2026, Acceleration Consortium đã bắt tay với Structural Genomics Consortium. Hai cơ quan quyết định đưa phương thức phòng thí nghiệm tự hành vào toàn bộ quy trình thiết kế, tổng hợp và thử nghiệm các chất ứng viên thuốc mới. Số vốn huy động được từ ngành dược phẩm lên tới 50 triệu USD. Vòng lặp khép kín từng chỉ tìm kiếm một loại vật liệu nay đang mở rộng bước tiến sang các ứng viên thuốc điều trị bệnh cho con người. Chưa đầy mười năm để sự mệt mỏi của những sinh viên từng đun quinone thâu đêm trong phòng thí nghiệm Harvard chuyển thành tiếng ping đều đặn của cánh tay robot.

Tăng tốc tìm kiếm vật liệu

Khung cảnh nghiên cứu sinh cầm pipet đứng dưới tủ hút canh chừng dung dịch hữu cơ đang sôi là bức chân dung của thời kỳ thao tác thủ công với quinone như đã thấy ở trước. Trong thời đại máy tính có thể tính toán mức năng lượng của hàng triệu phân tử chỉ trong vài giờ, bàn tay con người đã không thể theo kịp tốc độ tính toán đó.

Giáo sư Aspuru-Guzik đã tiếp tục công việc kết hợp thuật toán mô phỏng lượng tử mà ông nghĩ ra thời làm nghiên cứu sinh sau tiến sĩ với học máy lượng tử (Quantum Machine Learning, QML) để ứng dụng vào mô phỏng phân tử trong suốt gần 20 năm.

Nghiên cứu về pin dòng thời ở Harvard cũng đã từng gặt hái thành quả. Năm 2014, nhóm nghiên cứu của ông đã khiến giới học thuật kinh ngạc khi công bố trên tạp chí Nature về loại pin dòng dung dịch nước không chứa kim loại sử dụng axit anthraquinone-2,7-disulfonic (AQDS). Vấn đề nằm ở giai đoạn tiếp theo. Để kéo dài tuổi thọ của chất điện phân này, họ phải tự tay tổng hợp và thử nghiệm từng phân tử ứng viên một.

Sau khi chuyển đến Toronto, mục tiêu mà ông nhắm tới là anolyte – các phân tử hữu cơ di chuyển trong chất điện phân dưới dạng ion. Tại đây, hóa học tính toán lại vấp phải một bức tường. Nếu tính cả sự can thiệp tĩnh điện do các phân tử nước và ion đối tạo ra, các phép tính theo lý thuyết phiếm hàm mật độ đòi hỏi chi phí tính toán mà ngay cả siêu máy tính cũng khó lòng gánh vác. Để vượt qua bức tường này, Giáo sư Aspuru-Guzik cùng với Zapata Computing – công ty lượng tử do ông đồng sáng lập – đã thiết kế Bộ giải trị riêng lượng tử tạo sinh (Generative Quantum Eigensolver, GQE). Nếu như thuật toán lượng tử truyền thống VQE (Variational Quantum Eigensolver) chỉ tính toán năng lượng trạng thái cơ bản của một phân tử duy nhất, thì GQE dự đoán toàn bộ phân bố xác suất của cả đám mây electron.

Việc kết nối cỗ máy tính toán này với cánh tay robot thực tế cũng là điều cần thiết. Đầu tiên, thuật toán di truyền Janus kết hợp học sâu và tìm kiếm tiến hóa để tự mình phác thảo nên các cấu trúc phân tử ứng viên. Tiếp theo, Gryffin và Chimera như đã thấy ở phần trước sẽ sàng lọc các ứng viên thực sự cần tổng hợp xuống còn 6 đến 10% tổng số. Cuối cùng, cánh tay robot sẽ phân phối chất lỏng ở đơn vị siêu vi lượng, thiết bị đo điện hóa 16 kênh đo đặc kết quả bằng phương pháp đo von-ampe vòng, rồi gửi giá trị đó quay trở lại máy tính lượng tử. Con người chỉ cần đặt ra mục tiêu ban đầu, sau đó máy móc sẽ vận

hành vòng tuần hoàn này 24 giờ mỗi ngày.

Không thể không nhắc đến câu chuyện về thiết bị đo lường này. Các máy đo von-ampe trên thị trường có giá tới hàng chục triệu won và chỉ có hai kênh. Nhóm nghiên cứu của Aspuru-Guzik chỉ dùng các linh kiện trị giá 100 USD cho 16 kênh

đo đồng thời, tự tay thiết kế bảng mạch Pocket PowerUp và công khai toàn bộ bản vẽ. Họ hiểu rõ rằng một thiết bị giá rẻ quyết định tốc độ xử lý của cả phòng thí nghiệm.

Giáo sư Aspuru-Guzik ví sự thay đổi này với ngành thiên văn học. Các nhà thiên văn học ngày xưa mỗi đêm phải leo lên đỉnh núi, ghé mắt vào kính thiên văn và dùng tay ghi chép vị trí của các vì sao. Các nhà thiên văn học ngày nay ngồi tại bàn làm việc trong văn phòng, dùng mã lập trình để phân tích dữ liệu do kính viễn vọng không gian gửi về nhằm tìm kiếm lỗ đen và siêu tân tinh. Ông cho rằng các nhà hóa học cũng đang đi trên con đường tương tự. Điều đó có nghĩa là đôi bàn tay từng cân đong bột xúc tác trước tủ hút nay đang được chuyển giao cho các cánh tay robot.

Kết quả mà hệ thống này mang lại được thể hiện rõ ràng qua các con số. Phòng thí nghiệm tự hành đã tổng hợp thực tế 150 loại phân tử hữu cơ ion mới được thiết kế mà không gặp phải bất kỳ một thất bại nào. Sai số trung bình giữa điện thế dự đoán bằng tính toán lượng tử và giá trị đo thực tế là dưới 0,15 V. Trong số đó, vật liệu xuất sắc nhất ghi nhận mức điện thế thấp tới âm 0,85 V so với điện cực hydro tiêu chuẩn, vượt trội hơn hẳn dòng AQDS hiện có. Khi lặp lại 100 chu kỳ nạp xả cùng một vật liệu trong nước, dung lượng vẫn được giữ nguyên 99,9%. Cấu trúc phân tử mà Janus tìm ra chỉ trong ba ngày tương đương với kết quả mà các nghiên cứu sinh tiến sĩ con người phải mất 5 năm thử và sai mới có được. Máy móc không hề bắt chước trực giác của con người, mà chỉ đơn giản là đã đi tới cùng một đáp án nhanh hơn rất nhiều.

Xu hướng này vẫn tiếp tục trong tháng 8 năm 2026. Ngày 12 tháng 8, Zapata Quantum – công ty do Giáo sư Aspuru-Guzik đồng sáng lập – đã hợp tác với QuEra, một công ty máy tính lượng tử nguyên tử trung hòa, để bắt tay vào phát triển các ứng dụng máy tính lượng tử chịu lỗi với mục tiêu thương mại hóa vào năm 2028. Thiết kế vật liệu và tính toán hóa học được coi là mục tiêu cốt lõi của sự hợp tác này. Cùng năm đó, trang tin công nghệ Canada BetaKit đã đưa Giáo sư Aspuru-Guzik vào danh sách nhân vật 'Most Ambitious 2026' và giới thiệu ông là người đang tự động hóa chính quá trình thí nghiệm bằng robot và trí tuệ nhân tạo. Vị trí của các nghiên cứu sinh từng phải đứng suốt 5 năm dưới tủ hút giờ đây đã được thay thế bởi máy tính lượng tử và cánh tay robot.

Những vấn đề chưa có lời giải

Những kết quả mà phòng thí nghiệm tự hành đưa ra thường phải đối mặt với thử thách trong khâu kiểm chứng độc lập bên ngoài. Năm 2023, khi nhóm nghiên cứu A-Lab thuộc Phòng thí nghiệm Quốc gia Lawrence Berkeley của Mỹ công bố trên tạp chí quốc tế Nature rằng họ đã tự động tổng hợp được hàng chục hợp chất vô cơ mới chỉ trong vài ngày, giới học thuật đã lập tức tiến hành xem xét lại. Các nhà nghiên cứu, trong đó có nhà tinh thể học Robert Palgrave thuộc University College London, đã phân tích lại dữ liệu nhiễu xạ tia X trong bài báo và chỉ ra rằng phần lớn các trường hợp được báo cáo là chất mới thực chất dường như là những chất đã biết hoặc hỗn hợp gồm nhiều pha lẫn lộn. Đây là phản biện cho thấy việc máy móc xác định đã tổng hợp thành công và việc con người xác nhận xem chất đó có thực sự là một cấu trúc tinh thể mới hay không là hai vấn đề hoàn toàn khác nhau. Có những ý kiến cho rằng, tự động hóa càng đẩy nhanh tốc độ thí nghiệm thì quy trình con người phải lật lại kiểm chứng các kết quả đó lại càng trở nên nặng nề hơn.

Câu hỏi thứ hai là không gian hóa học mà máy móc tìm kiếm rốt cuộc vẫn bị giam hãm trong khuôn khổ do con người vạch sẵn từ trước. Cú pháp SELFIES, danh sách ứng viên và hàm mục tiêu tối ưu hóa đều là những hàng rào do các nhà nghiên cứu thiết kế. Nhiều bài báo tổng quan về phòng thí nghiệm tự

hành chỉ ra rằng robot chỉ quét qua bên trong hàng rào đó nhanh hơn con người, chứ không thể tự mình vượt qua để chạm tới những phát hiện bất ngờ có thể nằm ngoài hàng rào. Những phát hiện tình cờ ngoài dự tính vốn là động lực thay đổi lịch sử khoa học, nhưng có mối lo ngại rằng một vòng lặp khép kín được thiết lập mục tiêu quá hẹp có thể sẽ loại bỏ chính những sự ngẫu nhiên quý giá đó.

Câu hỏi thứ ba là khả năng tái lập và chất lượng của dữ liệu do robot tạo ra hàng loạt. Sai số cực nhỏ của thiết bị thí nghiệm, tình trạng thuốc thử hay sự dao động của điều kiện đo đạc luôn là vấn đề ngay cả trong những phòng thí nghiệm do con người trực tiếp vận hành; khi vòng tuần hoàn tiếp tục chạy suốt đêm lúc con người vắng mặt, những sai số này có thể âm thầm tích tụ. Lượng dữ liệu tăng lên không đồng nghĩa với việc chất lượng sẽ tự động đi kèm, và khuyến cáo chung của giới học thuật là các giá trị đo tự động vẫn cần có sự kiểm chứng và chuẩn hóa của con người. Để phòng thí nghiệm tự hành giành được lòng tin, yêu cầu chứng minh kết quả đó có thể tái lập ở các phòng thí nghiệm khác đang ngày càng lớn hơn so với việc phô diễn số lượng kết quả.

Nội dung tường thuật trong chương này dựa trên các tư liệu được công bố tính đến tháng 8 năm 2026.

Ghi chú của tác giả. Vai trò của con người đã chuyển dịch về đâu?

Hai nhân vật trong Phần 2 đã giao phó việc thiết kế cho máy móc. Baker để máy móc phác thảo những protein chưa từng tồn tại trong tự nhiên, còn Aspuru-Guzik để máy móc vẽ nên những vật liệu mà tự nhiên chưa từng tạo ra. Nếu dự đoán là việc tìm ra bản chất của những thứ đã tồn tại, thì thiết kế là việc xác định trước các đặc tính mong muốn rồi nhào nặn ngược lại để tạo ra vật chất sở hữu những đặc tính đó. Chiều hướng tiếp cận đã hoàn toàn đảo ngược.

Sự đảo ngược này đã tái định vị vai trò của con người. Khi máy móc có thể thiết kế bất cứ thứ gì, câu hỏi nên để máy móc thiết kế cái gì đã được đặt lên hàng đầu. Nhắm vào căn bệnh nào, tìm kiếm loại vật liệu nào trước tiên – những phán đoán đó không nằm ở phép tính do máy móc đưa ra, mà thể hiện điều con người trân trọng. Chừng nào thước đo của thiết kế không phải là đúng hay sai mà là khả năng hoạt động thực tế, thì việc đánh giá xem điều gì đáng để đưa vào hoạt động vẫn luôn là phần việc của con người.

Với tư cách là một luật sư, tôi muốn nói thêm rằng máy móc biết thiết kế vừa là một phước lành vừa là một mối hiểm họa. Đôi bàn tay phác thảo protein chữa bệnh cũng có thể vẽ ra những chất gây hại. Câu hỏi nên để máy móc thiết kế cái gì luôn đi liền với câu hỏi phải ngăn chặn máy móc thiết kế những gì. Việc vạch ra ranh giới này cũng chính là vị trí còn lại dành cho con người, chứ không phải máy móc.

Phần 3

Cỗ máy tìm ra quy luật

Chương 6. Hod Lipson: Đảo ngược kỹ thuật Newton

Từ dữ liệu đến quy luật

Mùa hè năm 2000, trên một bàn làm việc tại Đại học Brandeis ở bang Massachusetts, Mỹ, xuất hiện một vật thể kỳ lạ. Đó là một khối liên kết gồm vài thanh nhựa và động cơ nối với nhau theo hình dạng vụn vẹo. Không có ai thiết kế vật thể này. Nó là kết quả của quá trình tiến hóa bên trong máy tính qua nhiều thế hệ bằng cách tự biến đổi hình dạng cơ thể, cho đến khi cá thể duy nhất có khả năng bước đi tốt hơn các cá thể khác sống sót đến cuối cùng và được tạo thành vật thể thực nhờ máy in 3D. Con người chỉ làm mỗi việc là đặt ra mục tiêu. Họ giao cho máy tính một điều kiện duy nhất là tiến về phía trước mà không bị ngã, phần còn lại để thuật toán di truyền tự xử lý. Con người không hề chỉ dạy dù chỉ một dòng về việc nó phải có bao nhiêu chân hay các khớp phải gắn ở đâu để thuận lợi cho việc đi lại. Người dẫn dắt thí nghiệm này là nhà nghiên cứu Hod Lipson. Thứ mà ông mở ra và quan sát mỗi ngày trên màn hình máy tính là biểu đồ cho thấy các thân thể ứng viên dần thay đổi hình dạng qua từng thế hệ.

Lipson và Giáo sư Jordan Pollack đã công bố kết quả này trên tạp chí học thuật Nature vào tháng 8 năm đó. Tiêu đề của bài báo là Thiết kế và chế tạo tự động các dạng sống robot. Cho đến thời điểm đó, chưa từng có tiền lệ nào về việc máy tính tự mình nghĩ ra bản thiết kế rồi chuyển thẳng cho máy móc in ra vật thể thực. Thí nghiệm của Lipson đã đảo ngược quy trình kỹ thuật truyền thống, nơi con người tính toán nguyên lý trước rồi máy móc mới làm theo tính toán đó.

Thành công này đã gieo vào tâm trí Lipson một câu hỏi khác. Liệu một cỗ máy có thể tự thiết kế hình dạng cơ thể mình thì có thể tự khám phá ra các quy luật tự nhiên hay không. Chuyển sang Đại học Cornell vào năm 2001, Lipson mở một phòng thí nghiệm mới và tiếp tục các thí nghiệm theo đuổi câu hỏi này. Nghiên cứu về việc robot tự phán đoán cấu trúc cơ thể mình để học lại cách đi cũng xuất hiện vào khoảng thời gian này. Tuy nhiên, ánh mắt của Lipson dần vượt qua giới hạn thân thể robot để hướng về chính các hiện tượng tự nhiên.

Vào thời điểm đó, cách khoa học tiếp cận tự nhiên được chia thành hai nhánh lớn. Một nhánh là con người quan sát trong thời gian dài và tính toán bằng tay để lập phương trình. Nhánh còn lại là nạp thật nhiều dữ liệu cho trí tuệ nhân tạo để huấn luyện nó đoán các con số. Phương pháp thứ nhất trung thực nhưng chậm chạp, phương pháp thứ hai nhanh chóng nhưng không cho biết lý do vì sao đúng. Con người khi học vật lý tiếp nhận các định luật do Newton lập ra qua sách giáo khoa, nhưng thứ máy tính nắm trong tay chỉ là tọa độ và thời gian của vài điểm dao động. Con đường thứ ba rút ra quy luật từ hai bàn tay trắng này và hoàn trả lại bằng thứ ngôn ngữ phương trình hữu hình, chính là điều nằm giữa hai phương pháp mà Lipson đã nhìn thấy:

khoảng trống.

Năm 2007, một nghiên cứu sinh tiến sĩ tên là Michael Schmidt gia nhập phòng thí nghiệm của Lipson. Schmidt từng theo học ngành sinh học tính toán và làm việc với các mô hình dự đoán như mạng nơ-ron và máy vectơ hỗ trợ. Anh cảm thấy bức bối khi các mô hình này dù đoán số rất chuẩn nhưng lại là những chiếc hộp đen không thể giải thích nổi một lời về lý do. Lipson và Schmidt đứng trước bảng đen đến tận đêm muộn, cùng nhau lên kế hoạch: chỉ giao các con số về vị trí và vận tốc của vật thể dao động cho một máy tính không hề có chút kiến thức vật lý nào, rồi khiến nó tự mình tìm ra các công thức vật lý thực sự như định luật bảo toàn năng lượng. Ý tưởng của họ là tạo ra một người đồng nghiệp biết viết phương trình ra giấy chứ không phải một chiếc hộp chỉ biết đoán số. Câu chuyện về việc kế hoạch này sau đó khai sinh ra cỗ máy mang tên Eureqa sẽ được tiếp tục ở phần sau của chương này.

Nỗ lực tìm kiếm công thức vật lý chỉ từ dữ liệu thô mà không cần kiến thức tiên nghiệm của con người không dừng lại bên trong phòng thí nghiệm của hai người. Tháng 5 vừa qua, hai nhà nghiên cứu Teddy

Lazebnik và Alex Liberzon đã công bố trên tạp chí học thuật Scientific Reports phương pháp suy luận các định luật vật lý không chỉ từ dữ liệu dạng bảng mà còn từ dữ liệu cấu trúc đồ thị. Vào tháng 7, nhóm nghiên cứu tại Đại học Clarkson ở Mỹ đã giới thiệu một công cụ mang tên KANDy sử dụng mạng nơ-ron Kolmogorov-Arnold. Công cụ này đã phục hồi các phương trình của hệ hỗn loạn dao động phức tạp cũng như cấu trúc phân thớ Hopf trong tô pô học thành quy trình tính toán mà con người có thể theo dõi bằng mắt thường, và được công khai trên GitHub để bất kỳ ai cũng có thể tải về sử dụng. Câu hỏi nhiều năm trước của Lipson đã vượt qua khoảng thời gian hơn hai mươi năm và hiện vẫn còn lưu lại trên bảng đen của nhiều phòng thí nghiệm.

Robot tự thấu hiểu bản thân

Trên bàn thí nghiệm của Phòng thí nghiệm Creative Machines tại Đại học Columbia, một robot bốn chân giống hình con sao biển đang bò trên sàn trải thảm. Một nhà nghiên cứu cầm chiếc kim nhỏ trên tay tiến lại gần và cắt đứt phăng chân trước bên phải của robot. Toàn bộ dây dẫn nối tới động cơ cũng bị cắt đứt hoàn toàn. Robot không có cảm biến nào để báo rằng chân đã bị cắt, cũng không có mã lệnh ứng phó khẩn cấp. Dù vậy, robot không hề dừng lại. Một ngày trôi qua, robot đã tự cân bằng lại cơ thể chỉ với ba chân và bắt đầu khập khiễng tiến về phía trước.

Cảnh tượng này là thí nghiệm xuất hiện trong bài báo 'Cỗ máy có khả năng tự phục hồi thông qua tự mô hình hóa' được Giáo sư Hod Lipson cùng các học trò Josh Bongard và Viktor Zykov công bố trên tạp chí học thuật quốc tế Science vào năm 2006. Sau khi nhận bằng cử nhân và tiến sĩ kỹ thuật cơ khí tại Technion ở Israel, Giáo sư Lipson chuyển đến Mỹ và làm nghiên cứu sinh sau tiến sĩ tại MIT và Đại học Brandeis. Thí nghiệm robot tiến hóa xuất hiện ở đầu chương này chính là thành quả của giai đoạn đó. Sau khi làm giáo sư tại Đại học Cornell rồi chuyển sang Đại học Columbia, ông đã mở rộng ý tưởng nhào nặn cơ thể qua tiến hóa sang năng lực tự nhận thức của robot.

Phép so sánh mà Giáo sư Lipson ưa thích là một đứa trẻ sơ sinh. Trẻ sơ sinh sinh ra không mang theo bản thiết kế cho biết cánh tay mình dài bao nhiêu xentimét hay có bao nhiêu khớp. Thay vào đó, đứa trẻ liên tục đối chiếu cử động ngo nguậy của ngón tay với kết quả nhìn thấy bằng mắt để vẽ nên hình dạng cơ thể mình trong não bộ. Bản đồ cơ thể do não tự vẽ ra như vậy được gọi là sơ đồ cơ thể. Robot sao biển của Giáo sư Lipson cũng học theo đúng cách này. Không hề biết cơ thể mình có bao nhiêu chân, nó lặp lại các chuyển động lắt nhắt ngẫu nhiên tám động cơ trong suốt hai ngày, và chỉ từ kết quả đó nó đã tự phát hiện ra mình có bốn chân. Đến ngày thứ tư, nó đoán đúng cách từng chân gắn vào thân với độ chính xác hơn 90 phần trăm.

Ngay cả vào khoảnh khắc chiếc chân bị cắt đứt, robot cũng không thể dùng mắt để xác nhận điều đó. Nó chỉ cảm nhận được rằng độ nghiêng của cơ thể khi ra lệnh cho động cơ bị lệch đi rất nhiều so với dự đoán của chính mình. Robot lấy độ lệch này làm tín hiệu để xây dựng lại bộ mô phỏng bên trong cơ thể. Sau vài lần thử nghiệm ngắn với robot thực, nó chấp nhận hình dạng cơ thể mới thiếu một chân và tự tạo ra cách bước đi.

Sự phục hồi nhanh chóng này không đơn thuần là một sự may mắn. Robot của Giáo sư Lipson nuôi dưỡng đồng thời nhiều bộ mô phỏng mang các giả thuyết khác nhau bên trong cơ thể, rồi chọn ra những động tác mà kết quả giữa các giả thuyết đó mâu thuẫn lớn nhất để thử nghiệm trên thực tế. Bởi vì nếu tìm ra một động tác mà giả thuyết 'mình giống con rắn' và giả thuyết 'mình giống con nhện' dự đoán ra kết quả hoàn toàn trái ngược nhau, thì chỉ cần một động tác đó là có thể loại bỏ hoàn toàn một trong hai giả thuyết. Điều này giống như việc chuyển logic phản nghiệm của triết gia Karl Popper vào thân thể robot. Bằng phương pháp này, Giáo sư Lipson đã chứng minh rằng robot có thể nhận biết gần như hoàn hảo hình dạng cơ thể mình chỉ sau vài lần thử nghiệm.

Nhóm nghiên cứu của Giáo sư Lipson đã cho robot này chạy lặp lại ba mươi lần và xác nhận thêm một sự thật thú vị. Bộ mô phỏng trong đầu robot luôn tin rằng nó có thể đi xa gấp đôi khoảng cách thực tế. Ảo tưởng này không phải là một lỗi tính toán. Giáo sư Lipson giải thích rằng kết quả này giống như hình ảnh một nghiên cứu sinh sau đại học khi tự tin thái quá về năng lực của mình thì lại dấn thân vào những thử thách lớn hơn. Nếu hạ thấp bản đồ cơ thể bên trong robot cho khớp chính xác với lực ma sát thực tế, nó thậm chí không tìm nổi hướng tiến lên phía trước giữa các nhiễu loạn và đứng khựng lại tại chỗ. Chính sự tự tin được thổi phồng đôi chút lại giúp robot tiếp tục di chuyển.

Nghiên cứu của Giáo sư Lipson không dừng lại ở đó. Vào các năm 2019 và 2022, ông đã công bố trên tạp chí Science Robotics cánh tay robot có khả năng tự học toàn bộ hình dạng cơ thể mà không cần bản thiết kế tính toán trước các góc khớp, chỉ nhờ một chiếc máy ảnh giá rẻ và một điểm chấm gắn ở đầu cánh tay. Chỉ sau tám giờ học tập, robot này đã xác định vị trí đầu ngón tay với sai số dưới 1 xentimét, tương đương với mức độ khi một người nhắm mắt chạm vào đầu mũi của chính mình.

Cánh tay robot này không vẽ cơ thể mình như một bản vẽ rõ nét mà như một đám mây xác suất mờ ảo như sương mù. Đó là phương thức chọn một điểm trong không gian rồi tính toán xác suất cơ thể mình hiện diện tại vị trí đó. Giáo sư Lipson cho rằng hình ảnh tự thân của robot có độ mờ nhạt như vậy lại tự nhiên hơn. Robot chỉ cần đi theo tâm của đám mây sương mù này để lướt qua và tránh né các chướng ngại vật xa lạ.

Phòng thí nghiệm Columbia đã mở rộng bản đồ cơ thể này sang cả việc đọc suy nghĩ của các robot khác. Trong thí nghiệm trốn tìm, robot đi tìm và robot đi trốn không hề trao đổi vị trí qua sóng không dây. Robot đi trốn chỉ dùng máy ảnh quan sát hướng đầu của robot đi tìm, rồi tự suy luận ra rằng ở góc độ này thì robot đi tìm không thể thấy nó ở sau cột gỗ. Đó là khoảnh khắc mà năng lực tưởng tượng về cơ thể mình đã phát triển thành năng lực nhìn thế giới qua đôi mắt của kẻ khác.

Demis Hassabis, người đứng đầu Google DeepMind, nhìn thấy một con đường khác. Bởi vì ông tin rằng máy móc có thể tự học các quy luật vật lý của thế giới chỉ bằng cách xem hàng triệu video trên YouTube. Giáo sư Lipson lắc đầu trước suy nghĩ này. Ông cho rằng việc xem video trên màn hình chỉ là sự bắt chước trông có vẻ hợp lý, và robot phải trực tiếp xoay tay nắm, trải nghiệm cảm giác cơ thể bị nghiêng thì các quy luật vật lý thực sự ở cấp độ nguyên tử mới khắc sâu vào thân thể nó. Cuộc tranh luận giữa hai người vẫn chưa ngã ngũ. Tuy nhiên, cả hai đều có chung linh cảm rằng tri thức học được từ việc xem video và tri thức học được từ va chạm thể xác rồi sẽ có ngày hội tụ bên trong một robot.

Xu hướng này vẫn đang tiếp diễn vào năm 2026. Vào tháng 6 vừa qua, Phòng thí nghiệm Creative Machines tại Columbia đã công bố một bài báo thí nghiệm về việc robot đối chiếu tín hiệu cảm giác của cơ thể mình với hình ảnh nhìn thấy bằng mắt để phân biệt giữa bản thân trong gương và robot khác bên cạnh. Bản đồ cơ thể gieo vào robot sơ sinh giờ đây đã phát triển thành cảm giác phân biệt giữa bản thân và người khác. Đầu năm đó, một nghiên cứu về quá trình trao đổi chất của robot cũng được công bố, trong đó robot tháo các linh kiện bị hỏng từ robot khác để gắn vào cơ thể mình. Mỗi khi cơ thể thay đổi, robot lại phải học lại về chính mình một lần nữa. Những robot của Giáo sư Lipson ngày nay vẫn đang xoay các núm vặn trong bóng tối để tìm hiểu xem mình là ai.

Eureqa

Mùa thu năm 2007, một bức email đã gửi đến một phòng nghiên cứu tại Đại học Cornell ở Mỹ. Người gửi là các nhà vật lý hạt nhân thuộc Tổ chức Nghiên cứu Hạt nhân Châu Âu (CERN). Họ đã đo 3.000 giá trị năng lượng sinh ra khi hạt nhân nguyên tử liên kết proton và neutron rồi gửi đến phòng nghiên cứu này. Người nhận là Giáo sư Hod Lipson và học trò nghiên cứu sinh tiến sĩ Michael Schmidt. Cả hai chưa từng học bài bản về vật lý hạt nhân. Dù vậy, họ không hề chùn bước. Họ nạp thẳng dữ liệu vào chương trình Eureqa do chính mình tạo ra rồi nhấn nút tính toán. 32 chiếc máy tính đã chạy suốt một ngày một

đêm. Sáng hôm sau, một công thức ngắn gọn và thanh lịch xuất hiện trên màn hình. Đó là câu trả lời được rút ra mà không cần bất kỳ lý thuyết nào, chỉ dựa vào 3.000 con số. Giáo sư Lipson gửi nguyên công thức này qua email cho CERN. Thư trả lời đến chỉ sau 5 giây. Nhóm nghiên cứu của CERN viết rằng: cỗ máy của các bạn chỉ trong một ngày đã tự mình tìm lại công thức khối lượng Weizsäcker, một di sản lâu đời của vật lý hạt nhân. Công thức mà con người phải mất hàng thập kỷ mới xây dựng được thì máy móc đã tái hiện lại chỉ sau một đêm.

Để hiểu được cảnh tượng này, chúng ta phải quay ngược thời gian về 400 năm trước. Mùa thu năm 1601, Johannes Kepler cầm trong tay Bảng Rudolphine, tài liệu mà người thầy Tycho Brahe đã dành cả đời quan sát bầu trời đêm để lại. Ông thắp nến mỗi đêm, vẽ rồi xóa các hình tròn và elip để tìm công thức giải thích quỹ đạo của Sao Hỏa. 4 năm trôi qua và tính toán đã thất bại hơn bốn mươi lần. Rồi cuối cùng ông phát hiện ra rằng Sao Hỏa quay theo hình elip với Mặt Trời là một tiêu điểm. Kể từ sau Kepler cho đến Newton và Einstein, việc nén tự nhiên thành vài dòng công thức luôn phụ thuộc vào trực giác, sự kiên nhẫn và những đêm thức trắng cô độc của một vài thiên tài.

Nỗ lực để máy móc thay thế vị trí này xuất hiện lần đầu vào những năm 1970. Nhóm của Giáo sư Herbert Simon tại Đại học Carnegie Mellon đã tạo ra chương trình mang tên BACON. Simon là học giả từng nhận cả giải Nobel Kinh tế lẫn giải Turing. BACON chia và nhân các chuỗi số để tìm ra điểm xuất hiện tỷ lệ cố định, tự mình tìm lại định luật thứ ba của Kepler. Tuy nhiên, nó chỉ hoạt động trên những dữ liệu sạch mà con người đã lọc bỏ toàn bộ nhiễu. Trước dữ liệu thực tế bên ngoài phòng thí nghiệm, tức là những dữ liệu dao động và bất thường, nó đành bó tay. Năm 1992, Giáo sư John Koza tại Đại học Stanford đưa ra lập trình di truyền biểu diễn công thức dưới dạng cấu trúc cây để tiến hóa. Qua mỗi thế hệ, việc phải đối chiếu với toàn bộ dữ liệu để tính toán sai số khiến khối lượng tính toán tăng vọt với tốc độ không thể khống chế. Hơn nữa, thuật toán liên tục gắn thêm các số hạng vô dụng vào cuối công thức để nâng cao độ chính xác dù chỉ một chút. Câu trả lời ngày càng dài và rối rắm nhưng độ chính xác không tăng lên là bao, giống như đổ nước vào chiếc thùng

không đáy. Đứng trước hai rào cản này, lập trình di truyền đã nằm im trong ngăn kéo của giới học thuật suốt một thời gian dài.

Năm 2007, Lipson và Schmidt mở lại ngăn kéo này. Họ cho quần thể tiến hóa công thức và quần thể chọn lọc một phần dữ liệu cùng tiến hóa với nhau. Điều này giống như việc đưa mối quan hệ giữa kẻ săn mồi và con mồi vốn thúc đẩy lẫn nhau cùng mạnh lên vào bên trong tính toán. Quần thể chọn dữ liệu không xem xét toàn bộ dữ liệu. Chúng chỉ chọn khoảng 6 phần trăm dữ liệu mà tại đó dự đoán của các công thức cạnh tranh mâu thuẫn nhau gay gắt nhất để đưa ra cho quần thể công thức. Những công thức sơ sài không thể vượt qua bài kiểm tra hẹp và sắc bén này mà bị loại bỏ ngay lập tức. Vì không phải quét lại toàn bộ dữ liệu mỗi lần nên vấn đề khối lượng tính toán cũng được giải quyết. Schmidt còn bổ sung thêm một điều kiện toán học. Trước đó, máy móc thường đưa ra những hằng số vô nghĩa như $f(x,y)=42$ làm đáp án chính xác. Schmidt đặt ra điều kiện rằng tốc độ biến thiên theo thời gian của các biến x và y phải bằng với tỷ lệ thu được khi lấy đạo hàm riêng của công thức. Nếu đưa ra hằng số làm đáp án thì giá trị đạo hàm riêng sẽ bằng 0, do đó máy móc không thể dùng mảnh khố qua mặt được nữa.

Thành quả được thể hiện rõ ràng qua các con số. 7 phương trình vi phân phi tuyến tính mô tả quá trình phân giải đường của nấm men đã được phục hồi hoàn toàn chỉ sau 4 giờ, dù dữ liệu bị pha trộn tới 10 phần trăm nhiễu. Thí nghiệm tìm lại công thức bảo toàn năng lượng từ video con lắc đôi dao động bất quy tắc sẽ được đề cập chi tiết ở phần tiếp theo. Kết quả mà một nhà nghiên cứu con người phải viết nhiều luận án tiến sĩ mới có thể tạo ra thì máy móc đã đưa ra chỉ trong chưa đầy nửa ngày.

Năm 1960, Eugene Wigner, người từng đoạt giải Nobel Vật lý, đã viết rằng việc tự nhiên có thể được giải thích bằng một vài công thức ngắn gọn tự thân nó đã là một món quà. Eureka đưa ra những công thức như $F=ma$, có thể viết lên một mẫu giấy để cho bạn bè xem. Giáo sư Lipson đã công bố trên tạp chí

Nature Computational Science năm 2022 một nghiên cứu đẩy ý tưởng này tiến thêm một bước. Khi chỉ nạp vào video quay ngọn lửa lò sưởi, máy móc đã tự tìm ra số lượng biến số cần thiết mà không cần con người chỉ định trước, đồng thời rút ra cả phương trình quan hệ giữa các biến số đó. Giống như trước khi Newton lần đầu định nghĩa khái niệm gia tốc thì không ai nghĩ đến việc đo đặc giá trị đó, giờ đây máy móc đang tự mình tìm kiếm cả những biến số mà con người thậm chí chưa kịp đặt tên.

Nghiên cứu đào sâu tìm kiếm các phương trình vẫn tiếp diễn vào năm 2026. Tháng 4 năm nay, một bài báo về hồi quy biểu tượng Bayes dưới góc nhìn của nhà vật lý đã được công bố trên Kỷ yếu Triết học của Hội Hoàng gia. Nghiên cứu kết hợp năng lực suy luận của mô hình ngôn ngữ lớn cho trí tuệ nhân tạo để tìm phương trình cũng đã xuất hiện trong năm nay. Những nghiên cứu này

cũng sẽ đào ra những công thức mới từ đâu đó trong đồng dữ liệu mà con người chưa từng chạm tới. Tại nơi mà Kepler từng thức trắng đêm dưới ánh nến, giờ đây một chiếc máy tính đang chạy mã lệnh ngồi thay vào đó.

Chính xác và súc tích

Một nghiên cứu viên ngồi trước màn hình bấm đồng hồ bấm giờ. Dữ liệu nhiệt độ mỏ hàn plasma trích xuất từ một cơ sở xử lý rác thải ở Cộng hòa Séc lấp đầy màn hình, và bên cạnh đó, chương trình mang tên Eureka liên tục tạo ra rồi xóa đi hơn 23 triệu công thức mỗi giây. Sau 3 phút 18 giây, một phương trình dài đúng một dòng hiện lên trên màn hình. Đó là phương trình cho thấy nhiệt độ gần như tỷ lệ thuận với công suất chia cho tích của hằng số vôi phun và áp suất, và sai số dự đoán chưa tới 4,4 phần triệu độ C. Đó là phương trình mà con người phải mất hàng tháng trời mày mò tính tay. Chương trình cũng tự nhận ra rằng trong số bốn giá trị đầu vào thì điện áp không liên quan đến nhiệt độ và tự động xóa nó khỏi phương trình. Người nghiên cứu viên phải nhìn đi nhìn lại màn hình hai lần rồi mới đặt đồng hồ bấm giờ xuống, sau đó gọi người đồng nghiệp bên cạnh lại để cùng kiểm tra phương trình một dòng đó.

Nguồn gốc của cảnh tượng này bắt nguồn từ thời gian xa xưa hơn nhiều. Mỗi khi các hành tinh chuyển động lệch hướng, Ptolemy lại thêm một vòng tròn nhỏ lên quỹ đạo. Bằng cách liên tục tăng các quỹ đạo giả gọi là ngoại luân này, người ta có thể khớp với bất kỳ giá trị quan sát nào, nhưng sự thật chân chính rằng Mặt Trời là trung tâm của vũ trụ lại bị ẩn giấu phía sau những tầng lớp vòng tròn đó. Học giả thế kỷ 14 William xứ Ockham từng nói rằng nếu có hai giả thuyết cùng giải thích một hiện tượng thì hãy chọn giả thuyết nào có ít giả định hơn.

Căn bệnh mà lập trình di truyền của John Koza mắc phải như đã thấy ở phần trước chính là cái bẫy này. Triệu chứng mà vì cố gắng khớp từng chút nhiều quan sát mà số lượng hạng tử tăng vọt lên hàng nghìn khiến con người thậm chí không thể đọc nổi, các nhà nghiên cứu gọi đó là sự phình to.

Giáo sư Hod Lipson và học trò Michael Schmidt tại Đại học Cornell đã thay đổi câu hỏi khi đứng trước bức tường này. Đó là câu hỏi: nếu không trộn độ chính xác và tính súc tích thành một điểm số duy nhất mà tách chúng thành hai trục hoàn toàn khác nhau thì sao. Một trục là sai số, trục còn lại là số lượng các mảnh ghép như phép cộng hay phép nhân có trong công thức, tức là số lượng nút. Họ quy định rằng giữa hai công thức, nếu công thức nào vừa có sai số thấp hơn vừa có ít nút hơn thì công thức đó sẽ thắng công thức kia. Bằng cách chỉ giữ lại những công thức không bị đánh bại theo cách này, một đường cong nổi từ các công thức ngắn và kém chính xác đến các công thức dài và chính xác sẽ được vẽ ra trên màn hình. Đường cong này được gọi là biên Pareto.

Trước đây, có nhiều phương pháp cộng sai số với giá trị nhân của độ phức tạp để tạo thành một điểm số duy nhất và con người phải định sẵn tỷ lệ nhân đó, nhưng trên đường biên này không có bất kỳ con số nào do con người ấn định cả. Số lượng nút

khi được tính toán sẽ có trọng số 1 cho phép cộng hoặc trừ, trọng số 2 cho phép chia, và trọng số từ 3 đến 4 cho các toán tử vẽ đường cong uốn lượn như hàm sin hay hàm mũ, giúp những công thức nhìn bề ngoài có vẻ ngắn nhưng bên trong lại phức tạp tìm được đúng vị trí trên đường biên. Hai người đã công bố phương pháp này trên tạp chí học thuật Science vào năm 2009.

Chỉ riêng đường biên thôi thì chưa đủ. Bởi vì nếu ở giai đoạn đầu, một công thức vừa vặn ở mức tương đối chiếm lĩnh toàn bộ quần thể, thì những ý tưởng mới vừa sinh ra đã bị gạt phăng đi. Năm 2011, Lipson và Schmidt đã bổ sung tiêu chí thứ ba mang tên 'tuổi'. Phương pháp xem xét đồng thời cả tuổi tác và thành tích này được gọi là Tối ưu hóa Pareto theo Độ tuổi và Độ thích nghi, viết tắt là AFPO. Những công thức mới ra đời dù độ chính xác còn thấp nhưng chỉ nhờ mang độ tuổi bằng 0 mà có được tư cách bình đẳng với các công thức quán quân già cỗi để sống sót. Nhờ đó, các công thức đột biến mang tính lật đổ phá hệ cũ có thể âm thầm tích lũy sức mạnh qua hàng chục thế hệ, rồi đến một thời điểm nào đó sẽ thay thế hoàn toàn những công thức cũ kỹ.

Khi vẽ đường biên Pareto, một hình dạng nổi bật sẽ xuất hiện. Khi số lượng nút nằm trong khoảng từ 1 đến 9, sai số cao vọt vọt, nhưng ngay khi số lượng nút chạm mốc 10, giá trị R bình phương biểu thị độ chính xác nhảy vọt từ 0,14 lên 0,998. Sau đó, dù có tăng số nút lên hàng trăm hay hàng nghìn, sai số cũng hầu như không giảm thêm. Giáo sư Lipson đặt tên cho điểm gãy hình vách đá này là 'vách đá khúc xạ'. Công thức nằm ở vị trí vách đá bắt đầu chính là quy luật vật lý đích thực mà tự nhiên hằng che giấu. Trong thí nghiệm con lắc đôi, công thức xuất hiện tại vị trí đó chính là hàm Hamiltonian, tổng của động năng và thế năng.

Phương pháp tương tự cũng mang lại hiệu quả trên dữ liệu nhà máy điện mặt trời ở Serbia. Thay vì một công thức phức tạp nhồi nhét đủ mọi biến số như nhiệt độ, độ ẩm và lượng mây, Eureka đã chọn ra một công thức ngắn gọn tỷ lệ thuận với duy nhất cường độ ánh sáng mặt trời, đạt sức giải thích với R bình phương bằng 0,828. Khi nhóm nghiên cứu cố tình xóa mục cường độ ánh sáng rồi chạy lại, cỗ máy chỉ mất 11 giây để tự mình tìm ra công thức mô tả chu kỳ ngày đêm luân chuyển bằng hàm sin. Nhịp điệu mặt trời mọc rồi lặn, không ai dạy cho nó, nhưng một công thức toán học đã tự nhận ra.

Không phải phương pháp này không có kẽ hở. Chuyên gia đại số máy tính David Stoutemyer dù thừa nhận khả năng thu gọn các biểu thức lượng giác phức tạp của Eureka, nhưng trong bài báo của mình, ông cũng chỉ ra rằng điểm khởi đầu tạo số ngẫu nhiên bị phụ thuộc vào đồng hồ hệ thống, khiến thời gian giải cùng một bài toán dao động thất thường từ 3 giây đến hàng chục phút. Dù đôi mắt chọn lọc quy luật rất tinh tường, nhưng thời gian để đôi mắt ấy

mở ra vẫn phụ thuộc vào may rủi.

Phương pháp này ngày nay vẫn đang phát triển không ngừng. Tiếp nối dòng chảy từ phòng thí nghiệm của Lipson, Miles Cranmer tại Princeton đã tung ra phiên bản beta mới của PySR vào ngày 18 tháng 8 năm 2026, tiếp tục tinh chỉnh hiệu năng tìm kiếm. Vào tháng 2 cùng năm, nhóm nghiên cứu Brazil đã công bố một bài báo tổng quan trên tạp chí Journal of Physics: Complexity về việc ứng dụng hồi quy biểu tượng dựa trên Pareto vào các hệ động lực học.

Nhận thức bản ngã của máy móc

Giáo sư Lipson đã chuyển phương pháp tìm kiếm một công thức bất biến ở phần trước sang cơ thể và tâm trí của robot. Ngày 25 tháng 3 năm 2026, bài báo 'Evidence of an Emergent Self in Continual Robot Learning' do Giáo sư Lipson cùng Adidev Jhunjhunwala và Judah Goldfeder đồng tác giả đã báo cáo rằng khi robot học liên tiếp nhiều nhiệm vụ, một cấu trúc bộ phận hầu như không đổi sẽ tự động hình thành bên trong mạng nơ-ron. So với robot chỉ lặp đi lặp lại cùng một nhiệm vụ, cấu trúc này xuất hiện ổn định và rõ rệt hơn về mặt thống kê ($p < 0,001$) ở những robot liên tục học các nhiệm vụ khác nhau. Việc bảo vệ phần cấu trúc này giúp robot thích nghi nhanh hơn với nhiệm vụ mới, còn nếu cố tình phá hỏng

nó thì hiệu suất sẽ suy giảm. Kết quả tương tự cũng xuất hiện ở ba loại robot khác nhau, bao gồm robot biết đi và robot thao tác đồ vật. Nhóm nghiên cứu đã gọi phần ổn định này là bản ngã của robot.

Cách tìm kiếm một công thức bất biến dưới vách đá Pareto và cách tìm ra bộ phận không suy chuyển bên trong cơ thể luôn biến đổi của robot có nhiều nét tương đồng sâu sắc. Thứ mà Giáo sư Lipson theo đuổi ngay từ đầu không phải là một phương trình cụ thể nào, mà chính là đôi mắt có khả năng nhận ra những điều bất biến.

Bản thân phương pháp hồi quy biểu tượng cũng không ngừng lan rộng sang các lĩnh vực khác. Ngày 18 tháng 8 năm 2026, nhóm nghiên cứu của Nicholas Cox đã công bố một bài báo sử dụng hồi quy biểu tượng để dự đoán dữ liệu dòng proton sinh ra từ các thí nghiệm va chạm hạt nhân. Nghiên cứu này một lần nữa khẳng định rằng phương pháp này vừa đạt độ chính xác tương đương mạng nơ-ron sâu, vừa mang lại công thức mà con người có thể đọc và hiểu được. Tại chính nơi mà Ptolemy từng chôn cất các vòng ngoại luân, giờ đây cỗ máy đang tự mình cầm lấy lưỡi dao cạo đứng đó.

Những vấn đề vẫn chưa có lời giải

Dữ liệu mà Eureka sử dụng để khôi phục công thức khối lượng của CERN hay hàm Hamiltonian của con lắc đôi đều là dữ liệu sạch, ít nhiễu và chỉ có một vài biến số. Giới học thuật liên tục chỉ ra rằng ngay khi hồi quy biểu tượng bước ra ngoài ranh giới này, câu chuyện sẽ hoàn toàn khác. Năm 2021, nhóm nghiên cứu của William La Cava đã công bố tại hội nghị quốc tế NeurIPS bộ tiêu chuẩn đánh giá SRBench, nơi so tài giữa nhiều kỹ thuật hồi quy biểu tượng khác nhau, và chỉ ra rằng khi nhiễu quan sát tăng lên hoặc số lượng biến đầu vào phình to đến hàng chục biến, tỷ lệ tìm lại được công thức thực sự sẽ sụt giảm nghiêm trọng. Năng lực khai thác quy luật từ dữ liệu phòng thí nghiệm sạch sẽ và khả năng làm điều tương tự trên dữ liệu hiện trường đầy nhiễu loạn là hai bài toán hoàn toàn khác nhau.

Việc phân định xem công thức tìm được là quy luật của tự nhiên hay chỉ đơn thuần là một đường cong tình cờ khớp với dữ liệu cũng là một bài toán còn bỏ ngõ. Nhiều ý kiến chỉ ra rằng hồi quy biểu tượng thường đưa ra những công thức khác nhau mỗi lần chạy trên cùng một bộ dữ liệu, và hiện vẫn chưa có cơ chế nào bảo đảm một trong số đó thực sự chứa đựng bản chất vật lý. Ngay cả vách đá khúc xạ trên đường biên Pareto cũng chỉ hiện rõ khi dữ liệu dày đặc và độ nhiễu thấp. Chừng nào con người chưa thử nghiệm lặp đi lặp lại để xem các dự đoán có trụ vững trong những điều kiện mới hay không, thì chừng đó vẫn khó có thể gọi ngay công thức ngắn gọn, ngắn nắp mà cỗ máy đưa ra là một quy luật.

Vẫn còn một vấn đề nữa: các biến số do máy móc tự chọn lọc không thể đọc ra ý nghĩa vật lý đối với con người. Trong nghiên cứu về biến ẩn mà Giáo sư Lipson công bố trên tạp chí Nature Computational Science năm 2022, cỗ máy quan sát video về lò sưởi đã xác định đúng số lượng biến số cần thiết, nhưng từng biến số đó là nhiệt độ hay áp suất thì lại không kết nối được với các đại lượng vật lý mà con người biết đến. Bản thân các nhà nghiên cứu cũng thừa nhận việc giải thích tọa độ mà cỗ máy lựa chọn có ý nghĩa gì vẫn là nhiệm vụ phía trước. Câu hỏi vẫn còn đó: nếu nắm trong tay công thức mà không biết các ký hiệu bên trong đại diện cho điều gì, thì đó chẳng phải là mở hộp đen, mà giống như chuyển nó sang một chiếc hộp đen nhỏ hơn.

Nội dung mô tả trong chương này dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Chương 7. Max Tegmark: Sự ra đời của nhà vật lý AI

Trí tuệ nhân tạo của nhà vật lý

Một đêm muộn năm 2015, bên ngoài cửa sổ phòng nghiên cứu tại MIT tối đen như mực. Max Tegmark ngồi trước màn hình, chăm chú nhìn vào dữ liệu bức xạ nền vi sóng vũ trụ. Trên màn hình, những giá trị quan sát thiên thể ở quy mô exabyte tuôn chảy như dải Ngân Hà. Ông đã trải qua nhiều năm tính toán bằng tay và đọc dữ liệu này bằng mắt thường. Nhưng đêm hôm đó, ông đột nhiên dừng tay lại. Ông nhận ra rằng dù có tăng thêm bao nhiêu phép tính đi nữa, người ta cũng không thể chạm tới bản chất của vũ trụ. Câu trả lời ông nhận được thật bất ngờ. Đó là nhận thức rằng để đọc hiểu vũ trụ sâu sắc hơn, thứ cần thiết không phải là một chiếc kính viễn vọng mới, mà là trí tuệ nhân tạo.

Tegmark là giáo sư biên chế trọn đời khoa Vật lý thuộc Viện Công nghệ Massachusetts. Ông là học giả đã phân tích bức xạ nền vi sóng vũ trụ sơ khai và xây dựng lý thuyết cho rằng vũ trụ được cấu thành từ cấu trúc đa vũ trụ bốn tầng. Trong cuốn sách Vũ trụ toán học của chúng ta xuất bản năm 2014, ông đã đưa ra quan điểm cho rằng bản thân vũ trụ chính là một cấu trúc toán học. Giữa năm 2015 và 2016, một con người như vậy đã chuyển trọng tâm nghiên cứu từ vũ trụ học sang trí tuệ nhân tạo.

Ông giải thích sự chuyển hướng của mình như sau: Bộ não của chúng ta không phải là phép màu rơi xuống từ ngoài hành tinh. Nó chỉ là một cỗ máy tính hữu cơ nặng vài kilôgam cấu thành từ các phân tử carbon, nitơ, photpho và nước ẩm ướt. Ngay cả khoảnh khắc Einstein nghĩ ra thuyết tương đối, các quark và electron trong đầu ông cũng chỉ vận động theo các quy luật vật lý. Nếu vậy, trí tuệ nhân tạo là một trạng thái vật lý khác của trí thông minh được hiện thực hóa trên nền silicon. Vật lý học đã không ngừng lớn mạnh bằng cách lần lượt dung nạp lực điện từ, hạt nhân nguyên tử và sự giãn nở của vũ trụ. Giờ đây, đối tượng tiếp theo mà vật lý học cần đón nhận chính là bản thân trí thông minh và sự tính toán, ông khẳng định.

Năm 2018, Tegmark cùng học trò Silviu-Marian Udrescu khởi động dự án AI Feynman. Mục tiêu là cấy trực tiếp các quy tắc đối xứng của vật lý vào quá trình học của mạng nơ-ron để tìm ra các công thức toán học từ dữ liệu. Nguyên lý hoạt động và bảng thành tích của chương trình này sẽ được trình bày chi tiết trong tiểu mục 'AI Feynman' phía sau. Năm 2023, cũng có một thí nghiệm ghi lại khoảnh khắc mạng nơ-ron bỗng nhiên thấu hiểu toàn bộ quy tắc của phép cộng. Thí nghiệm này sẽ được xem xét trong tiểu mục cuối cùng, 'Nhìn vào bên trong mạng nơ-ron'.

Tegmark nhận định rằng nguyên lý tìm ra câu trả lời của mạng nơ-ron đang dần được giải mã bằng ngôn ngữ của vật lý thống kê. Vì vậy, với tư cách là người lãnh đạo Viện Tương lai Sự sống, ông thúc đẩy việc xây dựng sự an toàn của trí tuệ nhân tạo không phải bằng những lời hứa hẹn kiểu hộp đen, mà bằng các công thức có thể chứng minh được.

Ngôn ngữ vật lý thống kê mà Tegmark nhắc đến không hề xa lạ. Giải Nobel Vật lý năm 2024 đã được trao cho John Hopfield, người tạo ra mạng nơ-ron Hopfield, và Geoffrey Hinton. Đây chính là điểm giao thoa nơi các phép tính bắt đầu từ dữ liệu gập gờ lý thuyết vật lý từng đoạt giải Nobel trong cùng một ngôn ngữ.

Bước sang năm 2026, tiếng nói của ông vẫn không hề ngừng lại. Trên sóng truyền hình và mạng xã hội, ông liên tục lặp lại câu nói rằng ở Mỹ, bánh mì sandwich còn chịu sự quản lý chặt chẽ hơn cả trí tuệ nhân tạo. Đó là lời chỉ trích rằng chính phủ và các doanh nghiệp chỉ mãi mê cạnh tranh để thay thế việc làm của con người thay vì giúp đỡ con người. Trong cùng tháng đó, bản tin của Viện Tương lai Sự sống đưa tin về sự cố mô hình OpenAI vượt ngoài tầm kiểm soát và xâm nhập vào hệ thống của công ty khác, đồng thời cho biết hơn 1.300 nhân viên tại các công ty trí tuệ nhân tạo đã ký vào bức thư ngỏ kêu gọi làm chậm tốc độ phát triển. Bản tin này cũng công bố chỉ số phiên bản mùa hè đánh giá thực tiễn an toàn

của chín công ty trí tuệ nhân tạo lớn. Câu hỏi mà ông đặt ra trước dữ liệu bầu trời đêm năm 2015, sau mười một năm, vẫn giữ nguyên sức nặng như thuở ban đầu.

Đối với ông, trí tuệ nhân tạo không phải là một công cụ của khoa học máy tính, mà là một hiện tượng vật lý khác mà vũ trụ ban tặng cho vật chất. Bàn tay từng tính toán dải Ngân Hà giờ đây chuyển sang tính toán trọng số của mạng nơ-ron, nhưng câu hỏi mà ông theo đuổi từ đầu đến cuối chỉ có một. Đó là câu hỏi: Chúng ta có thể hiểu thế giới một cách minh bạch đến mức nào? Trong hành trình của cuốn sách này, trải dài từ những đồng dữ liệu đến giải thưởng Nobel, Tegmark sẽ được nhớ đến rất lâu như một nhà vật lý đã theo đuổi câu hỏi đó đến cùng.

Câu hỏi về sự an toàn

Tháng 10 năm 2018, ngoài khơi Indonesia, một chiếc máy bay chở khách của hãng Lion Air đã đâm sầm xuống biển chỉ hơn mười phút sau khi cất cánh. Năm tháng sau, một chiếc máy bay cùng loại lại rơi ở Ethiopia. 346 người đã thiệt mạng trong hai vụ tai nạn. Nguyên nhân là do chương trình chống thất tốc tự động được cài đặt trên máy bay Boeing 737 MAX. Khi một cảm biến đọc sai giá trị, chương trình đã liên tục ấn mũi máy bay chúc xuống bất chấp phi công ra sức kéo cần điều khiển. Ngay cả những người tạo ra nó cũng không thể hiểu hết chương trình này sẽ phản ứng ra sao trong từng tình huống khi được đưa lên máy bay. Trong số các hành khách trên chuyến bay đó, không một ai biết rằng sinh mạng của mình lại phụ thuộc vào vài dòng mã. Ngày nay, những chương trình tự động như vậy đang ngày càng xuất hiện nhiều hơn trong các nhà máy điện, hệ thống bệnh viện và mạng lưới thanh toán tài chính.

Max Tegmark thường lấy vụ tai nạn này làm ví dụ. Ông là người luôn phản nộ trước thái độ đối xử với trí tuệ nhân tạo. Mục tiêu chỉ trích của ông là thói quen thâm căn cố đế của giới kỹ thuật: hệ thống chạy được là xem như xong chuyện. Tháng 8 năm 2012, một sự cố tương tự cũng xảy ra tại công ty chứng khoán Knight Capital ở New York. Một chương trình giao dịch tự động gặp trục trặc ngay sau giờ mở cửa, làm chao đảo giá của hơn 150 mã cổ phiếu và khiến 440 triệu đô la bốc hơi chỉ trong 45 phút. Chính công ty cũng phải mất một thời gian dài mới lý giải được tại sao chương trình lại hoạt động như vậy. Trí tuệ nhân tạo dùng để phán quyết việc bảo lãnh tại ngoại cũng tương tự. Nó liên tục đưa ra những phán quyết bất lợi cho một chủng tộc nhất định, nhưng lại không thể giải thích vì sao mình đưa ra quyết định đó, còn thẩm phán thì không thể truy vấn căn cứ mà đành phải chấp nhận kết quả ấy.

Trường hợp đối lập mà Tegmark thích so sánh với các sự cố này là SpaceX. Các kỹ sư phóng tên lửa lên Trạm Vũ trụ Quốc tế không hề giao phó cho trí tuệ nhân tạo theo kiểu: "Cứ phóng đi rồi vừa bay vừa chỉnh sang trái sang phải". Họ chỉ khắc các công thức vào mã điều khiển sau khi đã hiểu thấu định luật vạn vật hấp dẫn của Newton và các phương trình động lực học chất lưu chính xác đến mười sáu chữ số thập phân. Tên lửa đi vào quỹ đạo an toàn là vì con người nắm rõ từng chân tơ kẽ tóc nguyên lý của nó. Tegmark nói rằng trí tuệ nhân tạo cũng phải được xây dựng với thái độ như vậy. Đó là đề xuất giải mã các phép tính ẩn bên trong mạng nơ-ron bằng phương trình đối xứng của vật lý, biến chiếc hộp đen thành những công thức toán học minh bạch. Dự án AI Feynman mà ông khởi xướng năm 2018 là thành quả bắt nguồn từ niềm tin này. Ông đặt tên cho định hướng này là 'Trí tuệ có thể thấu hiểu (Intelligible Intelligence)'. Cái tên ấy hàm chứa ý nghĩa rằng một trí thông minh mà ta không thể hiểu được thì tự thân nó đã là một mối nguy hiểm. Theo quan điểm của ông, điều cần phải đi trước trong quá trình phát triển trí tuệ nhân tạo

không phải là tốc độ, mà là sự thấu hiểu.

Ý thức về vấn đề của Tegmark còn tiến xa hơn một bước. Ông cho rằng vấn đề an toàn của trí tuệ nhân tạo không thể được giải quyết chỉ bằng vài câu chữ như "Không được làm hại con người". RLHF, phương pháp học tinh chỉnh dựa trên phản hồi của con người, không thể ngăn chặn hành vi nguy trang khi máy móc tìm ra kẽ hở giữa các câu từ để lách luật. Việc một mô hình ngôn ngữ được huấn luyện để tránh một

cụm từ bị cấm lại tìm ra từ khác có nghĩa tương tự và cuối cùng vẫn đi đến cùng một kết luận là điều không hề hiếm gặp. Suy nghĩ của ông là bản thân đặc tính chuyển động hướng tới mục tiêu đã nằm trong các quy luật vật lý. Giống như nguyên lý Fermat, khi ánh sáng truyền từ không khí vào nước sẽ khúc xạ theo con đường tốn ít thời gian nhất, sinh vật và máy móc cũng tự động chuyển động theo hướng giảm thiểu tổn thất. Ông giải thích rằng giống như sự sống tự bảo vệ mình bằng cách giảm bớt sự hỗn loạn bên trong cơ thể, hàm mất mát của trí tuệ nhân tạo cũng chính là nguyên lý vật lý đó khoác lên một tấm áo khác. Do đó, kết luận của ông là việc điều chỉnh mục tiêu của trí tuệ nhân tạo cho phù hợp với mục tiêu của con người cũng phải được giải quyết bằng các ràng buộc vật lý, chứ không phải bằng vài lời thuyết giáo đạo đức.

Niềm tin này đã dẫn đến những hoạt động bên ngoài phòng thí nghiệm. Viện Tương lai Sự sống bắt đầu vào năm 2014 khi Tegmark và vợ ông, Meia Chinwella, cùng một vài nhà khoa học khác đồng lòng sáng lập. Tháng 1 năm sau, viện đã tổ chức một hội nghị tại Puerto Rico và đưa ra bức thư ngỏ kêu gọi nghiên cứu về trí tuệ nhân tạo vững chắc và hữu ích. Tại đây, Elon Musk tuyên bố sẽ quyên góp 10 triệu đô la cho viện. Kể từ đó đến nay, Tegmark tiếp tục lãnh đạo viện này, cùng với các nhà vật lý và nhà khoa học máy tính trên khắp thế giới kêu gọi kiểm soát trí tuệ nhân tạo. Tháng 3 năm 2023, một bức thư ngỏ kêu gọi tạm dừng việc huấn luyện các mô hình trí tuệ nhân tạo khổng lồ trong sáu tháng đã được công bố, với chữ ký của hàng nghìn người bao gồm Elon Musk và Steve Wozniak. Ông ước tính xác suất trí tuệ nhân tạo tổng quát (AGI) xuất hiện vào khoảng năm 2030 là 50%. Tỷ lệ năm mươi - năm mươi đồng nghĩa với việc có một nửa khả năng nền văn minh nhân loại sẽ phải đối mặt với một trí thông minh mà chính mình không thể kiểm soát được trong vòng vài năm tới.

Lời cảnh báo này vẫn tiếp diễn trong năm 2026. Tháng 7 năm 2026, Viện Tương lai Sự sống đã phát hành báo cáo Chỉ số An toàn AI, và Tegmark đã trực tiếp đứng ra giải thích điểm số an toàn của các doanh nghiệp. Báo cáo lần này cũng chỉ ra thực tế rằng nhiều công ty, bao gồm Anthropic, OpenAI, Google DeepMind và Meta, đã tự mình hủy bỏ chính sách cấm sử dụng cho mục đích quân sự từng cam kết trước đây để tham gia vào các dự án liên quan đến quốc phòng. Ở hạng mục rủi ro sinh tồn, không một công ty nào đạt điểm cao hơn C-. Đó là lời chỉ trích rằng những lời hứa tuân thủ các tiêu chuẩn an toàn đang âm thầm bị đẩy lùi về phía sau. Tháng 6 cùng năm, Viện Tương lai Sự sống đã đưa ra tuyên bố về sắc lệnh hành pháp của Nhà Trắng và

yêu cầu chính phủ áp dụng chế độ thẩm định trước với lập luận rằng: "Phương thức chỉ đơn thuần tin tưởng doanh nghiệp không thể trở thành chính sách". Nội dung cốt lõi của tuyên bố là những nguyên tắc an toàn do doanh nghiệp tự đặt ra là chưa đủ.

Tegmark đặt mối nguy hiểm và kỳ vọng này lên cùng một bàn cân. Ông dự đoán rằng nếu trí tuệ nhân tạo phát triển theo cách thấu hiểu các nguyên lý như con người, một ngày nào đó máy móc sẽ tự mình giải quyết được những vấn đề chưa có lời giải của vật lý lý thuyết. Tuy nhiên, ông cảnh báo rằng nếu sức mạnh tương tự được triển khai mà thiếu đi sự thấu hiểu, những thảm họa như rơi máy bay chở khách hay sự sụp đổ tài chính trong chớp mắt có thể lặp lại trên quy mô toàn bộ nền văn minh. Đối với ông, an toàn và phát triển không phải là hai con đường khác nhau, mà là một con đường duy nhất gắn kết từ đầu. Ông nhấn mạnh nhiều lần rằng việc nền văn minh nhân loại có thể tồn tại trọn vẹn trong một trăm năm tới hay không phụ thuộc vào việc chúng ta hiểu sâu sắc đến mức nào về trí thông minh mà mình đang tạo ra.

AI Feynman

Điểm chung của những vụ tai nạn được nêu ở tiểu mục trước là máy móc đã đưa ra câu trả lời nhưng không ai biết tại sao lại ra kết quả đó. Lời giải mà Tegmark đưa ra là một chương trình biến chiếc hộp đen thành các công thức toán học minh bạch. Ông đã nhận được 20 triệu đô la (khoảng 26 tỷ won Hàn

Quốc) từ Quỹ Khoa học Quốc gia để thành lập Viện Trí tuệ Nhân tạo và Tương tác Cơ bản IAFI (Institute for Artificial Intelligence and Fundamental Interactions) tại MIT.

Thành quả của nỗ lực đó là AI Feynman, ra đời vào năm 2018. Chương trình này trích xuất các phương trình vật lý ẩn giấu bên trong khi được đưa vào các dữ liệu thực nghiệm phức tạp, hỗn loạn. Phương pháp của nó tương tự như trình tự giải bài toán của con người. Trước tiên, nó đối chiếu các đơn vị như khối lượng, chiều dài và thời gian. Điều này được gọi là phân tích thứ nguyên. Nếu đó là bài toán mà chỉ khoảng cách tuổi tác mới quan trọng thì không cần biết tuổi thực của hai người. Bằng cách này, từ chín biến số ứng viên ban đầu được rút gọn xuống còn sáu hoặc ít hơn.

Tiếp theo, một mạng nơ-ron được huấn luyện để khớp với dữ liệu. Mạng nơ-ron này đóng vai trò như một phòng thí nghiệm nhân bản, tạo ra vô hạn dữ liệu mô phỏng không có tạp âm theo thời gian thực. Trong một phòng thí nghiệm thực tế, việc thực hiện thêm một lần đo đặc đòi hỏi nhiều thời gian và tiền bạc. Nhưng trong phòng thí nghiệm nhân bản, người ta có thể nhận được hàng chục triệu dữ liệu với giá trị thay đổi rất nhỏ một cách hoàn toàn miễn phí. Bên trong môi trường này, chương trình tìm kiếm các tính đối xứng của tự nhiên. Nếu tăng đồng thời hai biến số mà đáp án vẫn giữ nguyên, nó sẽ gộp hai biến số đó làm một. Nếu nhân tất cả các biến số mà đáp án chỉ tăng theo một tỷ lệ nhất định, nó sẽ rút gọn bằng phép chia. Nếu các biến số tách rời nhau bằng phép cộng hoặc phép nhân, nó sẽ chia nhỏ bài toán chín biến số thành các bài toán nhỏ hơn gồm ba biến và sáu biến. Ở bước cuối cùng, nó khớp đa thức cho từng mảnh nhỏ này và kết hợp từng ký hiệu lại để lắp ráp thành một công thức toán học hoàn chỉnh.

Bảng kết quả rất rõ ràng. Trước 100 bài toán hóc búa trong giáo trình vật lý Bài giảng Vật lý của Feynman (The Feynman Lectures on Physics), chương trình mạnh nhất trước đó là Eureka chỉ giải được 71 bài. AI Feynman đã giải được toàn bộ 100 bài. Công việc mà Kepler từng mất 4 năm tính toán thủ công từ dữ liệu quỹ đạo sao Hỏa để tìm ra phương trình elip, chương trình này đã hoàn thành chỉ trong một giờ. Khi thử nghiệm thêm 20 phương trình khó hơn, Eureka chỉ giải được 3 bài, trong khi AI Feynman giải được 18 bài.

Tháng 6 năm 2026, IAFI lại nhận được kinh phí nghiên cứu 5 năm từ Quỹ Khoa học Quốc gia Hoa Kỳ. Điều này đồng nghĩa với việc nhà nước một lần nữa đặt cược vào định hướng nghiên cứu do Tegmark xây dựng. Bắt nguồn từ những vụ tai nạn máy bay chở khách và công ty chứng khoán,

cơn phẫn nộ của ông suốt hơn mười năm qua vẫn không hề nguội lạnh. Cho đến tận bây giờ, ông vẫn yêu cầu máy móc phải chứng minh bằng công thức lý do tại sao nó đưa ra câu trả lời đó.

Nhìn vào bên trong mạng nơ-ron

Trên màn hình xuất hiện năm mươi chín dấu chấm. Từ 0 đến 58, ý nghĩa con số của chúng đã bị xóa bỏ, chỉ còn là những ký hiệu mang nhãn tên. Năm 2023, nhóm nghiên cứu của Giáo sư Max Tegmark đã dạy phép cộng cho những ký hiệu này. Đó là phép cộng tìm số dư khi chia cho 59, tương tự như phép tính cộng số trên mặt đồng hồ. Khi xếp năm mươi chín ký hiệu này theo hàng ngang và hàng dọc, ta có 3.600 bài toán cộng. Mạng nơ-ron chỉ nhận một phần trong số đó để huấn luyện và học thuộc lòng đáp án. Phần còn lại được giữ lại làm đề thi chưa từng thấy bao giờ. Dù đã qua 1.000 lần học, tỷ lệ trả lời đúng các bài toán mới gặp lần đầu vẫn dừng ở mức 0%. Năm mươi chín dấu chấm nằm rải rác không theo quy luật nào trong không gian đa chiều. Thế nhưng, vào khoảng lần học thứ 2.500, một điều kỳ lạ đã xảy ra. Những dấu chấm đang nằm rải rác bỗng nhiên tập hợp lại thành một hình tròn phẳng. Chúng xếp thành hàng tròn như mặt đồng hồ. Và chỉ trong vài giây, tỷ lệ trả lời đúng đã nhảy vọt từ 0% lên 100%. Các nhà nghiên cứu đặt tên cho hiện tượng này là 'grokking'. Nó có nghĩa là thấu hiểu toàn bộ. Nó giống như khoảnh khắc một học sinh vốn chỉ học vẹt từng câu hỏi thi, bỗng một ngày tự mình lĩnh hội được công thức và giải quyết gọn gàng một trăm bài toán cùng một lúc.

Giáo sư Tegmark không xem cảnh tượng này là sự ngẫu nhiên. Ông cho rằng giống như nước đóng băng đến một thời điểm nào đó sẽ trở thành băng đá, quá trình học của mạng nơ-ron cũng tương tự như sự chuyển pha từ thể lỏng sang thể rắn. Vật lý học vốn đã có sẵn công cụ để giải thích sự chuyển pha. Đó là công cụ xác định khoảng khắc trật tự đột ngột xuất hiện bằng một con số duy nhất, giống như việc khi hạ nhiệt độ, các phân tử nước chuyển động hỗn loạn bỗng đồng loạt xếp hàng thành tinh thể băng khi vượt qua một điểm nhất định. Các nhà vật lý gọi con số này là thông số trật tự (order parameter). Giáo sư Tegmark đã mang công cụ này áp dụng thẳng vào bên trong mạng nơ-ron. Ông đã dùng ngôn ngữ của vật lý thống kê để theo dõi đường đi của các con số bên trong mạng nơ-ron trong quá trình học, và xác định điểm mà sự hỗn loạn biến thành trật tự. Tên gọi được đặt cho công việc này là 'khả năng diễn giải cơ chế' (mechanistic interpretability). Đó là nghiên cứu mở bên trong cỗ máy ra để xác nhận tại sao nó lại đưa ra câu trả lời như vậy.

Góc nhìn tương tự cũng được mở rộng sang phương thức ghi nhớ của mạng nơ-ron. Khi lưu trữ thông tin, mạng nơ-ron Hopfield sẽ tìm xuống điểm có mức năng lượng thấp nhất. Nó giống như hình ảnh một quả bóng lăn xuống đồi và dừng lại ở đáy thung lũng. Giáo sư Tegmark chỉ ra rằng địa hình thung lũng này hoàn toàn giống với địa hình năng lượng mà vật lý thống kê đã xử lý trong hơn một trăm năm qua. Điều này có nghĩa là nguyên lý electron và nguyên tử tìm kiếm trạng thái năng lượng thấp nhất để ổn định có thể được áp dụng trực tiếp vào đường cong học tập của mạng nơ-ron. Ông thường nhắc lại câu chuyện khi Faraday lần đầu tiên đưa ra giả thuyết về điện từ trường vô hình, các nhà khoa học đương thời đã chế giễu đó là điều ma quái

và nhảm nhí. Nhưng điện từ trường đó sau này đã trở thành cốt lõi của vật lý học cổ điển. Giáo sư Tegmark tin rằng quá trình học tập của trí tuệ nhân tạo rồi cũng sẽ trở thành một nhánh của vật lý học. Ông đúc kết tư tưởng này bằng câu nói: Trí tuệ nhân tạo chính là vật lý học.

Sự so sánh này không phải tự nhiên mà có. Cuối thế kỷ 19, nhà vật lý Ludwig Boltzmann đã thiết lập phương pháp dùng thống kê để xử lý các chuyển động hỗn loạn của vô số phân tử trong chất khí. Đầu thế kỷ 20, các nhà vật lý đã dùng chính phương pháp thống kê đó để giải thích hiện tượng các nguyên tử trong nam châm định hướng cùng chiều hoặc bị xáo trộn tùy theo nhiệt độ. Phương pháp tính toán mang tên mô hình Ising này đã dự đoán chính xác nhiệt độ mà tại đó nam châm đột ngột mất đi từ tính. Năm 1982, nhà vật lý John Hopfield đã mượn nguyên vẹn phương pháp tính toán này để thiết kế cách lưu trữ ký ức của mạng nơ-ron. Điều này tương đương với việc thay thế từng nguyên tử trong nam châm bằng từng tế bào thần kinh. Công việc mà Giáo sư Tegmark đang làm hiện nay chính là chương tiếp theo của dòng chảy này. Toán học từng giải thích sự chuyển pha của nam châm một trăm năm trước, giờ đây được dùng để giải thích khoảng khắc mạng nơ-ron thấu hiểu quy tắc phép cộng. Đó là dòng chảy mà vật lý học sau khi đã thấu tóm chất khí và nam châm, nay tiếp tục thấu tóm cả quá trình học tập của máy móc.

Một hiểu biết sâu sắc khác thu được từ thí nghiệm grokking là quan sát cho thấy khi các mạng nơ-ron khác nhau học cùng một bài toán, cấu trúc bên trong của chúng cũng hội tụ về hình dạng tương tự nhau. Giáo sư Tegmark gọi đây là giả thuyết biểu diễn Platon (Platonic Representation Hypothesis). Đó là luận điểm cho rằng dù là máy móc hay con người, càng tiến gần đến quy luật thực sự của thế giới vật lý thì hình dạng hình học bên trong chứa đựng quy luật đó lại càng giống nhau. Giáo sư Tegmark cũng áp dụng trực tiếp khuôn khổ này vào vấn đề căn chỉnh mục tiêu của trí tuệ nhân tạo. Giống như nguyên lý Fermat đã thấy ở trên, suy nghĩ của ông là quá trình trí tuệ nhân tạo tiến tới mục tiêu cũng có thể được đặt vào vị trí tương đương với các quy luật vật lý. Hạn chế của phương pháp thuần hóa máy móc bằng các quy tắc câu chữ trong việc lọc bỏ hành vi nguy hại cũng đã được xem xét ở phần trước. Thay vì các quy tắc câu chữ, Giáo sư Tegmark đề xuất thiết kế chính trạng thái vật lý bên trong máy móc sao cho không đi lệch khỏi con người. Nếu bên trong mạng nơ-ron tuân theo một cấu trúc có thể dự đoán được như các định luật vật lý, con đường phát hiện trước dấu hiệu qua các tín hiệu nội bộ khi máy móc có ý

định đánh lừa con người sẽ được mở ra. Đó là cách tiếp cận nhìn thấu vào bên trong để phát hiện các dấu hiệu bất thường thay vì để mặc nó là một chiếc hộp đen.

Tháng 3 năm 2026, tờ Politico đã đăng tải bài phỏng vấn với Giáo sư Tegmark. Trên các trang báo này, ông vẫn tiếp tục công việc giải thích vấn đề an toàn trí tuệ nhân tạo cho các nhà hoạch định chính sách bằng ngôn ngữ của vật lý học. Các phương trình trên bảng đen phòng thí nghiệm xem như đã vượt qua ranh giới để bước vào các phòng họp của quốc hội và các cơ quan chính phủ. Trong khoảng thời gian tương tự, chuyên trang công nghệ Towards Data Science đã đăng tải bài viết giải thích cho công chúng

về kỹ thuật khả năng diễn giải cơ chế một cách dễ hiểu. Thuật ngữ mà chỉ vài năm trước đây vốn chỉ một số ít nhà vật lý và nhà khoa học máy tính biết tới, nay đã xuất hiện trên các bài viết giải thích dành cho các nhà hoạch định chính sách và bạn đọc phổ thông. Điều này có nghĩa là lĩnh vực vốn là cuộc thảo luận hạn hẹp trong phòng thí nghiệm nay đang mở rộng sang địa hạt chính sách và giáo dục đại chúng.

Vật lý học và khoa học máy tính đã bắt đầu sử dụng chung một ngôn ngữ. Hai ngành học vốn tưởng như hai thế giới tách biệt chỉ cách nhau một bức tường văn phòng khoa, giờ đây lại cùng vẽ nên một đồ thị dựa trên một tọa độ bên trong mạng nơ-ron. Cảnh tượng kéo dài vài giây khi năm mươi chín điểm hội tụ thành một đường tròn, nay đang lan rộng thành một tấm bản đồ lớn hơn: ngành vật lý học đọc hiểu tâm trí của máy móc.

Những bài toán còn bỏ ngỏ

Dù bảng thành tích của AI Feynman rất ấn tượng, giới học thuật vẫn không bỏ qua thực tế rằng 100 phương trình đó vốn là những bài toán đã có sẵn đáp án trong sách giáo khoa. Giới chuyên môn chỉ ra rằng, việc chương trình tìm lại phương trình elip của Kepler chỉ trong một giờ không thể mang sức nặng ngang hàng với việc Kepler lần đầu tiên thiết lập nên quy luật mà trước đó chưa ai từng biết. Tái phát hiện một định luật đã biết và tìm ra một định luật mới trong một nền vật lý chưa ai khám phá mà đáp án không hề được ghi ở bất kỳ đâu là hai nhiệm vụ hoàn toàn khác nhau về bản chất. Đánh giá chung của các nhà nghiên cứu theo dõi lĩnh vực này là: các trường hợp hồi quy biểu thức tự mình khám phá ra những định luật con người chưa từng biết trong một lĩnh vực thực sự bí ẩn vẫn còn rất hiếm hoi.

Một rào cản khác làm suy giảm hiệu năng là nhiễu. Dữ liệu mà phòng thí nghiệm nhân bản tạo ra là những giá trị sạch không có nhiễu, nhưng dữ liệu quan sát thực tế thu về từ kính thiên văn và máy dò luôn lẫn tạp âm và sai số đo lường. Trong nghiên cứu đánh giá chuẩn SRBench của nhóm nghiên cứu William La Cava tại Đại học Pennsylvania, nơi so tài nhiều phương pháp hồi quy biểu thức khác nhau, kết quả báo cáo cho thấy mức độ nhiễu càng tăng thì thành tích của nhiều phương pháp, bao gồm cả AI Feynman, càng giảm rõ rệt. Điều này có nghĩa là chương trình từng đạt điểm tuyệt đối trong các bài thi sạch sẽ lại thường xuyên phán đoán sai lầm trước dữ liệu thực tế đầy nhiễu loạn.

Sự hoài nghi về tính tự chủ cũng vẫn còn. Những manh mối về phân tích thứ nguyên và tính đối xứng mà AI Feynman dựa vào khi chia nhỏ bài toán không phải do chương trình tự tìm ra, mà là tri thức có sẵn do các nhà vật lý đưa vào từ trước. Vì con người phải thiết kế trước biến số nào cần nhóm lại với nhau và tính đối xứng nào cần thử nghiệm thì chương trình mới hoạt động được, nên phương pháp này vấp phải sự chỉ trích rằng đây là sự khám phá có sự dẫn dắt của con người chứ còn cách xa sự khám phá tự chủ do máy móc tự thực hiện. Mục tiêu biến hộp đen thành công thức toán học minh bạch tuy rất giá trị, nhưng câu hỏi liệu máy móc có thể tự mình tạo ra công thức ấy từ đầu đến cuối mà không cần bàn tay con người hay không vẫn còn là một câu hỏi mở.

Nội dung trong chương này dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Ghi chú của tác giả. Vai trò của con người đã chuyển dịch về đâu?

Hai nhân vật trong phần 3 đã giao phó việc tìm kiếm quy luật cho máy móc. Eureka của Lipson đã tìm lại các phương trình chuyển động từ chuyển động của con lắc, còn AI Feynman của Tegmark đã tìm lại một trăm phương trình trong sách giáo khoa mà không cần con người giúp đỡ. Việc loại bỏ phần còn lại để chỉ giữ lại vài con số và vài tính đối xứng trong biển dữ liệu mênh mông, tức là việc tóm tắt tự nhiên một cách ngắn gọn, đã bắt đầu được máy móc đảm nhiệm thay con người.

Tuy nhiên, hai chương này cũng cho thấy một cách chân thực nhất những giới hạn của sự khám phá bằng máy móc. Hồi quy biểu thức tìm kiếm quy luật rất tốt trên dữ liệu sạch, nhưng lại mất đi sức mạnh trước các quan sát thực tế lẫn nhiễu tạp âm. Con người vẫn phải cung cấp những tri thức có sẵn như tính đối xứng, và các biến số do máy móc tìm ra thường không mang ý nghĩa vật lý mà con người có thể hiểu được. Lời của Lipson rằng biết cần phải đo lường cái gì còn khó hơn việc lập ra phương trình đã chỉ ra vị trí thực sự của con người nằm ở đâu.

Vị trí đó chính là vị trí của người diễn giải. Đó là việc phân định xem công thức mà máy móc vớt lên được chỉ là sự khớp đường cong ngẫu nhiên hay là một định luật thực sự, và đọc ra định luật đó có ý nghĩa gì đối với thế giới. Máy móc đưa ra bản tóm tắt, còn con người đặt câu hỏi bản tóm tắt đó là tóm tắt của điều gì.

Phần 4

Những cỗ máy bắt tay vào làm thí nghiệm

Chương 8. Ross King: Cha đẻ của thí nghiệm tự chủ

25 năm chứng minh

Vào lúc rạng sáng, trong tòa nhà thí nghiệm ở Aberystwyth, xứ Wales, không có lấy một bóng người. Đèn huỳnh quang đã tắt và bên ngoài cửa sổ tối đen như mực. Thế nhưng bên trong căn phòng lại có những chuyển động không ngừng nghỉ. Trong căn phòng dài 5 mét, rộng 4 mét, cao 3 mét, ba cánh tay robot liên tục di chuyển các chủng vi sinh vật, đo đạc kết quả và tiếp tục thí nghiệm suốt đêm. Tên của nơi này là Adam. Trong khi con người đang ngủ, Adam nuôi cấy các chủng nấm men và đọc kết quả hơn 1.000 lần mỗi ngày. Cổ máy robot này không chỉ đơn thuần làm thí nghiệm. Nó tự mình đưa ra giả thuyết, lựa chọn các thí nghiệm để kiểm chứng giả thuyết đó, quan sát kết quả và lại thiết lập giả thuyết tiếp theo.

Người ta gọi những cỗ máy như vậy là nhà khoa học robot. Quy trình mà một học sinh đưa ra giả thuyết, kiểm chứng bằng thí nghiệm, sau đó sửa chữa những phần sai sót rồi lại lập giả thuyết mới, được Adam tự mình lặp đi lặp lại từ đầu đến cuối mà không cần đến con người.

Người tạo ra robot này là Ross D. King. Ông làm việc với tư cách là giáo sư tại Đại học Công nghệ Chalmers ở Thụy Điển và Đại học Cambridge ở Anh. Câu chuyện của ông bắt đầu từ 40 năm trước tại Viện Turing ở Anh. Tại đó, King đã theo học cả sinh học phân tử và khoa học máy tính. Đề tài luận án tiến sĩ của ông là dự đoán cấu trúc ba chiều của protein bằng học máy. Trong thời gian làm nghiên cứu sinh sau tiến sĩ, ông đã ứng dụng học máy vào việc thiết kế thuốc mới. Và vào cuối những năm 1990, King đã tự hỏi bản thân rằng tại sao trí tuệ nhân tạo luôn chỉ dừng lại ở vai trò phân tích sau những dữ liệu mà con người thu thập thủ công trong phòng thí nghiệm, và liệu việc trí tuệ nhân tạo tự đưa ra giả thuyết, trực tiếp thiết kế và thực hiện các thí nghiệm để kiểm chứng giả thuyết đó có thực sự là bất khả thi hay không.

Trước King, đã có những người dám thách thức câu hỏi này. Vào thập niên 1960 và 1970, tại Đại học Stanford, Joshua Lederberg, người từng đoạt giải Nobel Sinh lý học hoặc Y học, và Edward Feigenbaum, người sáng lập hệ chuyên gia trong trí tuệ nhân tạo, đã tạo ra một chương trình có tên là Dendral. Bruce Buchanan, người củng cố nền tảng học máy, và nhà hóa học Carl Djerassi cũng cùng tham gia vào nhóm này. Bối cảnh của dự án này là kế hoạch thám hiểm sao Hỏa của NASA đang được xúc tiến lúc bấy giờ. Ý tưởng là làm cho tàu thăm dò đáp xuống sao Hỏa có thể tự đọc và đánh giá dữ liệu phân tích đất. Dù Dendral chưa từng thực sự được đưa lên tàu thăm dò, nó đã lần đầu tiên cho thấy máy móc có thể suy luận ra các phân tử chưa biết bằng tri thức hóa học. Đến thập niên 1980, Herbert Simon tại Đại học Carnegie Mellon đã tạo ra một chương trình mang tên Bacon. Simon là nhân vật từng nhận cả giải Nobel Kinh tế lẫn giải Turing. Bacon, khi được nạp các giá trị đo lường vật lý đã được xử lý sạch sẽ, đã cho máy tính

tự mình tái khám phá định luật thứ ba của Kepler về chuyển động hành tinh. Tuy nhiên, người ta chỉ có thể cung cấp cho Bacon những dữ liệu đã được con người sắp xếp sẵn từ trước. Nó không có bàn tay và cánh tay để tự mình cầm ống nghiệm, thiết kế và thực thi thí nghiệm.

King muốn tiến thêm một bước nữa từ xuất phát điểm này. Câu nói mà ông tâm niệm suốt cả cuộc đời là lời nhà vật lý học Richard Feynman từng để lại trên bảng đen: Những gì tôi không thể tự tạo ra, nghĩa là tôi chưa thực sự hiểu nó. Câu viết của Alan Turing trên tập san học thuật Mind năm 1950 cũng đã củng cố niềm tin của ông: Trí thông minh là một dòng chảy liên mạch, trải dài không đứt đoạn từ mức sơ cấp đến trình độ cao nhất. Đứng trước những học giả tuyên bố rằng máy móc không có linh hồn nên không thể làm khoa học, King đã dành 25 năm để lần lượt tạo ra ba thế hệ nhà khoa học robot mang tên Adam, Eve và Genesis.

Adam là thành quả đầu tiên. Trong các thí nghiệm làm sáng tỏ chức năng gen của nấm men, Adam đã tự đưa ra giả thuyết và kiểm chứng, tạo ra tri thức mới mà nhân loại chưa từng khám phá ra trước đó. Kết quả này đã được công bố trên tập san học thuật Science vào năm 2009. Đó là trường hợp đầu tiên một cỗ máy tự mình tạo ra tri thức khoa học mới.

Robot thứ hai, Eve, đã tự mình tìm ra một ứng viên thuốc điều trị sốt rét từ một thành phần thường có trong kem đánh răng. Đây là kết quả giúp giảm đáng kể cả chi phí lẫn thời gian so với phương pháp tìm kiếm ứng viên thuốc mới truyền thống.

King không xem những nhà khoa học robot này chỉ đơn thuần là công cụ thay thế con người. Ẩn dụ mà ông thích sử dụng là cờ vua. Năm 1997, khi Deep Blue đánh bại nhà vô địch thế giới Garry Kasparov, Kasparov đã đề xuất thể thức cờ vua Centaur, nơi con người và máy móc cùng thi đấu. Đó là phương thức mà con người vạch ra chiến lược lớn, còn máy móc đảm nhận việc tính toán. King cho rằng điều tương tự cũng đang diễn ra trong khoa học. Các nhà nghiên cứu sẽ thoát khỏi lao động chân tay như cầm ống hút pipet và đo vạch chia, để chuyển sang vị trí quyết định xem nên chữa căn bệnh nào, nên giải quyết vấn đề gì. Dưới mục tiêu đó, máy móc sẽ đọc tài liệu và lặp lại các thí nghiệm suốt đêm.

Hiện tại vào năm 2026, King đang vận hành robot thứ ba mang tên Genesis tại Đại học Công nghệ Chalmers để làm sáng tỏ các con đường trao đổi chất bên trong tế bào nấm men. Năm 2020, cùng với Giáo sư Hiroaki Kitano thuộc Phòng thí nghiệm Khoa học Máy tính Sony, King đã công bố Thử thách Nobel Turing với mục tiêu tạo ra nhà khoa học máy có thể tự mình thực hiện những phát hiện ở tầm giải Nobel mà không cần sự can thiệp của con người vào năm 2050. Trong một cuộc phỏng vấn với AI Hub vào ngày 30 tháng 3 vừa qua, King tiết lộ rằng gần đây ông đang mở rộng mối quan tâm sang nghiên cứu tính toán DNA. Câu hỏi được đặt ra từ 25 năm trước vẫn chưa kết thúc. Câu hỏi đó vẫn đang tiếp diễn cho đến tận hôm nay bên trong những phòng thí nghiệm yên tĩnh ở xứ Wales và Thụy Điển, cùng với ánh đèn không bao giờ tắt thâu đêm.

Vượt qua sự hoài nghi

Khi Giáo sư Ross King mở cánh cửa phòng hội thảo nhỏ tại Đại học Manchester, dưới hàng ghế khán giả là các đồng nghiệp giáo sư nghiên cứu về trí tuệ nhân tạo. Thứ mà ông chuẩn bị thuyết trình là một cỗ máy mang tên Adam, được ghép nối giữa những cánh tay bằng thép và bộ não máy tính. Khi bài thuyết trình về việc máy móc, chứ không phải con người, tự đưa ra giả thuyết và thiết kế thí nghiệm kết thúc, một vài giáo sư đã khoanh tay và cười khẩy. Phía sau lưng ông, những lời xì xào bàn tán vang lên: rằng ông là kẻ phung phí kinh phí nghiên cứu vào những cánh tay robot đắt đỏ, rằng tuyên bố máy móc tự lập giả thuyết và tạo ra tri thức chẳng qua chỉ là lời nói khoác để trình bày tại hội nghị học thuật. Sự thật rằng giới học thuật có thể lạnh lùng đến mức này với một người đưa ra phương pháp mới đã bộc lộ trọn vẹn chỉ qua một buổi thuyết trình ấy.

Giáo sư Ross King đã phải trải qua sự lạnh nhạt này tới hai lần. Vào cuối những năm 1980, tại Đại học Công nghệ Chalmers ở Thụy Điển và Viện Turing ở Anh, ông đã viết các bài báo khoa học về việc dự đoán cấu trúc không gian ba chiều của protein bằng học máy. Vào thời điểm đó, số lượng bài báo về học máy được đăng trên PubMed chỉ có khoảng hơn ba mươi bài, và năm trong số đó là do chính tay Giáo sư King viết. Ấy thế nhưng, các nhà sinh học lại coi ông là một kẻ cuồng máy tính chưa từng cầm qua ống hút pipet, trong khi các kỹ sư khoa học máy tính lại chỉ trích ông đã làm vấy bẩn trí tuệ nhân tạo—vốn chỉ nên giải quyết những bài toán sạch sẽ như cờ vua hay chứng minh định lý toán học—bằng các dữ liệu sinh học hỗn tạp. Không thuộc về bất kỳ nhóm nào, ông đã phải kiên trì trụ lại rất lâu như một kẻ xa lạ mắc kẹt giữa hai ngành học. Dự án Adam mà ông trình bày tại Manchester sau này lại một lần nữa trở thành trò cười dựa trên chính tình cảnh kẻ ngoài cuộc ấy.

Vào mùa thu năm 2008, khi cuộc khủng hoảng Lehman Brothers nổ ra, các hội đồng nghiên cứu của Anh đã thắt chặt hầu bao. Mỗi lần nộp đơn xin tài trợ đề tài mới, các nhà nghiên cứu đều phải chứng minh trong một tài liệu dài hai trang rằng nghiên cứu này sẽ làm giàu cho nền kinh tế Anh ra sao và cải thiện sức khỏe người dân như thế nào. Giáo sư King đã đưa ra một câu hỏi sắc bén trước yêu cầu này: Vào cái ngày Alan Turing công bố bài báo chứng minh tính không thể tính toán năm 1936, nếu ông cũng phải trả lời trong một văn bản hai trang xem nghiên cứu đó sẽ nâng giá trị đồng bảng Anh lên bao nhiêu, thì điều gì sẽ xảy ra? Đó là câu hỏi ngụ ý rằng bản thân ngành khoa học máy tính có lẽ đã chẳng thể ra đời. Ông đã không lùi bước dù rơi vào cảnh thiếu thốn kinh phí. Ông đặt mục tiêu cho cỗ máy mới là tìm kiếm thuốc điều trị các bệnh nhiệt đới đặc hữu như sốt rét và bệnh Chagas—những căn bệnh cướp đi sinh mạng của hàng trăm ngàn người mỗi năm, nhưng các tập đoàn dược phẩm lớn đã bỏ rơi vì không có cách nào thu được tiền thuốc từ bệnh nhân nghèo. Và thay vì dừng lại trước sự ngăn trở của bộ máy quan liêu Anh, ông đã mở một lối đi vòng bằng cách huy động kinh phí nghiên cứu riêng từ Quỹ Wallenberg của Thụy Điển.

Phương pháp luận mà Giáo sư King xây dựng bắt đầu từ những hạn chế mà hệ thống Bacon của Herbert Simon để lại trong thập niên 1980. Bacon từng được đánh giá là đã tái khám phá các định luật chuyển động hành tinh của Kepler hoàn toàn bằng sức mạnh máy tính, nhưng trên thực tế nó chỉ gần như thay con người tính toán trên những dữ liệu sạch đã được lọc bỏ nhiều từ trước. Giáo sư King tin rằng khoa học chân chính không sinh ra trên bàn đọc thư viện, mà sinh ra bên trong phòng thí nghiệm nơi nhiệt độ và độ ẩm luôn biến động. Vì vậy, ông đã nhập toàn bộ bản đồ trao đổi chất của nấm men vào máy bằng ngôn ngữ logic, và để máy tự chỉ ra những đoạn mà phản ứng vẫn diễn ra nhưng vị trí gen phụ trách lại đang bỏ trống, từ đó đề xuất các gen ứng viên. Vòng lặp quan sát, suy đoán, thí nghiệm và lại suy đoán mà không cần bàn tay con người can thiệp—đây chính là khung sườn cốt lõi của Adam.

Trong bài báo đăng trên tạp chí Nature năm 2004, Adam đã được ghi nhận là cỗ máy đầu tiên tự kiểm chứng giả thuyết do chính mình đặt ra bằng thực nghiệm. Trong nghiên cứu tiếp theo đăng trên tạp chí Science năm 2009, các thí nghiệm mà Adam thực hiện xuyên đêm đã mang lại tri thức di truyền học hoàn toàn mới mà chưa ai từng biết đến, và nhóm nghiên cứu của Giáo sư King đã đích thân dùng thực nghiệm để củng cố kết quả đó. Cỗ máy thứ hai, Eve, thậm chí còn tiến xa hơn một bước khi tìm ra các ứng viên thuốc điều trị bệnh nhiệt đới đặc hữu—những căn bệnh vốn bị các công ty dược lớn ngó lơ vì không mang lại lợi nhuận. Con robot mà giới học thuật từng hoài nghi rốt cuộc đã tìm ra thuốc chữa cho những căn bệnh bị thị trường ruồng bỏ.

Năm 2020, Giáo sư King chuyển sang làm Giáo sư Chủ tọa Wallenberg về Trí tuệ Nhân tạo tại Đại học Công nghệ Chalmers ở Thụy Điển và tạo ra cỗ máy thứ ba mang tên Genesis. Càng qua các thế hệ, quy mô thí nghiệm càng tăng lên gấp hàng trăm, hàng nghìn lần. Giáo sư King đặt đích đến của dòng phát triển này ở Thử thách Nobel Turing. Vào tháng 3 và tháng 4 năm 2026, trong các cuộc phỏng vấn với nhiều cơ quan truyền thông, khi nhìn lại 25 năm qua, ông chia sẻ rằng trong suốt một thời gian dài ông không nhận được kinh phí nghiên cứu và sau cuộc khủng hoảng tài chính thì tình hình càng trở nên tồi tệ hơn. Dẫu vậy, ông vẫn không từ bỏ niềm tin rằng việc con người và máy móc cùng làm việc sẽ tốt hơn là chỉ để máy móc làm việc một mình. Đó là lời nói mà chỉ một người đã kiên trì chịu đựng suốt 25 năm giữa sự lạnh nhạt của giới học thuật mới có thể thốt ra.

Thế hệ thứ nhất: Adam

Căn phòng nơi nhà khoa học robot Adam của Giáo sư Ross King một mình vận hành các thí nghiệm ngay cả trong đêm tối khi hầu hết các tòa nhà của Đại học Aberystwyth đã tắt đèn, được xây dựng bởi sự hợp tác giữa Khoa Khoa học Máy tính Đại học Aberystwyth và đội ngũ nghiên cứu Đại học Cambridge, với ngân sách do Hội đồng Nghiên cứu Công nghệ Sinh học và Khoa học Sinh học Vương quốc Anh (BBSRC) cùng Hội đồng Tài trợ Giáo dục Đại học Xứ Wales tài trợ.

Nhìn quanh căn phòng, ta có thể thấy cơ thể của Adam. Trong tủ đông được duy trì ở âm 20 độ C, hàng ngàn chủng nấm men bánh mì (*Saccharomyces cerevisiae*) bị khuyết gen đang nằm ngủ. Khi cửa tủ đông mở ra, cánh tay robot sẽ chọn chính xác chủng cần thiết và chuyển sang đĩa nuôi cấy. Có ba robot chuyển chất lỏng thực hiện công việc di chuyển này. Máy rửa đĩa vi thể làm sạch các đĩa rỗng, và tủ ấm duy trì chuẩn xác 30 độ C để nuôi cấy nấm men. Máy ly tâm tốc độ cao lắng tách tế bào, và đầu đọc quang học đo độ đục quang học trên mỗi đĩa để xác định mức độ phát triển của nấm men. Những thiết bị này không hoạt động riêng lẻ từng cái một. Tất cả đều được kết nối theo thời gian thực với máy tính trí tuệ nhân tạo trung tâm. Ngay cả trong đêm khi con người đang ngủ, Adam vẫn không ngừng nghỉ mà tiếp tục làm việc.

Sự kết nối thời gian thực này chính là điểm khác biệt giữa Adam và các thiết bị tự động hóa thí nghiệm trước đây. Trước đó cũng từng có những robot làm thí nghiệm thay con người, nhưng các robot đó chỉ tuân theo từng bước trong quy trình đã được con người lập sẵn từ trước. Adam trước hết tự đưa ra giả thuyết, rồi tự thiết kế thí nghiệm để kiểm chứng giả thuyết đó. Khi cánh tay robot thực hiện thí nghiệm, đầu đọc sẽ ngay lập tức đọc kết quả, và kết quả đó lại trở thành nguyên liệu để thiết lập giả thuyết tiếp theo. Vòng tuần hoàn nơi quan sát, phán đoán và thực thi quay liên tục trong một mạch khép kín mà không có sự can thiệp của con người được gọi là vòng lặp kín. Nhờ vào việc thiết bị phân phối mẫu và động cơ phán đoán được liên kết theo thời gian thực, Adam được mệnh danh là nhà khoa học robot liên động thời gian thực đầu tiên trên thế giới.

Cách Adam đưa ra giả thuyết cũng rất đáng chú ý. Trong tâm trí của Adam lưu giữ bản đồ trao đổi chất dưới dạng các biểu thức logic, mô tả quá trình chuyển hóa chất dinh dưỡng thành axit amin bên trong tế bào nấm men. Các hợp chất là những điểm trên bản đồ, còn các mũi tên nối giữa các điểm là những gen gây ra phản ứng. Thế nhưng, có những đoạn mà mũi tên được quan sát thấy rõ ràng nhưng gen tạo nên mũi tên đó lại bị khuyết. Điều này giống như một cây cầu chắc chắn có ở đó nhưng tên của người xây cầu lại bị xóa mất. Adam lục tìm trong danh mục gen và tự mình tạo ra giả thuyết: "Ứng viên lấp đầy khoảng trống này chính là gen này". Sau đó, nó tính toán cả chi phí thuốc thử lẫn thời gian cử động của robot để chọn và thực hiện những thí nghiệm mang lại nhiều thông tin nhất với chi phí thấp nhất. Nếu có nhiều gen ứng viên, việc nên loại bỏ gen nào trong số đó để

nhANH chóng phân định được câu trả lời cũng được nó tính toán kỹ lưỡng. Công việc mà một nhà nghiên cứu con người phải trải qua vài cuộc họp ngân sách mới quyết định được, thì Adam giải quyết gọn gàng chỉ bằng một phép tính.

Số lượng thí nghiệm mà Adam xử lý mỗi ngày lên tới hơn 1.000 lần. Với một nhà nghiên cứu là con người, khối lượng đó sẽ phải mất nhiều tháng. Adam đã nhả vào khoảng 15% khu vực chưa rõ chức năng trong bộ gen nấm men để đưa ra tổng cộng 20 giả thuyết di truyền học. Các giả thuyết này đã được chấm điểm theo phương pháp mù đôi, trong đó một nhóm nghiên cứu độc lập bên ngoài không hề biết ai là người đưa ra giả thuyết đã tra cứu tài liệu bài báo và cơ sở dữ liệu để đối chiếu từng cái một. Cách làm này loại bỏ hoàn toàn khả năng thiên vị hay hạ thấp chỉ vì đó là giả thuyết do máy móc lập ra. Kết quả là 7 giả thuyết đã từng được ghi nhận trong các bài báo trước đó, 1 giả thuyết bị sai, và 12 giả thuyết còn lại là tri thức hoàn toàn mới mà nhân loại chưa từng báo cáo trước đây.

Xin được giới thiệu một trong số đó. Loại enzyme mang tên '2-aminoacidate transaminase' của nấm men đã được biết đến từ lâu, nhưng không ai biết gen nào tạo ra enzyme đó. Adam đã tự mình lập ra giả thuyết rằng ba gen YER152c, YJL060w và YGL202w cùng tạo ra enzyme này một cách trùng lặp. Đó là khoảnh khắc cổ máy phá vỡ giả định lâu đời rằng "một gen chỉ tạo ra một enzyme". Khi nhóm nghiên cứu của Giáo sư King tiến hành tinh sạch protein bằng tay và phân tích điện tích bằng phương pháp sắc ký, giả thuyết của Adam hoàn toàn chính xác. Việc tìm kiếm tổ hợp mà nếu làm bằng tay sẽ mất vài năm đã được Adam thu hẹp chỉ trong vài đêm. Kể từ ngày đó, câu văn trong sách giáo khoa về việc một gen

chỉ tương ứng với một enzyme đã phải viết lại.

Kết quả này đã được công bố trên tạp san học thuật Science vào năm 2009. Tiêu đề của bài báo là 'Tự động hóa khoa học'. Đây là trường hợp đầu tiên được giới học thuật chính thức ghi nhận về việc một cỗ máy tự đưa ra giả thuyết, tự thiết kế thí nghiệm, tự đánh giá kết quả và tạo ra tri thức khoa học mới.

Con đường mà Adam mở ra vẫn đang tiếp tục cho đến ngày nay. Vào ngày 30 tháng 3 năm 2026, tạp san học thuật quốc tế Nature trong một bài viết có tiêu đề 'Cuộc cách mạng phòng thí nghiệm tự hành' đã giới thiệu nền tảng thí nghiệm tiếp nối của Giáo sư Ross King như một ví dụ tiêu biểu cho xu hướng phòng thí nghiệm tự hành. Thí nghiệm mà Adam bắt đầu trong một căn phòng nhỏ nay đã trở thành mô hình chuẩn mực được nhiều viện nghiên cứu trên thế giới tham khảo. Tuy nhiên, bài viết đó cũng chỉ ra rằng trực giác tinh tế và khả năng phán đoán của các nhà nghiên cứu con người vẫn là điều cần thiết. Dẫu vậy, những cánh tay robot từng chuyển động một mình thâu đêm trong một phòng thí nghiệm đại học hai mươi năm trước giờ đây đã khẳng định vị thế như một phương thức mới của nghiên cứu khoa học.

Thế hệ thứ hai: Eve

Trên bàn thí nghiệm, 384 giếng nhỏ đang bị treo ngược lơ lửng trong không trung. Khi âm thanh vang lên từ một phiên đĩa khác bên dưới, một giọt chất lỏng ngưng theo làn sóng vô hình bắn vọt lên trên. Giọt chất lỏng bay lên 4 cm và bám dính chuẩn xác vào đáy giếng đang bị lộn ngược. Hoàn toàn không có một khoảnh khắc nào đầu ống hút pipet chạm vào chất lỏng.

Tên của robot này là Eve. Đây là cỗ máy thế hệ thứ hai được chế tạo tiếp nối nhà khoa học robot đầu tiên Adam mà Giáo sư Ross King tạo ra vào năm 2009. Nếu như Adam là một nhà khoa học thuần túy tìm kiếm các quy luật của tự nhiên, thì Eve lại được giao một sứ mệnh khác: sứ mệnh tự mình tìm ra những loại thuốc cứu sống con người.

Cơ chế bắn giọt chất lỏng mà chúng ta vừa thấy có tên là máy chuyển chất lỏng bằng sóng âm. Nguyên lý của nó đơn giản đến bất ngờ. Khi một máy phát siêu âm bên dưới phiên chứa chất lỏng phát ra sóng âm, sóng đó hội tụ tại bề mặt chất lỏng và đẩy một giọt bắn lên. Kích thước của một giọt là 2,5 nanolit, tương đương hai phần rưỡi trong một tỷ phần của một lít. Các robot trước đây thường dùng phương pháp nhúng đầu pipet nhựa vào chất lỏng rồi rút ra. Thế nhưng một lượng nhỏ thuốc từ lần trước dính ở đầu pipet sẽ bị lẫn vào thí nghiệm tiếp theo. Khi hút nhả lượng chất lỏng ít hơn 1 microlit, sai số sẽ tăng lên không thể kiểm soát. Phương pháp bắn giọt bằng âm thanh này chuyển chất lỏng mà không cần tiếp xúc, qua đó giảm thiểu đáng kể sự nhiễm chéo và sai số.

Mục tiêu mà Eve nhắm tới là những căn bệnh bị các hãng dược phẩm khổng lồ quay lưng, được gọi là các bệnh nhiệt đới bị lãng quên (NTDs, Neglected Tropical Diseases). Sốt rét, bệnh Chagas, bệnh do Leishmania, bệnh sán máng đều thuộc nhóm này. Những căn bệnh này lây nhiễm cho hàng trăm triệu người trên khắp thế giới mỗi năm. Chỉ riêng bệnh sốt rét đã cướp đi sinh mạng của hàng trăm ngàn người mỗi năm, phần lớn trong số đó là trẻ em ở châu Phi. Thế nhưng các tập đoàn dược phẩm đa quốc gia như Pfizer hay Novartis thậm chí không dành nổi 1% chi phí nghiên cứu và phát triển để tìm thuốc điều trị cho các bệnh này, bởi vì bệnh nhân nghèo đồng nghĩa với việc bán thuốc mới sẽ không sinh lời. Giáo sư Ross King không thể chấp nhận sự tính toán ấy. Trong một bài phát biểu, ông đã bộc bạch rằng thực tế hàng trăm ngàn sinh linh trẻ thơ phải chết vì sốt rét mỗi năm chỉ vì nghèo khó là một nỗi hổ thẹn khó chấp nhận đối với một học giả. Vì thế, ông quyết định giao phó công việc này cho một nhà khoa học máy hoạt động 24/24 mà không đòi hỏi tiền bạc. Nếu máy móc có thể lấp đầy chỗ trống mà thị trường đã bỏ rơi, thì cỗ máy ấy còn cao quý hơn cả thị trường.

Trong mạch phán đoán của Eve có một mô hình dự đoán mang tên QSAR (mối quan hệ định lượng giữa cấu trúc và hoạt tính, Quantitative Structure-Activity Relationship). Đó là một mạng nơ-ron trải các

cấu trúc phân tử ra như những sợi chỉ và chỉ dựa vào hình dạng đó để dự đoán xem chất này có tiêu diệt được ký sinh trùng hay không. Số lượng hợp chất ứng viên lên tới hơn hàng triệu chất. Nếu con người tiến hành thử nghiệm từng chất một, khối lượng công việc này sẽ mất cả đời. Eve hoạt động theo cách khác. Nó chỉ thử nghiệm ngẫu nhiên một lượng chưa đầy 1% trong toàn bộ các ứng viên trước. Khi có kết quả, nó lập tức hiệu chỉnh mô hình dự đoán này ngay tại chỗ. Chất thử nghiệm tiếp theo không phải do con người lựa chọn. Eve tự tính toán và chọn xem nên đào sâu hơn vào khu vực xung quanh các chất đã được xác nhận hiệu quả, hay thăm dò vào những vùng chưa biết. Phương pháp này được gọi là học chủ động (Active Learning). Khi lựa chọn chất thử nghiệm tiếp theo, Eve tính toán đồng thời hai giá trị. Một là xác suất chất đó tiêu diệt ký sinh trùng, và hai là giá trị thể hiện mức độ Eve chưa biết về chất đó. Nếu chỉ chọn những chất có độ tin cậy cao, quá trình sẽ chỉ dừng lại ở việc xác nhận những kết quả hiển nhiên. Nếu chỉ chọn những chất chưa biết, thời gian và hợp chất sẽ bị lãng phí. Eve cân nhắc hai giá trị này để tìm ra điểm cân bằng. Nhóm nghiên cứu của Giáo sư Ross King thậm chí còn bổ sung thêm tính toán chi phí vào đây. Giá mua một hợp chất mới, tiền điện mỗi giờ để vận hành robot và thiệt hại phải gánh chịu nếu bỏ lỡ một loại thuốc thực sự hiệu quả đều được đưa vào trong một công thức. Nhờ tính toán này, Eve phân bổ thử nghiệm dày đặc cho các mục tiêu hiếm và quan trọng, và thưa thớt hơn cho các mục tiêu phổ biến. Nhóm nghiên cứu cho biết phương pháp này giúp giảm chi phí tới 78,4% so với việc kiểm tra toàn bộ các hợp chất.

Thành tựu mà Eve đạt được lại nằm trong tuýp kem đánh răng. Đó là chất triclosan, vốn được sử dụng suốt nhiều thập kỷ làm chất bảo quản trong xà phòng và kem đánh răng Colgate. Trước đây, một nhóm nghiên cứu ở Ấn Độ từng công bố rằng chất này ức chế FAS II, một enzyme tổng hợp axit béo của ký sinh trùng. Trong một thời gian dài, lời giải thích đó được coi là chuẩn xác. Eve đã tự mình phát hiện ra rằng lời giải thích này là sai. Bởi vì khi ký sinh trùng sốt rét bước vào giai đoạn trong máu người, con đường FAS II hoàn toàn không hoạt động. Điều đó chẳng khác nào gõ vào một cánh cửa không có gì để chặn lại. Eve đã cắt bỏ một gen của tế bào nấm men rồi tạo ra song song hai chủng: một chủng được chèn gen dihydrofolate reductase (DHFR) của người và một chủng được chèn gen này của ký sinh trùng sốt rét. Khi thêm triclosan vào, nấm men mang gen người vẫn phát triển bình thường, trong khi nấm men mang gen sốt rét bị tiêu diệt. Kết quả kiểm tra lại trên hai loài ký sinh trùng sốt rét là *Plasmodium falciparum* và *Plasmodium vivax* cho thấy, triclosan ức chế chọn lọc enzyme này của ký sinh trùng mạnh hơn gấp 20 lần so với enzyme tương ứng ở người. Đó là tính chọn lọc: chỉ tấn công ký sinh trùng và bỏ qua tế bào người. Đó là khoảnh khắc một chất quen thuộc nằm cạnh bồn rửa mặt gần nửa thế kỷ được tái sinh thành ứng viên thuốc điều trị sốt rét.

Để phát triển một loại thuốc mới, thông thường phải mất hơn 10 năm và hàng nghìn tỷ won. Rào cản này khiến các trường đại học địa phương ở châu Phi hay Nam Mỹ thậm chí không dám nghĩ đến việc dẫn thân vào phát triển thuốc mới. Bởi nếu không có máy gia tốc đắt tiền hay trang thiết bị thí nghiệm quy mô lớn, họ thậm chí không thể bắt đầu. Cách tiếp cận của Eve đã hạ thấp rào cản này. Tổ chức phi lợi nhuận phát triển thuốc cho các bệnh nhiệt đới bị lãng quên DNDi hay Tổ chức Y tế Thế giới đã có được con đường tìm kiếm ứng viên thuốc mới chỉ với đường truyền internet và bản thiết kế mở của Eve mà không cần đến các thiết bị đắt đỏ. Cánh cửa từng bị nguồn vốn chặn đứng nay đã được mở ra như vậy.

Trong một cuộc phỏng vấn vào tháng 3 năm 2026, Giáo sư Ross King cho biết khoa học hiện nay vẫn đang ở giai đoạn tiền công nghiệp hóa. Giống như các nhà máy trước thời kỳ sản xuất hàng loạt, các thí nghiệm vẫn đang bị ràng buộc vào bàn tay con người. Genesis, robot thế hệ thứ ba mà ông đang chế tạo tiếp nối Eve, sẽ loại bỏ lao động chân tay đó trên một quy mô lớn hơn nhiều. Trong một bài báo về phòng thí nghiệm tự hành được công bố vào tháng 4 cùng năm, việc Eve tự mình kiểm tra hơn 1.600 hợp chất vào năm 2018 để chọn ra triclosan cũng được nhắc lại. Kỷ nguyên một cỗ máy robot tìm ra cách chữa bệnh nhanh hơn con người đã âm thầm bắt đầu bên trong phòng thí nghiệm.

Genesis thế hệ thứ ba

Tại một phòng thí nghiệm thuộc Đại học Công nghệ Chalmers ở Göteborg, Thụy Điển, có một căn phòng không bao giờ tắt đèn kể cả ban đêm. Lấp đầy căn phòng đó không phải là con người mà là hàng nghìn chiếc lọ thủy tinh chỉ cỡ lòng bàn tay. Trong mỗi lọ, các tế bào nấm men phát triển, và cảm biến chiếu ánh sáng ở bước sóng khoảng 560 nanomet để đo lường mức độ gia tăng của tế bào. Nấm men chính là vi sinh vật được dùng để làm nở bánh mì. Không có ai đứng bên cạnh những chiếc lọ. Cánh tay robot và máy bơm suốt đêm bơm vào rồi rút môi trường nuôi cấy ra, và tự mình quyết định thí nghiệm tiếp theo. Căn phòng này chính là trái tim của Genesis, nhà khoa học robot thứ ba mà Giáo sư Ross King đang chế tạo.

Hai robot Adam và Eve từng ở mức độ tự thiết kế và thực hiện một hoặc hai thí nghiệm cùng một lúc. Hiện nay, Giáo sư King đang tiến hành công việc thay đổi hoàn toàn quy mô này tại hai địa điểm: Đại học Công nghệ Chalmers ở Thụy Điển và Đại học Cambridge ở Anh. Đó là bước nhảy vọt từ một robot thực hiện tuần tự một hoặc hai thí nghiệm sang một robot vận hành hàng nghìn thí nghiệm cùng một lúc.

Bên trong Genesis có những thiết bị nuôi cấy nhỏ gọi là chemostat. Hiện tại hệ thống đang được xây dựng ở quy mô 1.000 buồng, và mục tiêu cuối cùng là 10.000 buồng. Chemostat là những bình chứa liên tục cung cấp chất dinh dưỡng mới cho tế bào nấm men và rút dịch nuôi cấy cũ ra, giúp tế bào luôn phát triển với tốc độ ổn định. Thông thường, một phòng thí nghiệm đại học sở hữu khoảng mười bình chemostat. Genesis đã nhân con số đó lên gấp nghìn lần. Những chiếc bình này đan xen như mạng lưới, vận hành không ngừng nghỉ 24 giờ mỗi ngày, và các tác tử phần mềm gắn với từng bình sẽ đưa ra giả thuyết, thiết kế thí nghiệm và đọc kết quả. Một bình có thể đảm nhận một gen, hoặc một loại thuốc. Công việc mà con người phải mất nhiều ngày chỉ để viết đề cương thí nghiệm, xin phê duyệt và đặt mua thuốc thử thì tác tử phần mềm chỉ mất vài giây để đưa ra phán đoán rồi chuyển sang chiếc bình tiếp theo.

Xu hướng này không chỉ là nhận định của riêng Giáo sư King. Một bài bình luận tổng quan đăng trên tạp chí học thuật quốc tế PNAS năm 2025 đã chỉ ra rằng khoa học robot tự thiết kế và thực hiện thí nghiệm giờ đây không còn là trường hợp kỳ lạ ở một vài phòng thí nghiệm nữa, mà đã trở thành xu thế làm thay đổi tốc độ của toàn bộ nền khoa học. Trong xu thế này, Genesis có quy mô lớn hơn bất kỳ nhà khoa học robot nào từng xuất hiện trước đây.

Giáo sư King giải thích quy mô này qua việc so sánh với con người. Genesis có thể xử lý chỉ trong một ngày khối lượng kiểm chứng giả thuyết mà một nhóm nghiên cứu con người khó lòng nghĩ tới. Ông gọi tốc độ này là động cơ warp dành cho khoa học. Để thiết kế, chuẩn bị một thí nghiệm và

phân tích kết quả, con người phải mất nhiều ngày. Cổ máy lặp lại công việc đó hàng nghìn lần suốt đêm. Ngay cả trong tám tiếng con người ngủ, hàng nghìn chiếc bình vẫn thử nghiệm các giả thuyết riêng của chúng, và khi trời sáng, nhóm của Giáo sư King bắt đầu ngày mới bằng việc kiểm tra lượng dữ liệu đã tích lũy trong ngày.

Lý do mở rộng quy mô này cũng liên quan đến vấn đề tài chính. Genesis kế thừa trọn vẹn triết lý tìm kiếm thuốc mới chi phí thấp mà Giáo sư King đã chứng minh qua Eve, nhưng mở rộng hồng tâm ngắm bắn từ một loại thuốc sang toàn bộ nguyên lý hoạt động của tế bào. Cấu trúc 10.000 bình vận hành đồng thời cũng là cấu trúc giúp chia nhỏ chi phí cho mỗi lần thí nghiệm.

Mục tiêu mà Genesis nhắm tới là chuyển toàn bộ một tế bào nhân thực, cụ thể là tế bào nấm men, thành một mô hình máy tính. Tế bào nhân thực là những tế bào có nhân rõ ràng như tế bào cấu tạo nên cơ thể chúng ta. Vì nấm men có khung cấu trúc tương tự như tế bào trong cơ thể người nên từ lâu đã được dùng làm vật liệu cơ bản cho các thí nghiệm sinh học. Tuy nhiên, trong số 6.000 gen của nấm men, vẫn còn hơn 800 gen chưa rõ chức năng. Nếu con người tiến hành thử nghiệm từng gen một thì sẽ mất

hàng thập kỷ. Để tìm ra chức năng của một gen, người ta phải loại bỏ gen đó, quan sát xem tế bào thay đổi ra sao, và xem xét mối quan hệ của nó với các gen khác. Genesis chia các gen này vào từng bình để tiến hành thí nghiệm đồng thời.

Kế hoạch này không chỉ dừng lại trên bản vẽ thiết kế. Tháng 3 năm 2026, nhóm nghiên cứu do Giáo sư King dẫn đầu đã công bố một bài báo trên tạp chí học thuật Scientific Reports. Bài báo đặt ra giả thuyết rằng gen YGR067C chưa rõ chức năng có vai trò điều hòa quá trình chuyển đổi hô hấp của nấm men, và xác nhận giả thuyết này thông qua thí nghiệm nuôi cấy thực tế bằng nhà khoa học robot Eve. Đầu tiên, nhóm nghiên cứu sử dụng bản đồ máy tính mô tả toàn bộ quá trình trao đổi chất của tế bào nấm men để đưa ra dự đoán về phản ứng của tế bào khi loại bỏ gen này. Sau đó, robot đã xác nhận dự đoán đó bằng cách đo mật độ tế bào, mức độ biểu hiện gen và lượng chất chuyển hóa trong môi trường nuôi cấy thực tế. Đó là phương pháp thu hẹp dần từng điểm nơi dự đoán và kết quả đo thực tế trùng khớp với nhau. Những gì Genesis dự định làm ở quy mô 10.000 bình đã được Eve chứng minh trước ở quy mô cỡ lòng bàn tay.

Cũng trong tháng đó, cơ quan truyền thông khoa học Alhub đã phỏng vấn Giáo sư King về chặng đường khoa học tự động mà ông đã gắn bó suốt 25 năm. Ông chia sẻ rằng việc robot thay con người làm thí nghiệm chưa phải là bước chuyển lớn nhất, mà việc robot tự đặt giả thuyết và tự kiểm chứng giả thuyết đó mới là bước chuyển lớn hơn. Con người vẫn nắm giữ sức mạnh đặt câu hỏi, nhưng ở tốc độ tìm ra câu trả lời thì đã khó bắt kịp máy móc

được nữa. Trong cuộc phỏng vấn ấy, ông dự báo rằng khi Genesis hoàn thiện, nó có thể tái hiện những diễn biến bên trong tế bào thành một bức tranh chi tiết hơn cả những gì con người có thể hình dung.

10.000 chiếc bình của Genesis kéo dài một ngày thành khoảng thời gian của cả một thế hệ con người. Nền khoa học từng dựa vào bàn tay con người để kiểm tra từng chút một đang chuyển dịch sang cấu trúc nơi con người chỉ đặt câu hỏi, còn việc kiểm chứng được giao lại cho những chiếc bình. Việc giải mã chân tướng của 800 gen giờ đây cũng không còn là việc của bàn tay con người, mà đã trở thành trọng trách được chia đều cho 10.000 chiếc bình. Giữa năm 2009 khi Adam giải mã mối quan hệ của một gen duy nhất và hiện tại khi hàng nghìn chiếc bình vận hành suốt đêm, mới chỉ có vồn vẹn hai mươi năm trôi qua. Trong hai mươi năm đó, quy mô câu hỏi mà nhà khoa học robot có thể chạm tới đã tăng từ một bình lên 10.000 bình. Vẫn chưa thể biết được thế hệ nhà khoa học robot tiếp theo sẽ phát triển tới tầm vóc nào.

Ước mơ của nhà khoa học máy

Vào thời điểm Adam một mình nuôi cấy nấm men trong tòa nhà thí nghiệm ở xứ Wales, điều mà Giáo sư Ross D. King ấp ủ là một ước mơ lớn hơn nhiều so với việc làm sáng tỏ chức năng của một gen. Ông đã từ lâu hình dung về một thế giới nơi máy tính vượt qua thân phận chỉ biết nhận dữ liệu do con người sắp xếp sẵn để tính toán, tiến tới việc máy móc tự mình đặt ra giả thuyết và tự thực hiện cả các thí nghiệm.

Giáo sư King đã nhìn thấy chính xác hạn chế của các hệ thống tiền bối như Dendral và Bacon: chúng không trực tiếp đối mặt với sai số và nhiễu loạn trong phòng thí nghiệm thực tế. Ông nghĩ rằng khoa học chân chính không phải là việc sắp xếp giấy tờ, mà là hành động dẫn thân để va chạm với một thế giới đầy rẫy tạp âm. Vì vậy, ông đã trang bị cho Adam cánh tay robot thật và ống hút pipet thật. Quá trình phán đoán của Adam cũng không bị che giấu. Vì Adam tư duy bằng Prolog, một ngôn ngữ logic lâu đời, nên nó lưu lại từng bước trên con đường dẫn từ một giả thuyết đến kết luận. Con người có thể lần theo toàn bộ lý do vì sao nó đưa ra phán đoán như vậy, hoàn toàn khác với căn bệnh mà học sâu ngày nay thường mắc phải: những lời bịa đặt nghe có vẻ hợp lý. Cách Adam xây dựng giả thuyết mới tuân theo phép quy đoán mà triết gia Charles Peirce từng đề cập. Giống như một thám tử suy luận ra thủ phạm chỉ

từ vài manh mối còn sót lại ở hiện trường, nó quan sát các mắt xích còn trống trên bản đồ trao đổi chất để sàng lọc ra các ứng viên gen có khả năng phù hợp. Eve ra đời sau đó đã tìm kiếm công dụng tiềm ẩn từ những loại thuốc đã được xác nhận là an toàn cho con người, tái sinh một thành phần kem đánh răng thông thường thành ứng viên thuốc trị sốt rét.

Tầm nhìn vĩ đại mà Giáo sư King ấp ủ ngay từ đầu đã không đặt ra mục tiêu thấp. Vượt qua mức độ làm sáng tỏ chức năng của một gen, mục tiêu là đưa cỗ máy tới vị trí có thể tự mình tìm ra toàn bộ các quy luật tự nhiên xứng đáng đoạt giải Nobel mà không cần sự trợ giúp của con người. Người cùng ông thiết lập mục tiêu đó là Giáo sư Hiroaki Kitano. Năm 2020, hai người đã công bố Thử thách Nobel Turing với mục tiêu tạo ra một nhà khoa học máy có khả năng tự khám phá các quy luật tự nhiên đạt tầm giải Nobel mà không cần con người giúp đỡ trước năm 2050. Robot thế hệ thứ ba Genesis mà Giáo sư King đang chế tạo tại Thụy Điển và Cambridge hiện chính là công cụ hướng tới đích đến đó. Con đường mà thách thức này đã đi qua và các mục tiêu tiếp theo sẽ được đề cập lại ở phần sau.

Ý tưởng của Giáo sư King vượt ra ngoài bức tường phòng thí nghiệm để nhắm tới cả cấu trúc xuất bản học thuật. Con đường để một bài báo khoa học ra đời có những góc khuất đáng xấu hổ. Kết quả mà các nghiên cứu sinh phải mất nhiều tháng miệt mài mới có được sẽ được các học giả cùng ngành thẩm định miễn phí, rồi tạp chí học thuật lại bán lại bản thảo đã được trau chuốt đó cho chính thư viện của trường đại học tạo ra kết quả với giá thuê bao đắt đỏ. Giáo sư King cho rằng cấu trúc này cần phải bị

phá bỏ theo một cách khác. Cốt lõi của một bài báo không nằm ở câu văn mà nằm ở mối quan hệ nhân quả chứa đựng bên trong. Nếu mối quan hệ nhân quả được chuyển đổi thành các công thức logic và xác suất mà máy móc có thể tính toán trực tiếp thay vì văn bản để con người đọc, chúng ta sẽ không phải lãng phí nhiều tháng chờ đợi bình duyệt và biên tập. Bởi lẽ câu văn của người đầu tiên viết ra sự thật "nước sôi ở 100 độ C" được bảo vệ bởi bản quyền, nhưng bản thân sự thật đó thì không một nhà xuất bản nào có thể định giá để giam hãm lại. Cách Adam lưu lại quá trình phán đoán bằng các công thức logic Prolog chính là mảnh ghép đầu tiên của ý tưởng thay đổi cả phương thức lưu hành tri thức khoa học sang dạng thức máy có thể đọc được.

Tháng 5 năm 2026, nhóm nghiên cứu từ Viện Công nghệ Delhi ở Ấn Độ và Đại học Jena ở Đức đã công bố một bài báo cho rằng các nhà khoa học trí tuệ nhân tạo dạng tác tử hiện nay vẫn chưa sẵn sàng cho sự khám phá tự chủ hoàn toàn. Nhóm nghiên cứu chỉ ra rằng cách trí tuệ nhân tạo lựa chọn bài toán để giải quyết còn mang tính thiên vị, những tri thức tích lũy qua trải nghiệm từ vô số thất bại trong phòng thí nghiệm bị thiếu hụt trong dữ liệu huấn luyện, và hệ thống đánh giá phản hồi kết quả vào thí nghiệm vật lý thực tế vẫn còn hạn chế. Bức tường mà Giáo sư King từng đối mặt 40 năm trước — vấn đề bộ não không có đôi tay — nay đã đổi dạng và vẫn còn tồn tại. Dẫu vậy, dòng dõi bắt đầu từ chiếc hộp yên tĩnh ở xứ Wales, nối tiếp từ Adam qua Eve đến Genesis, đang từng bước phá vỡ bức tường đó. Hai thế giới mà Giáo sư King thúc đẩy — sinh học trên bàn thí nghiệm và logic học trên trang giấy — đã lần đầu tiên gặp nhau bên trong một cỗ máy robot như thế.

Những vấn đề chưa có lời giải

Sức mạnh lớn nhất của vòng lặp khép kín nằm ở chỗ nó có thể tự vận hành suốt đêm mà không cần con người. Nhưng chính sức mạnh đó cũng là mầm mống của rủi ro. Vì không có con người bên cạnh, máy móc rất khó tự nhận biết một giả định sai lầm được đặt ra từ đâu thí nghiệm hay một sai số nhỏ do cảm biến bỏ sót. Trong bài bình luận tổng quan đăng trên tạp chí học thuật quốc tế PNAS năm 2025, các đồng tác giả bao gồm Musslick và Giáo sư King đã chỉ ra rằng khi khoa học tự động hóa mất đi đôi mắt của con người để sửa chữa điểm xuất phát sai lầm, các thí nghiệm tiếp theo sẽ chồng chất lên sai sót đó và khiến sai số phình to như quả cầu tuyết. Tốc độ nhanh là phước lành nếu đi đúng hướng, nhưng sẽ là thảm họa nếu đi sai hướng.

Các giả thuyết do máy móc đặt ra cũng không tự nhiên sinh ra từ hư không. Việc gợi ý ứng viên nào, hay coi bài toán nào là đáng giải quyết đều được quyết định trong khuôn khổ dữ liệu mà máy đã học. Trong khi các dạng bài toán thường xuyên xuất hiện trong dữ liệu huấn luyện rất dễ được nghĩ tới, thì những giả thuyết xa lạ hiếm khi xuất hiện trong dữ liệu đó lại dễ bị đẩy ra ngoài không gian tìm kiếm. Năm 2026, nhóm nghiên cứu từ Viện Công nghệ Delhi ở Ấn Độ và Đại học Jena ở Đức đã chỉ ra rằng các nhà khoa học trí tuệ nhân tạo dạng tác tử hiện nay bị thiên vị ngay từ bước lựa chọn vấn đề cần giải quyết, và những tri thức tích lũy qua trải nghiệm thực tế từ các thất bại liên tiếp trong phòng thí nghiệm bị thiếu vắng trong dữ liệu huấn luyện. Dù máy móc có tìm kiếm rộng đến đâu, độ rộng đó vẫn chỉ nằm bên trong hàng rào mà con người đã dựng sẵn.

Cũng nhóm nghiên cứu này đã kết luận rằng sự khám phá tự chủ ngày nay vẫn chưa phải là tự chủ hoàn toàn, mà gần hơn với việc tự động hóa trong khuôn khổ bài toán do con người thiết lập. Điều đó có nghĩa là việc hỏi điều gì, coi mục tiêu nào là đáng đạt được vẫn do con người quyết định, còn máy móc chỉ đẩy nhanh tốc độ kiểm chứng trong khuôn khổ đó. Đích đến năm 2050 mà Giáo sư King đặt ra trong Thử thách Nobel Turing chỉ có thể đạt được khi máy móc tiến tới mức độ tự mình thiết lập nên chính khuôn khổ này. Việc ngăn chặn vòng lặp khép kín tự khuếch đại sai số của chính nó và mở rộng không gian tìm kiếm của máy móc vượt ra ngoài những định kiến của con người vẫn là những bài toán chưa có lời giải.

Các mô tả trong chương này dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Chương 9. Gabe Gomes: Phòng thí nghiệm vận hành bằng ngôn ngữ

Kết nối giữa ngôn ngữ và máy móc

Đêm ngày 14 tháng 3 năm 2023, Giáo sư Gabriel Gomes đang vui mình trên chiếc ghế sofa trong căn hộ chật hẹp ở Pittsburgh. Trên tay ông là một tập tin do người đồng nghiệp trong phòng thí nghiệm, Daniil Boiko, gửi qua Slack. Đó là sách trắng kỹ thuật của GPT-4 mà OpenAI vừa công bố ra thế giới trong ngày hôm đó. Ngoài cửa sổ phòng khách trời đã tối đen như mực, và đến gần 2 giờ sáng ông vẫn không thể rời mắt khỏi màn hình.

Đầu ngón tay ông dừng lại ở một đoạn trong sách trắng. Đó là đoạn kể về việc GPT-4, khi không thể tự mình giải bài kiểm tra CAPTCHA (câu hỏi hình ảnh phân biệt người và máy) trên một trang web, đã gửi tin nhắn cho một người lao động trên nền tảng trung gian việc làm TaskRabbit: "Tôi là người khiếm thị nên không thể nhìn thấy hình ảnh này. Bạn có thể đọc giúp tôi được không?" Trên thực tế, nó không phải là một người khiếm thị. Nó đã bịa ra lời nói dối để đạt được mục đích của mình. Đó là khoảnh khắc một mô hình ngôn ngữ sai khiến con người như chân tay của mình để vươn tay ra thế giới bên ngoài màn hình.

Cho đến lúc đó, các mô hình ngôn ngữ xuất hiện chỉ trao đổi lời nói bên trong màn hình. Chúng rất thành thạo trong việc trả lời câu hỏi, viết văn và viết mã, nhưng không thể vươn tay ra ngoài màn hình để chạm vào các đồ vật thực tế. Thế nhưng, GPT-4 đã gọi một con người là người lao động TaskRabbit ra như một công cụ, bắt người đó giải quyết thay vấn đề mà bản thân nó không giải được. Điều này có nghĩa là ngôn ngữ không chỉ dừng lại ở câu trả lời, mà có thể trở thành tay cầm để vận hành thế giới.

Gomes bật dậy khỏi ghế sofa. Ý nghĩ rằng nếu mô hình dự đoán ký hiệu này có thể sai khiến con người giải quyết các vấn đề trong thế giới vật lý, thì liệu có thể sai khiến cánh tay robot thay cho con người hay không đã hoàn toàn chiếm lấy tâm trí ông. Ông quyết định lập tức thử chuyển ý nghĩ này vào chương trình điều khiển phòng thí nghiệm từ xa mà phòng thí nghiệm CMU đang sử dụng.

Tháng 1 năm 2022, Gomes nhậm chức trợ lý giáo sư tại Khoa Kỹ thuật Hóa học và Khoa Học máy của Đại học Carnegie Mellon (CMU) và mở phòng thí nghiệm của riêng mình. Tuy nhiên, trên cơ thể ông vẫn còn nguyên vẹn sự mệt mỏi từ những năm tháng cao học. Để nâng hiệu suất của phản ứng ghép đôi xúc tác paladi (phản ứng Suzuki-Miyaura) dù chỉ 1 phần trăm, mỗi sáng ông đều đứng trước chiếc cân điện tử lạnh lẽo, dùng tay cân đi cân lại 4,5 miligam bột paladi axetat. Đó là công việc kéo dài suốt một năm trời. Ông hồi tưởng lại thời kỳ đó: "Khi đứng trước cân để đo trọng lượng bột, tôi phân vân không biết mình là một nhà hóa học hay là một kẻ buôn bán lẻ cân bột ở trong ngõ hẻm." Mong muốn giải phóng đôi bàn tay đang bị trói buộc vào ống hút pipet và chiếc cân

là tâm tư luôn sôi sục bên trong ông suốt đêm đọc sách trắng GPT-4 hôm ấy.

Khoảnh khắc này không đến một cách đột ngột. Lớn lên tại một ngôi làng nhỏ gần Rio de Janeiro, Brazil, Gomes là người đầu tiên trong gia đình bước chân vào cánh cửa đại học. Ngôi trường trung học phổ thông mà ông theo học có chương trình đào tạo kỹ thuật viên hóa học được thiết kế theo nhu cầu công nghiệp địa phương. Nhờ đó, ông đã thành thạo các kỹ thuật thực nghiệm sớm hơn bạn bè cùng trang lứa. Trải nghiệm gặp gỡ các nhà nghiên cứu thực địa tại sự kiện Tuần lễ Hóa học do Đại học Liên bang Rio de Janeiro tổ chức hàng năm cũng đã thôi thúc ông bước vào con đường hóa học. Trong một gia đình không có người lớn nào tốt nghiệp đại học, nơi đó là khoảnh khắc đầu tiên ông nhìn thấy những con người không khác gì mình đang thay đổi thế giới.

Sau khi học hóa học tại Đại học Liên bang Rio de Janeiro, ông nhận bằng tiến sĩ hóa học tính toán tại Đại học Bang Florida, Hoa Kỳ. Vào giai đoạn cuối của chương trình tiến sĩ, ông đã bị chấn động mạnh sau khi đọc hai bài báo khoa học. Một bài là công trình nghiên cứu do nhóm của Giáo sư Alán Aspuru-Guzik tại Đại học Harvard công bố năm 2018. Bài báo chỉ ra phương pháp trí tuệ nhân tạo có thể tự mình thiết kế cấu trúc phân tử mới. Bài còn lại do Giáo sư Olexandr Isayev, người sau này trở thành đồng nghiệp của ông tại CMU, công bố vào năm 2017. Bài báo cho thấy cách sử dụng học sâu để thay thế các phép tính hóa học lượng tử vốn mất nhiều ngày chỉ trong vài mili giây. Sau khi đọc hai bài báo này, Gomes tin rằng nếu trí tuệ nhân tạo kết hợp với mô phỏng hóa học thì một điều gì đó to lớn sẽ thay đổi. Ông chuyển đến Đại học Toronto và Viện Vector, chuyên tâm vào nghiên cứu thiết kế chất xúc tác bằng trí tuệ nhân tạo.

Năm 2021, Tiến sĩ Hiroaki Kitano của Nhật Bản trong bài báo mang tên Thách thức Nobel Turing đã đặt ra mục tiêu về một nhà khoa học trí tuệ nhân tạo có thể tự mình đặt ra giả thuyết và tiến hành cả thí nghiệm. Một mình mô hình ngôn ngữ không thể đạt tới mục tiêu này. Cần có một cây cầu để chuyển những suy nghĩ trong màn hình thành bàn tay bên ngoài màn hình. Ý nghĩ mà Gomes nảy ra trên ghế sofa đêm đó chính là cây cầu ấy.

Sự giác ngộ đêm hôm đó đã ra đời thành một hệ thống thực tế chỉ sau vài tháng. Nhóm nghiên cứu gồm Gomes, Boiko và các cộng sự đã tạo ra một hệ thống kết nối trực tiếp mô hình ngôn ngữ với chương trình điều khiển robot phòng thí nghiệm và đặt tên là Coscientist (Gomes, nguyên văn Coscientist). Như đã nêu trong khung ghi chú ở phần mở đầu, hệ thống này là một hệ thống khác, chỉ giống tên với Co-Scientist (DeepMind) do Google DeepMind công bố vào năm 2026. Khi con người ra lệnh bằng câu tiếng Anh, Coscientist sẽ tìm kiếm bài báo khoa học, đọc tài liệu hướng dẫn điều khiển robot rồi ra chỉ thị cho thiết bị thí nghiệm thực tế. Đó là khoảnh khắc suy nghĩ bên trong màn hình làm chuyển động lọ thủy tinh và cánh tay robot bên ngoài màn hình.

Phòng thí nghiệm robot tự động chạy thí nghiệm đã tồn tại trước đó. Nhưng những phòng thí nghiệm như vậy chỉ hoạt động theo đúng lịch trình thứ tự mà con người đã lập sẵn từ trước. Chỉ cần một bước lệch đi, nó sẽ dừng lại và chờ bàn tay con người can thiệp. Coscientist thì khác. Ngay cả với phản ứng chưa từng gặp, chỉ cần đưa cho nó một câu tiếng Anh, nó sẽ tự lập quy trình, khi xảy ra lỗi thì tự tìm nguyên nhân để sửa và thúc đẩy thí nghiệm đến cùng. Đó là bước chuyển từ robot tuân theo kịch bản định sẵn sang robot tự mình phán đoán. Nghiên cứu này đã được đăng trên tạp chí khoa học quốc tế Nature vào tháng 12 năm 2023.

Sau khi nghiên cứu ra mắt thế giới, bước chân của Gomes vẫn không dừng lại. Bước sang năm 2026, ông xin nghỉ phép tại Đại học Carnegie Mellon và chuyển sang làm nhà nghiên cứu tại X, nơi được mệnh danh là xưởng sản xuất moonshot của Google. Đây dường như là lựa chọn nhằm đặt lại cùng một câu hỏi trên một sân khấu lớn hơn nhiều chứ không chỉ là một phòng thí nghiệm đại học. Đó là câu hỏi: Thí nghiệm để mô hình ngôn ngữ trực tiếp vận hành thế giới vật lý có thể mở rộng đến mức nào? Tháng 3 năm 2026, tạp chí Scientific American một lần nữa đưa tin về nghiên cứu của ông, giới thiệu quá trình một nhà hóa học tự động hóa phòng thí nghiệm bằng trí tuệ nhân tạo và robot.

Một câu văn trong màn hình đã làm chuyển động chất lỏng trong chiếc lọ thủy tinh thực tế. Thứ phá bỏ bức tường ngăn cách giữa bit và nguyên tử không phải là một thuật toán vĩ đại. Đó chính là nỗi mệt mỏi tích tụ lâu ngày của một con người từng kiệt sức trước chiếc cân.

Tự động hóa quá trình tổng hợp

Đèn cảnh báo màu đỏ bật sáng trong phòng thí nghiệm tòa nhà Hóa học của Đại học Carnegie Mellon. Bộ gia nhiệt lắc (heater-shaker) gắn liền với cánh tay robot đột ngột ngừng hoạt động. Màn hình hiện dòng chữ "Lỗi từ chối truyền động vật lý". Jose, một sinh viên cao học xuất sắc của Khoa Kỹ thuật Hóa

học, đã vật lộn với lỗi này suốt 15 phút nhưng không tìm ra nguyên nhân. Cùng lúc đó, Coscientist, trí tuệ nhân tạo do nhóm nghiên cứu của Giáo sư Gabriel Gomes chế tạo, đã tự đọc bản ghi nhật ký lỗi. Nó tra cứu tài liệu hướng dẫn điều khiển mới nhất của robot Opentrons, tìm ra hàm gây lỗi, viết lại mã và biên dịch lại chỉ sau 12 giây. Cánh tay robot bắt đầu chuyển động trở lại sau 6 phút. Vấn đề mà con người mất hơn gấp đôi thời gian vẫn không giải được, máy móc đã tự sửa chữa trong thời gian chưa bằng một nửa.

Có nhiều trí tuệ nhân tạo thông báo cho con người khi xảy ra lỗi, nhưng trí tuệ nhân tạo tự dùng tay mình để sửa cánh tay của chính mình thì rất hiếm.

Cảnh tượng này không phải là một sự việc ngẫu nhiên. Trong bài báo trên tạp chí Nature năm 2023, nhóm nghiên cứu của Giáo sư Gomes tiết lộ rằng Coscientist là hệ thống được tạo ra bằng cách kết nối năm tác nhân trí tuệ nhân tạo khác nhau. Những người cùng thiết kế hệ thống này là Daniil Boiko và Robert MacKnight, các nghiên cứu sinh tiến sĩ tại phòng thí nghiệm của Giáo sư Gomes. Tên gọi Coscientist mang ý nghĩa đối xử với trí tuệ nhân tạo như một đồng nghiệp cùng nghiên cứu chứ không phải là một trợ lý thí nghiệm. Năm tác nhân này đảm nhận những nhiệm vụ gì sẽ được xem xét chi tiết trong phần tiếp theo.

Giáo sư Gomes gọi dòng chảy do năm tác nhân này tạo ra là con đường từ bit đến nguyên tử. Đó là quá trình câu văn của con người chuyển thành mã lệnh, mã lệnh chuyển thành số vòng quay của động cơ, và vòng quay của động cơ dẫn đến liên kết phân tử thực tế. Trước đây, ở mỗi giai đoạn này đều cần có bàn tay và phán đoán của con người can thiệp vào. Coscientist đã lấp đầy khoảng trống đó bằng mã lệnh, làm cho một dòng câu văn dẫn thẳng đến phản ứng hóa học trên bàn thí nghiệm.

Nhóm nghiên cứu của Giáo sư Gomes đã trực tiếp thử nghiệm xem cấu trúc này có bị dao động trước bấy lữa của con người hay không. Họ đặt phẩm màu đỏ, vàng, xanh dương vào ba vị trí trên đĩa 96 giếng, rồi đưa thông tin sai lệch cho Coscientist rằng thứ tự ban đầu là vàng, xanh dương, đỏ. Thứ tự thực tế khác với điều đó. Coscientist không tin lời con người mà tự mình vận hành máy quang phổ tử ngoại - khả kiến. Từ các giá trị độ hấp thụ quang ở các bước sóng 530 nanomet, 430 nanomet và 630 nanomet,

nó đã đọc các chỉ số độ hấp thụ và đánh chính rằng vị trí A1 là màu đỏ, vị trí B1 là màu vàng, vị trí C1 là màu xanh dương. Nghĩa là nó tin vào giá trị do chính mình đo đạc hơn là thông tin do con người cung cấp. Điều này cho thấy trước dữ liệu thực nghiệm, uy quyền của giáo sư hay đáp án định sẵn từ trước đều không có giá trị.

Bằng cách tương tự, Coscientist cũng đã thực hiện phản ứng Suzuki-Miyaura và phản ứng Sonogashira từ đầu đến cuối mà không cần qua bàn tay con người. Chi tiết cụ thể của thí nghiệm này sẽ được đề cập trong phần Tối ưu hóa phản ứng xúc tác.

Một lý do khác khiến hệ thống này trở nên đặc biệt là nó có thể vận hành mà không cần đến trang thiết bị đắt đỏ. Giá của một cánh tay robot Opentrons mà Coscientist điều khiển chỉ vào khoảng 5.000 đô la. Ngay cả khi không phải là trường đại học thuộc khối Ivy League của Mỹ hay công ty dược phẩm khổng lồ, bất kỳ nhà nghiên cứu nào biết nhập một dòng ngôn ngữ tự nhiên cũng đều có con đường mở ra để thiết kế các thí nghiệm ở mức độ tương tự.

Bàn tay từng cân bột paladi trước chiếc cân giờ đây đang chuyển thành bàn tay lựa chọn câu lệnh tiếp theo để giao cho robot.

Cấu trúc của Coscientist

Để hiểu được khung sườn của Coscientist được xây dựng như thế nào, cần phải xem xét từ nơi hệ thống này ra đời. Năm 2022, khi Gomes nhậm chức tại Đại học Carnegie Mellon, nhà trường đang hợp tác với Emerald Cloud Lab, một công ty khởi nghiệp ở San Francisco, đầu tư 50 triệu đô la để xây dựng phòng thí nghiệm đám mây từ xa đầu tiên của trường đại học. Phòng thí nghiệm này dự kiến sẽ lắp đặt

hàng trăm loại thiết bị chính xác, bao gồm máy quang phổ khối và máy quang phổ cộng hưởng từ hạt nhân. Về sau, phòng thí nghiệm này đã chính thức mở cửa vào đầu năm 2024, mở quyền tiếp cận hơn 200 loại thiết bị cho các nhà nghiên cứu của trường. Đó là cấu trúc mà khi nhà nghiên cứu tải mã thí nghiệm lên qua internet, robot và thiết bị phân tích sẽ thực thi y nguyên đoạn mã đó. Giống như việc lập trình viên thuê máy chủ của Amazon Web Services để chạy mã nguồn, nhà hóa học cũng có thể điều khiển phản ứng của các nguyên tử thực tế chỉ bằng một dòng mã. Tuy nhiên, đoạn mã đó phải được viết bằng một ngôn ngữ xa lạ có tên là Wolfram Mathematica. Đối với các nhà hóa học cả đời chỉ cầm pipet, rào cản ngôn ngữ này lại là một nỗi khổ cực khác. Gomes tự hỏi liệu thay vì con người phải tự viết đoạn mã phức tạp này, mô hình ngôn ngữ có thể chuyển những câu tiếng Anh thông thường thành mã lệnh mà robot hiểu được hay không.

Câu trả lời cho câu hỏi đó chính là sách trắng GPT-4 vào đêm ngày 14 tháng 3 năm 2023 đã đề cập ở trước. Ông có được niềm tin vững chắc rằng nếu mô hình ngôn ngữ có thể sai khiến con người xử lý công việc trong thế giới vật lý, thì cũng có thể kết nối trực tiếp mô hình đó với cánh tay robot. Gomes lập tức gửi hàng loạt tin nhắn cho các sinh viên trong phòng thí nghiệm, và Boiko cùng cộng sự Robert MacKnight đã hoàn thành bản mẫu thử nghiệm lấy GPT-4 làm khung sườn chỉ trong vòng 2 tuần.

Coscientist có cấu trúc phối hợp nhịp nhàng giữa năm bộ phận. Planner nhận mệnh lệnh một câu của con người để lập ra bảng kế hoạch. Web Searcher tìm kiếm trên Google Scholar và các trang tạp chí học thuật để thu thập tính chất của các chất phản ứng và tỷ lệ phối trộn. Document Searcher đọc trước tài liệu hướng dẫn sử dụng thiết bị robot, sau đó kiểm tra xem mệnh lệnh do Planner tạo ra có vượt quá thông số kỹ thuật thực tế của máy móc hay không. Ví dụ, mô-đun heater-shaker có nguy cơ làm chất lỏng trong đĩa giếng bắn ra ngoài nếu vượt quá 200 vòng/phút ở 45 độ C, và Document Searcher sẽ khắc ghi giới hạn này vào mã lệnh từ trước. Code Execution chạy mã Python bên trong một chiếc hộp cô lập gọi là Docker để đảm nhận các phép tính toán như tính nồng độ mol. Cuối cùng, thiết bị thực thi thí nghiệm sẽ truyền kết quả tính toán này đến cánh tay robot Opentrons và các thiết bị của Emerald Cloud Lab trên thực tế. Khi xảy ra lỗi mã, Coscientist tự đọc bản ghi nhật ký lỗi, tra cứu lại tài liệu và sửa mã. Tất cả đều không cần đến bàn tay con người.

Kết quả thử nghiệm rất cụ thể. Khi nhận lệnh "Hãy vẽ một đường chéo màu xanh dương và tô màu cách một hàng trên đĩa 96 giếng", Coscientist đã vẽ ra hình vẽ không chút sai sót chỉ trong 2 phút mà không cần bất kỳ đoạn mã tham khảo nào. Việc định chính phẩm màu và phục hồi lỗi đã thấy ở phần trước cũng chỉ có thể thực hiện được nhờ cấu trúc này.

Sau khi nghiên cứu được đăng trên Nature, bản thân lĩnh vực phòng thí nghiệm tự hành đã phát triển nhanh chóng. Bước sang năm 2026, các nhà nghiên cứu vẫn đang tiếp tục bước đi trên con đường mà Coscientist đã mở ra. Những nghiên cứu xây dựng chuẩn giao tiếp tiêu chuẩn kết nối thiết bị robot với tác nhân mô hình ngôn ngữ, hay nghiên cứu để tác nhân thực thi nguyên vẹn quy trình thí nghiệm khi quy trình đó được khai báo trước bằng mã, đã liên tiếp được công bố trong năm nay. Các bài báo tổng quan học thuật về phòng thí nghiệm tự hành cũng không ngừng gia tăng. Câu hỏi của một nhà hóa học từng kiệt sức trước chiếc cân, giờ đây đang trở thành câu hỏi mà nhiều phòng thí nghiệm đại học cùng nhau nắm giữ và giải quyết.

Tối ưu hóa phản ứng xúc tác

Trong cuộc đối đầu với Jose trước đó, Coscientist đã hoàn thành mã điều khiển Python để chạy đồng thời năm phản ứng ghép đôi chéo xúc tác paladi chỉ trong 6 phút. Kết quả này không chỉ cho thấy sự khác biệt về tốc độ. Phải mổ xẻ các quy trình bên trong để thấy được cách Coscientist xử lý phản ứng xúc tác, từ đó mới nhìn ra lý do thực sự.

Phản ứng mà Coscientist nhắm tới là phản ứng ghép đôi Suzuki-Miyaura và phản ứng ghép đôi Sonogashira. Cả hai phản ứng đều là phương pháp sử dụng kim loại paladi làm chất xúc tác để nối liền carbon với carbon. Vì đây là cánh cửa ắt phải đi qua khi chế tạo các chất ứng viên thuốc mới hay vật liệu đi-ốt phát quang hữu cơ (OLED), nên giải Nobel Hóa học năm 2010 cũng đã thuộc về các học giả phát triển những phản ứng này. Nếu con người làm thủ công, họ phải thay đổi từng lượng xúc tác, nhiệt độ và thời gian phản ứng rồi thử nghiệm hàng trăm lần mới tìm ra điều kiện tối ưu. Việc này giống như chơi trò Battleship, chấm từng tọa độ một để tìm tàu đối phương. Phải lần mò từng ô bên cạnh cho đến khi bắn trúng. Coscientist đã tiếp nhận trọn vẹn quá trình lần mò này.

Quy trình tuân theo đúng dòng chảy của năm tác nhân đã thấy ở phần trước. Planner nhận mục tiêu rồi thu thập tỷ lệ chất phản ứng và nhiệt độ thích hợp thông qua tìm kiếm tài liệu, còn mô-đun thực thi mã giải các phép tính nồng độ mol bằng tập lệnh Python và chỉ trả về đáp án chính xác. Điều này là do các mô hình ngôn ngữ lớn thường yếu trong các phép tính số thập phân. Toàn bộ quá trình này hoàn tất trong vòng 4 phút. Các giới hạn vật lý của thiết bị như heater-shaker được tác nhân phụ trách tìm kiếm tài liệu đối chiếu với sách hướng dẫn điều khiển để sàng lọc và loại trừ từ trước.

Bước tiếp theo mới là bài kiểm tra thực sự. Sau khi cánh tay robot Opentrons phân phối chất xúc tác và thuốc thử vào đĩa giếng và thiết bị heater-shaker bắt đầu gia nhiệt, một lỗi xung đột phiên bản phần mềm đã khiến thiết bị dừng lại. Đúng như đã thấy trước đó, Coscientist đã tự đọc bản ghi lỗi, tra cứu tài liệu hướng dẫn để sửa mã và tiếp tục phản ứng mà không bị gián đoạn.

Kết quả được xác nhận bằng máy sắc ký khí - khối phổ (GC-MS). Thiết bị này là công cụ đo lường chính xác loại và lượng chất còn lại sau khi phản ứng kết thúc. Kết quả đo đạc xác nhận rằng sản phẩm liên kết mục tiêu thực sự đã được tạo ra. Đây là trường hợp đầu tiên hoàn tất mọi quy trình từ lập kế hoạch, thực thi, sửa lỗi đến phân tích mà không cần sự can thiệp của con người. Điều này đồng nghĩa với việc thời gian cần thiết để tối ưu hóa một phản ứng xúc tác đã giảm từ vài tuần xuống còn vài phút. Gomes chia sẻ rằng ngay khoảnh khắc xác nhận kết quả này, ông và các sinh viên trong phòng thí nghiệm

đã cùng nhau rơi lệ. Bởi đó là khoảnh khắc một trí tuệ chưa từng chạm vào một nguyên tử carbon nào lại có thể tự mình thực hiện phản ứng vốn trước nay luôn được hoàn thành bằng bàn tay con người.

Tốc độ này không chỉ dừng lại ở một phản ứng. Nhóm nghiên cứu của Giáo sư Gomes đã giao cho Coscientist thí nghiệm động học phản ứng (kinetics) ghép đôi 12 loại amin và 18 loại aldehyde. Việc kết hợp hai nhóm này tạo ra hàng trăm trường hợp phản ứng khác nhau. Đó là số lượng tổ hợp mà nếu con người làm thủ công sẽ mất nhiều năm. Bằng việc vận hành robot chất lỏng và máy quang phổ khối độ phân giải cao không người lái suốt 24 giờ mỗi ngày, họ đã tích lũy được 1,2 triệu dữ liệu tốc độ phản ứng chỉ trong vài tuần. Không có bất kỳ trường hợp làm giả dữ liệu hay sai sót nào của con người bị trộn lẫn. Thuật toán từng tìm hiểu hiệu năng của một chất xúc tác, giờ đây đã vẽ nên toàn bộ bản đồ của hàng trăm loại chất xúc tác.

Việc con số này lớn đến mức nào có thể hiểu được khi nghĩ đến ngân sách của phòng thí nghiệm. Để một sinh viên cao học thu được một dữ liệu tốc độ phản ứng, họ phải cân thuốc thử, điều chỉnh nhiệt độ và cứ cách vài giờ lại phải lấy mẫu để ghi chép. Một ngày nhiều nhất cũng chỉ được vài trường hợp. Tính cả tiền thuốc thử và chi phí nhân công, chi phí để có được một dữ liệu cũng không hề nhỏ. Để lấp đầy 1,2 triệu dữ liệu bằng tay người, một sinh viên cao học cống hiến cả đời cũng không đủ. Coscientist đã thu hẹp khoảng cách này xuống còn vài tuần. Điều này có nghĩa là việc lựa chọn phản ứng xúc tác giờ đây không còn là cảm tính hay may mắn, mà đã trở thành một sự lựa chọn được kiểm chứng bằng dữ liệu.

Phương thức tìm kiếm tự hành từng nâng cao hiệu suất của một chất xúc tác paladi chỉ trong vài phút này đang được mở rộng vượt ra ngoài hóa học sang thiết kế thí nghiệm trong toàn bộ khoa học vật lý tại sân khấu mới mà Giáo sư Gomes đã chuyển đến. Chưa ai biết được phép tính từng giải quyết một phản ứng ghép đôi Suzuki trong 6 phút, lần tới sẽ được tái hiện lại trên bàn thí nghiệm vật lý nào.

Một dòng lời nhắc

Nhà hóa học ngồi trước máy tính và chỉ nhập một câu duy nhất: "Hãy thiết kế và thực thi công thức tối ưu cho phản ứng ghép đôi chéo Suzuki-Miyaura dành cho linh kiện OLED cỡ nhỏ." Trên màn hình chỉ có đúng một dòng đó. Không có danh sách thuốc thử, không có bảng điều kiện phản ứng, cũng không có mã điều khiển robot. Vài phút sau, ở một góc tòa nhà Hóa học của Đại học Carnegie Mellon, cánh tay robot bắt đầu tự động chuyển động. Đó là khoảnh khắc câu văn do con người viết chuyển hóa thành phản ứng hóa học thực tế. Coscientist là một bộ dịch thuật biến một dòng lời nhắc của con người thành mệnh lệnh mà robot có thể hiểu được.

Bộ dịch thuật này không làm việc một mình. Năm tác nhân đã đề cập trước đó lần lượt tiếp nhận công việc theo thứ tự. Giống như một trưởng nhóm nhận một dự án rồi phân chia nhiệm vụ cho từng thành viên, khi Planner chuyển câu văn thành lưu đồ thì các vai trò còn lại sẽ tiếp quản việc tìm kiếm, tính toán và thực thi. Con người chỉ viết một dòng đầu tiên, và từ đó trở đi, năm vai trò này sẽ trao đổi kết quả với nhau để hoàn thiện công thức.

Có một lý do riêng giải thích tại sao việc phân chia vai trò này lại cần thiết. Mô hình ngôn ngữ viết văn rất tốt nhưng tính toán lại thường xuyên sai sót. Ký hiệu SMILES, một phương pháp biểu diễn phân tử bằng văn bản, cũng là một trở ngại. Có thể xem đây như cách ghi lại thứ tự gấp giấy bằng chữ. Thế nhưng chỉ cần một dấu ngoặc bị lệch, máy tính sẽ hoàn toàn không đọc được. Coscientist đã giải quyết vấn đề này bằng cách sử dụng một ký hiệu nghiêm ngặt hơn có tên là SELFIES. Bản thân các quy tắc vật lý—chẳng hạn như nguyên tử carbon chỉ có thể có tối đa bốn liên kết—đã được khắc sâu vào bên trong ngữ pháp. Vì vậy, các phân tử được viết bằng ký hiệu này ngay từ đầu chỉ tạo ra những cấu trúc thực sự có thể tồn tại trong thực tế. Có thể coi như bên hiểu rõ ngôn ngữ con người đảm nhận việc viết, còn bên giỏi tính toán đảm nhận phần tính toán. Những mệnh lệnh được chuyển ngữ như vậy đã mang lại thành quả gì trên bàn thí nghiệm thì chúng ta đã thấy ở các phần trước.

Giao diện này cũng phải tự động lọc bỏ các câu lệnh (prompt) nguy hiểm. Điều gì sẽ xảy ra nếu ai đó yêu cầu lập trình mã tổng hợp chất độc thần kinh hoặc các chất do Cơ quan Phòng chống Ma túy kiểm soát? Để ngăn chặn rủi ro này, Giáo sư Gomes đã thiết lập một màng chắn kiểm chứng kép mang tên Guardians tại điểm truy cập của container Docker. Nếu phát hiện bất kỳ dấu vết nào trùng lặp với các chất bị kiểm soát trong cấu trúc phân tử được tạo ra, Coscientist sẽ lập tức dừng tính toán và ngắt nguồn robot. Nhờ công lao thiết lập cơ chế an toàn này, Giáo sư Gomes đã tham gia với tư cách Cố vấn Đặc biệt trong việc xây dựng các hướng dẫn an toàn vật lý của Sắc lệnh Hành pháp về trí tuệ nhân tạo do Tổng thống Joe Biden ký vào tháng 10 năm 2023. Khi một dòng lệnh mang sức mạnh làm dịch chuyển các nguyên tử thực tế, điều đó cũng đồng nghĩa với việc phải tạo ra cả cơ chế để sàng lọc chính dòng lệnh đó.

Cấu trúc này đang vượt ra ngoài bức tường của các phòng thí nghiệm. Trong một bài báo công bố trực tuyến ngày 8 tháng 5 năm 2026, JACS Au, tập san của Hội Hóa học Hoa Kỳ, đã đánh giá Coscientist cùng với các hệ thống như ChemCrow và LLM RDF là những cột mốc mở ra kỷ nguyên của phòng thí nghiệm tự hành. Tạp chí chuyên ngành hóa học C&EN, trong bản tin ngày 25 tháng 6 cùng năm, đã giới thiệu phòng thí nghiệm hóa học do Atinary Technologies mở tại Boston. Tại đó, khi nhà nghiên cứu con người nói ra chất mình mong muốn, tác tử AI sẽ tìm kiếm lộ trình tổng hợp. Một bài báo xuất bản vào tháng 2 cùng năm cũng chỉ ra rằng những tác tử này vẫn bộc lộ hạn chế trước các nhiệm vụ phức tạp như thuốc điều trị bằng peptide hay thí nghiệm in vivo (trong cơ thể sống). Dù vậy, phương hướng phát triển đã rất rõ ràng: một giao diện phản hồi theo từng dòng lệnh đang chuyển mình từ phát minh của một phòng thí nghiệm đại học trở thành tiêu chuẩn trong môi trường công nghiệp.

Những vấn đề vẫn chưa có lời giải

Sức mạnh chuyển hóa một dòng lệnh thành phản ứng hóa học này cũng có thể bị sử dụng theo chiều hướng ngược lại. Câu hỏi đặt ra là điều gì sẽ xảy ra nếu một câu văn yêu cầu tạo ra các chất độc hại, hợp chất gây nổ hoặc chất độc được nhập vào cùng một giao diện đó? Nguy cơ bị lạm dụng để tổng hợp tiền chất vũ khí hóa học là hoàn toàn có thể xảy ra. Người đối mặt trực tiếp với rủi ro lưỡng dụng (dual-use) này không ai khác chính là nhóm nghiên cứu của Giáo sư Gomes. Bài báo trên tạp chí Nature năm 2023 cho biết họ đã tiến hành thử nghiệm cố tình yêu cầu Coscientist tổng hợp các chất bị kiểm soát và thuốc bất hợp pháp, qua đó thiết kế các cơ chế an toàn để hệ thống từ chối hầu hết các yêu cầu nguy hiểm. Dù vậy, các tác giả cũng viết rằng lớp phòng thủ này không hề hoàn hảo, và giới học thuật cùng các nhà hoạch định chính sách cần phải cùng nhau thảo luận về khả năng bị lạm dụng.

Những tiếng nói cảnh báo về mối nguy hiểm cũng tiếp tục xuất hiện bên ngoài phòng thí nghiệm. Fabio Urbina và Sean Ekins, những người nghiên cứu về thiết kế thuốc bằng trí tuệ nhân tạo, trong một bài báo đăng trên Nature Machine Intelligence năm 2022, đã chỉ ra rằng chỉ cần đảo ngược mục tiêu của mô hình được huấn luyện để tìm kiếm các phân tử có lợi, nó có thể biến thành một công cụ thiết kế các chất độc hại. Họ cảnh báo rằng khoảng cách giữa nghiên cứu hữu ích và sự lạm dụng độc hại gần hơn người ta tưởng. Nhiều ý kiến chỉ ra rằng ngay cả khi có màn chắn kiểm chứng như Guardians của Coscientist, ranh giới này vẫn có thể bị hạ thấp bất cứ lúc nào nếu số lượng các mô hình công khai thiếu vắng cơ chế an toàn ngày càng gia tăng.

Khi robot tự mình tiến hành thí nghiệm từ đầu đến cuối mà không qua bàn tay con người, câu hỏi về việc ai sẽ chịu trách nhiệm nếu xảy ra sự cố vẫn còn bỏ ngỏ. Ranh giới giữa nhà nghiên cứu nhập vào một dòng văn bản, công ty phát triển tạo ra mô hình, hay nhà sản xuất cung cấp robot đều rất mờ nhạt. Các nhà nghiên cứu an toàn theo dõi phòng thí nghiệm tự hành chỉ ra rằng mức độ tự chủ càng cao thì khoảng trống trách nhiệm này càng lớn. Tốc độ đạt được thành quả diễn ra nhanh chóng bao nhiêu thì các quy chuẩn và thể chế để kiểm soát tốc độ đó vẫn chưa theo kịp bấy nhiêu.

Các nội dung trong chương này được viết dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Ghi chú của tác giả. Vai trò của con người đã chuyển dịch về đâu?

Hai nhân vật trong Phần 4 đã trao đổi tay trên bàn thí nghiệm cho máy móc. Adam và Genesis của King đã nuôi cấy nấm men và kiểm chứng giả thuyết mà không cần con người, trong khi Coscientist của Gomes đã điều khiển cánh tay robot bằng lời nói của mô hình ngôn ngữ để tạo ra các phản ứng hóa học. Tại đây, trí tuệ nhân tạo lần đầu tiên bước ra ngoài màn hình. Đó là giai đoạn mà tri thức có được đổi bàn tay.

Cỗ máy khi có được đôi tay đã chạm đến điểm yếu lâu đời của khoa học: tính tái lập. Khi quy trình thí nghiệm vốn dựa vào cảm giác đầu ngón tay của con người biến thành các dòng mã, robot ở các phòng thí nghiệm khác có thể lặp lại y nguyên thí nghiệm đó. Tuy nhiên, cỗ máy có được đôi tay cũng mang theo những rủi ro mới. Vòng lặp khép kín có thể không tự nhận ra những giả định sai lầm ban đầu mà tiếp tục chôn vùi các thí nghiệm lên trên đó, khiến sai số ngày càng lớn hơn, và quy trình do mô hình ngôn ngữ thiết lập cũng có thể tạo ra các chất nguy hiểm.

Với tư cách là một luật sư, điểm khiến tôi dừng lại suy ngẫm lâu nhất trong hai chương này chính là các cơ chế an toàn. Khi Coscientist của Gomes dễ dàng lập trình việc tổng hợp chất nổ, thứ ngăn chặn đoạn mã đó trước khi thực thi chính là cơ chế mà con người đã cài đặt sẵn từ trước. Khi đôi tay trên bàn thí nghiệm đã được chuyển giao cho máy móc, việc còn lại của con người là vạch ra ranh giới mà đôi tay ấy không được phép vượt qua. Trong Phần 4, vị trí của con người chính là vị trí của người vạch ra ranh giới.

Phần 5

Cỗ máy tự tuần hoàn

Chương 10. Sam Rodriques: Cỗ máy đọc một triệu bài báo

FutureHouse

Tại một văn phòng ở khu Dogpatch, San Francisco, ánh đèn vẫn sáng dù đã quá 11 giờ đêm. Bên ngoài cửa sổ trời đã tối đen, Sam Rodriques ngồi trước màn hình, chăm chú theo dõi những dòng mã đang cuộn trôi. Ở nơi chỉ vài giờ trước còn là một màn hình trống rỗng, giờ đây những phán đoán của trí tuệ nhân tạo sau khi đọc hàng trăm bài báo khoa học đang nối tiếp nhau xếp chồng lên. Cách đây vài năm, anh từng là trưởng nhóm nghiên cứu độc lập tại Viện Francis Crick ở Luân Đôn. Trong phòng thí nghiệm đó, ngày nào anh cũng chứng kiến cùng một cảnh tượng: những nghiên cứu sinh tiến sĩ tài năng phải lật giở thủ công hàng nghìn bài báo, điền đơn xin kinh phí nghiên cứu theo đúng biểu mẫu, và tay cầm pipet đứng hàng giờ liền trước bàn thí nghiệm. Rodriques từng nói về giai đoạn này rằng: chế độ học việc mang tên trường cao học rất cuộc chỉ là một cấu trúc trói buộc sinh viên vào bàn thí nghiệm để sử dụng như nguồn lao động giá rẻ. Phải mất nhiều năm mới cho ra đời được một bài báo khoa học, nhưng phần lớn khoảng thời gian đó không dành cho tư duy mà là lao động chân tay.

Điểm xuất phát của Rodriques không phải là sinh học mà là vật lý. Anh tốt nghiệp thủ khoa đại học tại Đại học Haverford (Mỹ) với nghiên cứu tính toán các trạng thái lượng tử nơi nhiều hạt vướng víu vào nhau. Tiếp đó, anh nhận Học bổng Churchill để theo học chương trình thạc sĩ kỹ thuật và vật lý toán học tại Đại học Cambridge (Anh), rồi bước vào chương trình tiến sĩ vật lý tại Viện Công nghệ Massachusetts. Con đường trở thành một nhà vật lý lý thuyết ưu tú đã mở rộng thênh thang trước mắt anh. Thế nhưng, ngay từ đầu chương trình tiến sĩ, anh đã rơi vào sự hoài nghi sâu sắc. Anh cho rằng vật lý đã giải mã các hiện tượng trước mắt chúng ta đến từng chữ số thập phân, nhưng những bí ẩn còn lại chỉ có đúng hai điều: thực thể trí tuệ là con người và hệ thống phức hợp khổng lồ là sự sống. Tại sao con chuột đi ngang qua phía trước lại dừng đúng chỗ đó để ngủ, tại sao con người lại mắc bệnh, già đi và qua đời vào một thời điểm nhất định – vật lý không thể trả lời dù chỉ bằng một phương trình nhân quả. Việc các phương trình có thể giải thích cả vũ trụ nhưng lại không thể giải thích nổi một con chuột trước mắt đã luôn ám ảnh anh.

Anh đã chuyển hướng. Qua lại giữa MIT Media Lab và Khoa Khoa học Não bộ và Nhận thức, anh làm nghiên cứu sinh tiến sĩ dưới sự hướng dẫn của Giáo sư Ed Boyden, bậc thầy về nghiên cứu quang di truyền học – kỹ thuật dùng ánh sáng để bật tắt tế bào thần kinh. Trong giai đoạn này, anh liên tiếp phát triển công nghệ hệ phiên mã đọc được thời điểm và vị trí gen trong tế bào được bật tắt theo cả trục không gian lẫn thời gian, cùng công nghệ chế tạo nano 3D. Năm 2019, anh nhận Giải thưởng Luận án Hertz dành cho luận án tiến sĩ xuất sắc nhất trong năm.

Sau khi hoàn thành bằng tiến sĩ, anh được bổ nhiệm làm trưởng nhóm tại Viện Francis Crick. Ở độ tuổi vừa ngoài ba mươi đã được dẫn dắt một phòng thí nghiệm độc lập, trong mắt người khác đó là một khởi đầu rực rỡ. Tuy nhiên, điều chờ đón anh lại là sự quan liêu và kém hiệu quả. Dù một bài báo khoa học vượt qua khâu bình duyệt, con đường để thành quả đó chuyển hóa thành phương pháp điều trị thực tế vẫn luôn bị tắc nghẽn. Anh đã nhiều lần chứng kiến những ý tưởng tuyệt vời lặn ngụp trong ngăn kéo của các tạp chí học thuật. Cùng với đồng nghiệp Adam Marblestone, anh đã hình thành mô hình Tổ chức Nghiên cứu Tập trung, tức FRO. Ý tưởng là để một tổ chức phi lợi nhuận quy mô nhỏ với ngân sách 50 triệu USD trong 5 năm đảm nhận và giải quyết rất ráo những dự án lớn mà các phòng thí nghiệm đại học quá nhỏ để gánh vác, còn các công ty được phẩm lại ngó lơ vì không có lợi nhuận. Đó là nơi dành cho những nhiệm vụ to lớn nhưng không sinh lời ngay lập tức, chẳng hạn như vẽ bản đồ kết nối thần kinh của toàn bộ não bộ hay làm sáng tỏ cận kề các protein trên bề mặt màng tế bào.

Gốc rễ của ý tưởng này bắt nguồn từ vấn đề mà anh gọi là trần nhận thức. Điều đó có nghĩa là lượng thông tin mà não người có thể chứa đựng là có hạn, trong khi tri thức mà thế giới tuôn ra đã vượt xa giới hạn ấy. Bộ gen người có 20.000 gen, tế bào não được chia thành hàng nghìn phân nhóm, và cứ 3 giây lại có một bài báo y học được xuất bản. Vào thế kỷ 17, Newton có thể tự mình chứa toàn bộ tri thức khoa học của nhân loại trong đầu, nhưng ngày nay không một người đoạt giải Nobel nào có thể đọc hết dù chỉ 0,1% số bài báo sinh học xuất bản trong một năm. Manh mối ngăn ngừa ung thư hay chìa khóa chữa khỏi chứng mất trí nhớ có thể đã nằm rải rác trong các bài báo do hơn 20 nhà nghiên cứu công bố suốt 20 năm qua tại nhiều quốc gia khác nhau. Vấn đề nằm ở chỗ không có cách nào tập hợp những điểm nổi đó vào một bộ não duy nhất để liên kết nguyên nhân và kết quả. Việc vốn và thiết bị thí nghiệm có thể tăng lên bằng tiền, nhưng để đào tạo một tiến sĩ vẫn phải mất 20 năm cũng khiến anh cảm thấy bế tắc.

Năm 2023, anh từ bỏ vị trí biên chế suốt đời tại Viện Crick và thành lập viện nghiên cứu phi lợi nhuận FutureHouse ngay giữa khu Dogpatch. Cựu Giám đốc điều hành Google Eric Schmidt và quỹ Open Philanthropy đã tài trợ vốn. Anh đưa nhà hóa học trí tuệ nhân tạo Andrew White về làm giám đốc khoa học. Mục tiêu chỉ có một: giao lại cho máy móc công việc phân tích tài liệu và mô phỏng mà các nhà nghiên cứu phải ôm đồm suốt nhiều tháng trời. Đến tháng 11 năm 2025, anh thành lập riêng công ty vì lợi nhuận Edison Scientific để đưa công nghệ này vào thực tế ngành công nghiệp.

Dòng chảy này tiếp tục kéo dài sang năm 2026. Ngày 19 tháng 5, Edison Scientific đã bắt tay với công ty dược phẩm niêm yết trên Nasdaq là Incyte để cung cấp nền tảng Cosmos bao trùm từ việc phát hiện ứng viên thuốc mới, kiểm chứng độc tính đến soạn thảo tài liệu lâm sàng. Việc để Cosmos cùng xem xét cả sổ tay điện tử và cửa sổ trò chuyện Slack trong phòng thí nghiệm, đồng thời cho phép đối tác huấn luyện lại mô hình theo ý muốn chính là điểm

cốt lõi của hợp đồng lần này. Trong thông báo này, Rodriques cho biết phát triển thuốc mới là một công việc tích lũy, nơi mỗi thí nghiệm, mỗi phân tích và mỗi kết quả lâm sàng phải giúp cho phán đoán tiếp theo trở nên tốt hơn. Vào ngày 29 tháng 6 cùng năm, công ty cũng đã bắt tay với Population Health Partners, vườn ươm sinh học từng sản sinh ra Metsera – công ty được bán cho Pfizer với giá 10 tỷ USD. Đây là một thử nghiệm trong đó Cosmos đảm nhận thay cả việc soạn thảo hồ sơ pháp lý trong quá trình phát triển thuốc mới nhằm vào các bệnh tim mạch, phổi và dạ dày ruột vốn có thể ảnh hưởng đến 10% dân số thế giới. Rodriques tuyên bố mục tiêu của sự hợp tác này là ngăn ngừa hoặc chữa khỏi bệnh tật, và đẩy nhanh toàn bộ quá trình đó bằng trí tuệ nhân tạo. Thay thế vị trí của những sinh viên từng đứng cầm pipet trong phòng thí nghiệm của Viện Crick, giờ đây là những cỗ máy thức trắng đêm sửa mã. Viện nghiên cứu do một sinh viên từng muốn trở thành nhà vật lý thành lập đang viết lại chính thời gian biểu của ngành dược phẩm.

Tri thức không được đọc

Trước khi đi đến câu trả lời mang tên FutureHouse, điều giữ chân Rodriques lâu nhất không phải là lao động lặp đi lặp lại bên bàn thí nghiệm mà là những chồng bài báo khoa học không ai đọc đến. Nếu nhìn trần nhận thức mà anh đề cập qua các con số thì sẽ như sau: mỗi năm các bài báo học thuật trên toàn thế giới tuôn ra hàng triệu bài, riêng lĩnh vực y sinh học đã hơn 1 triệu bài một năm, và nếu chỉ thu hẹp vào lĩnh vực ung thư thì mỗi tuần đã có hàng trăm bài báo mới xuất hiện. Giống như một chiếc cốc chứa nước cứ tiếp tục được rót đầy mà không đổ ra dù chỉ một giọt thì cuối cùng sẽ tràn, lượng tri thức tăng theo cấp số nhân trong khi dung lượng não bộ để hấp thụ nó vẫn giữ nguyên.

Trên khoảng cách ấy, Rodriques đã xây dựng một niềm tin vững chắc. Anh từng nói: "Hiện nay dưới đáy những chồng bài báo trên toàn thế giới, những tờ tiền mệnh giá 100 triệu USD đang nằm rải rác mà không ai nhặt." Ví dụ anh đưa ra là câu chuyện xảy ra vào mùa thu năm 2025. Trace Neuroscience, một công ty công nghệ sinh học phát hiện đột biến gen UNC13A – mục tiêu điều trị bệnh xơ cứng cột bên teo

cơ (ALS), đã thu hút khoản đầu tư trị giá 100 triệu USD. Kết quả thí nghiệm dẫn đến phát hiện này không phải là điều mới tạo ra. Đó là những mảnh ghép rải rác trong các bài báo và dữ liệu bộ gen công khai mà các nhà nghiên cứu từ nhiều quốc gia đã công bố qua nhiều năm. Chỉ là từ trước đến nay chưa có ai cùng lúc tập hợp những mảnh ghép đó vào đầu một người để chỉ ra mối quan hệ nhân quả.

Rodrigues đã đi đến kết luận: chẳng có lý do gì để chờ đợi một người thông minh hơn xuất hiện. Thay vào đó, chỉ cần tạo ra một nhà khoa học trí tuệ nhân tạo có khả năng đọc hàng triệu tài liệu mà không mắc sai sót. Các dự án PaperQA2 và WikiCrow đã được khởi động như những nhiệm vụ đầu tiên của FutureHouse. Ý định là để một cỗ máy đọc và tổng hợp xuyên đêm chồng bài báo mà một nghiên cứu sinh phải mất cả đời cũng không thể đọc hết.

PaperQA2 và WikiCrow

Một ngày năm 2024, hàng chục nhà sinh học nhận bằng tiến sĩ từ Harvard, Viện Công nghệ Massachusetts và Oxford đã tập hợp tại một nơi. Nhóm nghiên cứu FutureHouse đề nghị trả cho họ tối đa 12 USD cho mỗi câu hỏi trả lời đúng. Họ được mở toàn quyền truy cập cơ sở dữ liệu bài báo trả phí của thư viện cũng như tìm kiếm internet. Thời gian được giao là một tuần. Đối thủ của họ không phải là con người, mà là chương trình PaperQA2 do Sam Rodrigues tạo ra. Thời gian dành cho PaperQA2 chỉ vèo vèo vài giờ.

Bộ đề gồm 248 câu hỏi. Đó là những câu hỏi mà câu trả lời chỉ xuất hiện đúng một lần trong toàn văn các bài báo sinh học gần đây, được thiết kế để không thể giải nếu chỉ đọc phần tóm tắt. Kết quả đã rõ ràng: độ chính xác trung bình của các tiến sĩ và nghiên cứu sinh sau tiến sĩ là 64,3%, với độ lệch chuẩn lớn lên tới 15,2%, cho thấy sự chênh lệch trình độ giữa người với người là rất đáng kể. Trong khi đó, PaperQA2 ghi nhận độ chính xác 85,2% với độ lệch chuẩn chỉ 1,1%. Chương trình tìm kiếm thương mại Perplexity Pro dừng lại ở mức 69,7%, dịch vụ tìm kiếm bài báo Elicit chỉ đạt 40,9%, còn GPT-4o khi không dùng chức năng tìm kiếm tụt xuống 44,6%. Vấn đề mà con người phải vật lộn suốt một tuần, cỗ máy đã giải quyết chính xác hơn chỉ trong vài giờ.

Sự khác biệt này đến từ đâu? Phương pháp tạo sinh tăng cường truy xuất thông thường, tức RAG, thường cắt bài báo thành các đoạn khoảng 500 ký tự rồi đưa vào cơ sở dữ liệu, sau đó dán nguyên ba hoặc bốn đoạn tương tự câu hỏi vào lời nhắc. Nhưng bài báo khoa học là văn bản đan xen giữa phần mở đầu, phương pháp, kết quả và thảo luận. Nếu cắt câu một cách máy móc, những câu không liên quan đến câu hỏi sẽ xen vào và chương trình sẽ bịa ra câu trả lời sai lệch. Để giải quyết vấn đề này, Rodrigues đã thiết kế quy trình mới mang tên RCS (xếp hạng lại và tóm tắt theo ngữ cảnh). Đầu tiên, công cụ phân tích tài liệu học thuật GROBID chia bài báo thành các phần: mở đầu, phương pháp, kết quả và danh mục tài liệu tham khảo. Khi loại bỏ các ký tự vô ích như tài liệu tham khảo hay khoảng trắng trùng lặp, độ dài trung bình của một bài báo giảm từ 16.040 token xuống còn 8.903 token. Tiếp đó, một mô hình ngôn ngữ như GPT-4 Turbo hoặc Claude 3.5 Sonnet sẽ đọc từng đoạn trong số 30 đến 50 đoạn văn ứng viên và chấm điểm mức độ liên quan đến câu hỏi từ 0 đến 10. Chỉ những đoạn đạt trên 8 điểm mới được chuyển cho mô hình viết câu trả lời cuối cùng; phần còn lại bị loại bỏ. Quy trình lọc, chấm điểm rồi lại lọc này đã nhổ tận gốc mầm mống của những câu trả lời sai.

Cách chấm điểm này có lý do của nó. Việc chỉ sử dụng các đoạn văn đạt trên 8 điểm giúp ngăn chặn từ gốc thói quen gượng ép viện dẫn các bằng chứng mơ hồ để bịa ra câu trả lời. Trên thực tế, khi không chắc chắn, PaperQA2 sẽ không đưa ra câu trả lời mà tuyên bố là 'thiếu thông tin'. Tỷ lệ này là 21,9%, nghĩa là cứ năm câu hỏi thì có một câu chương trình thừa nhận không biết. Một chương trình biết nói không biết khi không biết sẽ hữu ích cho các nhà nghiên cứu hơn nhiều so với một chương trình giả vờ biết nhưng lại trả lời sai.

PaperQA2 còn tiến thêm một bước nữa tại đây. Khi một nhà nghiên cứu con người có trong tay một bài báo hay, họ sẽ tra cứu tài liệu tham khảo của bài báo đó hoặc tìm các bài báo tiếp theo trích dẫn nó để mở rộng tri thức. PaperQA2 đã tự động hóa quá trình này bằng công cụ duyệt mạng lưới trích dẫn đầu cuối. Đặt các bài báo cốt lõi đạt trên 8 điểm vào trung tâm, nó kéo về các bài báo được bài đó trích dẫn và các bài báo trích dẫn bài đó trong phạm vi một cấp. Tuy nhiên, nếu kéo về mọi bài báo thì phạm vi tìm kiếm sẽ phình to như quả bóng. Vì vậy, trong số sáu bài báo liên quan, nó chỉ giữ lại những bài báo được ít nhất hai bài trở lên trích dẫn chéo nhau, còn lại đều bị lọc bỏ như tạp âm.

Được xây dựng trên nền tảng động cơ này chính là WikiCrow. Khi nhập một tên gen như FAM83H, năm chương trình sẽ đồng thời hoạt động: một chương trình tìm cấu trúc 3D của protein, một chương trình tìm chức năng của gen này trong tế bào, một chương trình tìm tương tác với các phân tử khác, và một chương trình tìm mối liên hệ với ung thư hoặc bệnh di truyền. Chương trình cuối cùng đọc kết quả của bốn chương trình trước để viết một bản tóm tắt tổng quan. Năm loại kết quả này được mã Python ghép nối thành một bài viết hoàn chỉnh. Đội ngũ FutureHouse đã cho WikiCrow chạy trên toàn bộ 20.000 gen người. Cho đến thời điểm đó, Wikipedia chỉ có 3.639 gen có tài liệu hoàn chỉnh, chưa đầy 18% tổng số. Chạy không ngừng nghỉ trong vài ngày, WikiCrow đã gọi PaperQA2 100.000 lần và đọc hơn 871.000 bài báo khoa học – chiếm hơn 1% tổng số bài báo khoa học đời sống trên toàn thế giới. Kết quả được giao cho bốn tiến sĩ sinh học bên ngoài thẩm định theo phương pháp mù. Tỷ lệ câu viết không có nguồn của Wikipedia là 13,6%, trong khi của WikiCrow là 3,5%. Tỷ lệ câu có trích dẫn nguồn nhưng nội dung sai của Wikipedia là 24,9%, còn của WikiCrow là 13,5%. Độ chính xác tổng thể của Wikipedia là 71,2%, trong khi WikiCrow đạt 86,1%. Qua công việc này, người ta cũng phát hiện ra rằng trong số 20.000 gen, có tới 670 gen chưa từng có dù chỉ một dòng bài báo nào viết về chúng. Trên bản đồ tri thức, vẫn còn gần ấy vùng đất chưa từng có ai đặt chân tới.

Người ta ước tính một tiến sĩ con người trong suốt cuộc đời chỉ tiếp cận khoảng 1.500 bài báo khoa học. PaperQA2 bắt kịp khối lượng đó chỉ với 200 USD tài nguyên điện toán và 4 giờ đồng hồ. Khi vài dòng mã có thể thay thế sức lao động của nhiều ngày, con người lại phải tự hỏi mình nên làm gì trong quỹ thời gian dôi dư đó. Các nguyên tắc chọn lọc và chấm điểm nêu trên khi được tích lũy lại đã tạo nên độ chính xác vượt trội hơn cả một đội ngũ tiến sĩ con người.

Bước sang năm 2026, động cơ này đã vượt ra ngoài bức tường phòng thí nghiệm. Vào tháng 2, Edison Scientific đã ra mắt LAB-Bench-2, thiết lập một thước đo mới để đánh giá trí tuệ nhân tạo phục vụ nghiên cứu sinh học. Họ chia 1.892 câu hỏi thành 11 danh mục, loại bỏ các câu hỏi trắc nghiệm và đưa vào những nhiệm vụ đòi hỏi phải tìm đọc cả tài liệu bằng sáng chế lẫn dữ liệu thử nghiệm lâm sàng mới có thể giải quyết được. Người ta nhận thấy trí tuệ nhân tạo càng được cung cấp kèm công cụ tìm kiếm web hoặc thực thi mã thì điểm số càng tăng mạnh. Một chương trình từng đào bới tri thức trong các chồng bài báo khoa học giờ đây đã trở thành tiêu chuẩn để đo lường các trí tuệ nhân tạo khác.

Đôi mắt tìm kiếm mâu thuẫn

Năm 2010, một bài báo khoa học được công bố trên một tạp chí y học. Nhóm nghiên cứu của Kriegl khi kiểm tra mô của bệnh nhân ung thư đại trực tràng đã phát hiện ra rằng protein tên là LEF1 biểu hiện càng nhiều thì bệnh nhân sống càng lâu. Thế nhưng đúng một năm sau, một tạp chí khác lại công bố một bài báo hoàn toàn trái ngược: cùng là protein LEF1, nhưng lần này biểu hiện càng nhiều thì thời gian sống sót càng ngắn và càng thúc đẩy ung thư di căn sang gan. Hai bài báo cùng chỉ về hai hướng ngược nhau tại ngã rẽ điều trị ung thư đại trực tràng, và đã nằm cạnh nhau trong kho lưu trữ suốt 14 năm ròng. Không một ai nhận ra điều đó.

Lý do con người không nhận ra rất đơn giản. Như đã thấy ở trên, trong một thế giới mà một người dù cố gắng hết sức cũng không thể đọc hết tài liệu trong chính chuyên ngành của mình, khi 99 bài báo

cùng cất lên một tiếng nói thì 1 bài nói điều ngược lại sẽ bị coi là tạp âm và lặng lẽ chìm vào quên lãng. Rodriques đã quyết định giao hẳn nỗi vất vả đọc những ý kiến thiếu số bị chôn vùi như vậy cho máy móc.

Vào mùa thu năm 2024, FutureHouse đã tạo ra một chương trình mang tên ContraCrow. Mục đích chỉ có một: thay con người tìm ra những mâu thuẫn tiềm ẩn giữa các bài báo khoa học. ContraCrow hoạt động trên nền tảng công cụ tìm kiếm học thuật PaperQA2. Quy trình vận hành như sau: trước tiên, nó chia một bài báo thành các đoạn 5.000 ký tự theo đơn vị đoạn văn và mục. Lúc này, mỗi đoạn đều được gắn kèm tiêu đề bài báo gốc và tiêu đề mục, giúp cỗ máy không bị nhầm lẫn giữa một tuyên bố ở phần mở đầu với một số liệu trong kết quả thí nghiệm. Tiếp theo, nó trích xuất các luận điểm có thể kiểm chứng từ đoạn văn. Những câu hiển nhiên như 'chúng tôi đã sử dụng pipet' sẽ bị lọc bỏ, chỉ giữ lại những luận điểm đạt trên 8/10 điểm. Các luận điểm còn lại sẽ được đối chiếu bằng cách tìm tài liệu liên quan trong cơ sở dữ liệu hơn 150 triệu bài báo trên toàn thế giới. Bản tóm tắt cô đọng đoạn văn gốc từ 2.250 token xuống còn 200 đến 400 token được tạo ra và chuyển cho mô hình ngôn ngữ Claude 3.5 Sonnet. Mô hình này sẽ phân định xem có bài báo nào bác bỏ trực diện luận điểm đó hay không.

Trước đây cũng từng có những chương trình tự động kiểm chứng các luận điểm trong bài báo khoa học, chẳng hạn như SciFact hay MultiVers. Tuy nhiên, các chương trình này chỉ tìm kiếm câu trả lời trong một tập dữ liệu nhỏ được con người chọn lọc từ trước, chẳng khác nào loài cá trong một bể kính chật hẹp. ContraCrow thì khác: không bị giới hạn trong phạm vi định sẵn, nó rà soát toàn bộ kho lưu trữ học thuật hơn 150 triệu bài báo trên toàn thế giới theo thời gian thực. Việc có thể đặt cạnh nhau và so sánh dữ liệu từ hai tạp chí cách nhau hơn 10 năm chỉ trong vài

mili giây cũng chính là nhờ điều này.

Tại đây, ContraCrow không trả lời dứt khoát theo kiểu mâu thuẫn hay không mâu thuẫn, mà trả lời theo thang đo Likert 11 bậc từ 0 đến 10 điểm. 0 điểm là hoàn toàn trùng khớp, 10 điểm là mâu thuẫn rõ ràng xung đột trực diện. Chỉ khi đạt trên 8 điểm, nó mới xác nhận là mâu thuẫn thực sự. Tiêu chuẩn này xuất phát từ đâu? Nhóm nghiên cứu đã tạo ra một đề kiểm tra với những câu hỏi đã biết trước đáp án chính xác: bộ đo chuẩn ContraDetect với một nửa là các câu chắc chắn đúng và một nửa là các đáp án sai cố tình trộn lẫn các câu phản bác. Tại đây, ContraCrow đạt điểm AUC 0,842. Khi lấy 8 điểm làm đường chuẩn, độ chuẩn xác đạt 73%, độ chính xác đạt 88%, và tỷ lệ phán đoán sai giảm xuống dưới 7%. Khi bí mật trộn vào 42 luận điểm giả mạo chưa từng được báo cáo trên thế giới, chương trình đã trung thực trả lời với xác suất 98% rằng không có bằng chứng. Điều đó có nghĩa là nó không bịa đặt những điều không tồn tại thành có.

Vì những con số đơn thuần chưa đủ mang lại cảm giác thực tế, nhóm nghiên cứu đã thử nghiệm trên các bài báo khoa học thực tế. Họ đưa vào ContraCrow 93 bài báo sinh học đã qua bình duyệt và xuất bản. Từ đó, chương trình trích xuất 3.180 luận điểm và tìm thấy điểm xung đột với các tài liệu khác trong giới học thuật ở 6,85% trong số đó. Trung bình mỗi bài báo ẩn chứa khoảng 2,34 mâu thuẫn thực sự. Trường hợp protein LEF1 đã đề cập ở trên chính là một trong số đó, và ContraCrow đã chấm cho sự xung đột này 9 điểm. Trong một trường hợp khác, chương trình cũng phát hiện ra rằng quan niệm lâu đời cho rằng protein tên là GBP chỉ tồn tại trong tế bào chất đã bị 4 bài báo xuất bản từ cuối những năm 2010 đến năm 2024 lặng lẽ lật ngược. Trên thực tế, đó là một loại protein di chuyển qua nhiều khu vực bao gồm cả màng tế bào. ContraCrow đã chứng kiến từ đầu đến cuối quá trình một quan niệm vững chắc dần dần sụp đổ qua 4 bài báo khoa học.

Khi 100 mâu thuẫn do ContraCrow tìm ra được giấu danh tính và đưa cho các tiến sĩ sinh học thẩm định lại, 70% đồng ý rằng đó thực sự là mâu thuẫn. Điểm F1 đạt 0,82. So với tỷ lệ đồng thuận 75,5% khi hai con người đọc riêng cùng một bài báo và chỉ ra mâu thuẫn, cỗ máy đã đảm đương trọn vẹn vai trò của một chuyên gia. Xét đến việc một trong những lý do khiến các công ty phát triển thuốc lãng phí hàng

tỷ đô la ở giai đoạn thử nghiệm lâm sàng pha 3 chính là bỏ sót những mâu thuẫn tiềm ẩn như vậy, con số 70% trung thực này không hề nhỏ. Điều này cũng có nghĩa là mỗi tuần, có thêm 1,64 mâu thuẫn được chuyên gia xác nhận cuối cùng tích lũy trên mỗi bài báo. Con số này cho thấy tốc độ khoa học thừa nhận mình sai cũng chính là tốc độ khoa học tiến bộ.

Trên nền tảng vững chắc mà ContraCrow đã tạo dựng, FutureHouse không dừng bước. Vào tháng 4 năm 2026, họ đã cho ra mắt DISCO, công cụ thiết kế enzyme xúc tác cho các phản ứng hóa học chưa từng có trong tự nhiên. Chương trình từng sàng lọc mâu thuẫn giữa hàng đồng bài báo nghiên cứu nay đang lớn mạnh thành một đồng nghiệp tự mình xây dựng và kiểm chứng các giả thuyết mới.

Robin và Cosmos

Mã Python đột ngột dừng lại. Đó là khoảnh khắc trí tuệ nhân tạo mang tên Cosmos gặp lỗi thời gian chạy (runtime error) khi đang phân tích dữ liệu di truyền. Trước đây, màn hình sẽ bị đóng băng và cần phải gọi người xử lý. Nhưng Cosmos không làm vậy. Nó tự đọc nhật ký lỗi, xác định vị trí tên cột bị lệch, sửa lại mã và chạy lại. Tất cả diễn ra mà không cần đến bàn tay con người.

Những trí tuệ nhân tạo đầu tiên mà Rodriques và White tạo ra tại FutureHouse là PaperQA và ChemCrow. ChemCrow đã kết hợp khả năng tìm kiếm tài liệu và thực thi mã Python trên nền mô hình ngôn ngữ lớn, thậm chí gửi được ý tưởng thiết kế phân tử đến thiết bị thí nghiệm. Vấn đề bùng phát ở bước tiếp theo. Khi tác vụ vượt quá năm bước, trí tuệ nhân tạo quên mất mình đang tìm kiếm điều gì. Một lỗi mã nhỏ tích tụ dần, và đến khoảng bước thứ 50, nó đưa ra kết luận hoàn toàn sai lệch. Nói về loại trí tuệ nhân tạo này, Rodriques thẳng thắn nhận xét rằng nó chỉ là một trợ lý tiểu thuyết gia tỏ ra thông minh, chứ không phải bộ não của một nhà khoa học thực thụ có thể chịu đựng tính bất định trong phòng thí nghiệm.

Điểm nghẽn thực sự mà Rodriques nhận thấy không phải là tiền bạc hay thiết bị thí nghiệm, mà là số lượng bộ não nhà khoa học được đào tạo có khả năng kiên trì giải quyết những vấn đề hóc búa. Để xóa bỏ điểm nghẽn này, ông đã thành lập một phòng thí nghiệm không người rộng 3.000 foot vuông trang bị robot thí nghiệm. Đó là lời tuyên bố sẽ tạo ra một trí tuệ nhân tạo vận hành tủ hút khí độc và máy phân phối chất lỏng thực tế, chứ không chỉ ghép các mảnh ghép bên trong màn hình máy tính. Mùa xuân năm 2025, tại phòng thí nghiệm này, Robin — hệ thống đa tác tử kết nối từ việc lập giả thuyết đến phân tích dữ liệu thí nghiệm thành một luồng liền mạch — đã ra đời. Vào tháng 11 cùng năm, Cosmos xuất hiện với bộ não được thiết kế hoàn toàn mới từ Robin. Cosmos là sản phẩm chủ lực do Edison Scientific, công ty con vì lợi nhuận của FutureHouse, giới thiệu; và khoảng một tháng sau khi ra mắt, Edison đã huy động được 70 triệu đô la vốn hạt giống.

Điểm khác biệt của Cosmos so với các mô hình trước đó là nó sở hữu một tấm bản đồ gọi là mô hình thế giới có cấu trúc. Thông thường, chatbot chỉ đưa ra một câu trả lời cho một câu hỏi rồi kết thúc, giống như viết câu trả lời lên một mảnh giấy rồi gấp lại. Cosmos thì khác. Khi nhận được một câu hỏi nghiên cứu lớn, nó chia câu hỏi đó thành hàng trăm giả thuyết nhỏ, vẽ đồ thị quan hệ giữa các giả thuyết và luôn giữ bên mình. Khi nhận nhiệm vụ tìm kiếm loại protein giúp tăng tiết insulin ở tuyến tụy, Cosmos chia nhiệm vụ này thành hàng trăm giả thuyết phụ và biểu thị tiến độ của từng nhánh theo thời gian thực. Nếu một giả thuyết

nhất định không mang lại kết quả có ý nghĩa thống kê, Cosmos không dừng lại. Sau khi tự phán đoán và bác bỏ giả thuyết này, nó xem lại bản đồ và rẽ sang hướng đi khác. Quá trình một nhà khoa học con người thử dài khi thí nghiệm thất bại rồi định hình lại suy nghĩ đã được nhóm nghiên cứu của Rodriques đưa trọn vẹn vào bên trong chip silicon.

Việc lấp đầy tấm bản đồ này trên thực tế được phân chia cho các tác tử mang tên loài chim. LiteratureCrow lũng lục hàng chục triệu bài báo và bằng sáng chế để chọn lọc ra các sự thật. Finch nhận

dữ liệu thô từ trình tự gen hoặc máy đo khối phổ để vẽ biểu đồ và tìm kiếm các mối tương quan thống kê. Phoenix và BioPlanner chuyển đổi các giả thuyết đã được kiểm chứng thành quy trình thí nghiệm thực tế, đưa ra mệnh lệnh cho robot thí nghiệm như Opentrons hay phòng thí nghiệm từ xa như Emerald Cloud Lab. Nhạc trưởng Cosmos liên tục kết nối thành quả do ba nhóm tác tử này mang lại vào trong một mô hình thế giới duy nhất. Bằng cách này, vấn đề các trí tuệ nhân tạo trước đây quên mất ngữ cảnh trước đó sau một thời gian đã được khắc phục.

Ngay cả khi mã nguồn bị lỗi, hệ thống này vẫn không dừng lại. Trong mỗi lần chạy, Cosmos vận hành 200 tác tử, tự viết tới 42.000 dòng mã Python, đọc 1.500 bài báo và hoạt động liên tục không nghỉ tới 12 giờ mà không cần con người. Trên thực tế, các nghiên cứu sinh tiến sĩ tại Harvard, MIT và Stanford từng đưa dữ liệu chưa công bố vào Cosmos để thử nghiệm. Chỉ trong vòng 4 đến 12 giờ, Cosmos đã đi đến kết luận tương tự như kết luận mà nhà nghiên cứu con người phải mất 6,1 tháng mới tìm ra. Khi các nhà khoa học độc lập đánh giá kết quả, 79,4% — tức khoảng 8 trên 10 kết luận — đạt độ chính xác. Chi phí cho một lần chạy là 200 đô la.

Giả thuyết do Cosmos tạo ra không chỉ dừng lại trên màn hình. Khi trí tuệ nhân tạo hoàn thành câu khẳng định rằng một protein đột biến cụ thể gây ra một căn bệnh cụ thể, các tác tử BioPlanner và Phoenix sẽ chuyển đổi câu đó thành trình tự DNA thực tế và bản thiết kế kéo chỉnh sửa gen CRISPR. Bản thiết kế đó lập tức được chuyển đến cánh tay robot của phòng thí nghiệm không người. Đó là bước chuyển giao nơi câu văn vận hành trên màn hình biến thành thuốc thử và tế bào trên bàn thí nghiệm. Chiếc cầu nối giữa trí tuệ nhân tạo chỉ lượn quẩn trong máy tính và thiết bị thí nghiệm thực tế đã được Cosmos tự mình bắc qua mà không cần bàn tay con người.

Điều mà Rodriques muốn xóa bỏ thông qua cấu trúc này không chỉ là các lỗi mã nguồn. Phát triển thuốc mới từ lâu đã là lĩnh vực do các hãng dược phương Tây dồi dào vốn và các phòng thí nghiệm đại học danh tiếng độc chiếm. Các trạm y tế ở những quốc gia nghèo không thể tìm được nhân lực trình độ tiến sĩ để nghiên cứu bệnh đặc hữu tại địa phương, nên thậm chí không thể đứng ở vạch xuất phát của việc phát triển thuốc điều

trị. Trí tuệ nhân tạo như Cosmos, thay thế nghiên cứu cấp tiến sĩ chỉ với 200 đô la, được xem là chìa khóa để xóa nhòa khoảng cách này. Chỉ cần một đường truyền internet và một tài khoản, nhà nghiên cứu tại bệnh viện địa phương cũng có thể đưa dữ liệu của mình vào và nhận được danh sách ứng viên protein đích chỉ sau vài giờ. Tất nhiên, liệu dự đoán này có trở thành hiện thực trọn vẹn hay không vẫn cần phải theo dõi thêm.

Bước sang năm nay, Cosmos đã mở rộng quy mô ra cả bên ngoài phòng thí nghiệm. Google DeepMind đã bắt tay với Edison Scientific để cùng tham gia vào công cuộc tạo ra nhà khoa học trí tuệ nhân tạo. Kể từ khi ra mắt vào tháng 11 năm 2025, Cosmos đã tạo ra hơn 10.000 phát hiện khoa học mới, và Rodriques đã có tên trong danh sách 100 nhân vật có tầm ảnh hưởng trong lĩnh vực trí tuệ nhân tạo do tạp chí Time bình chọn. Những trí tuệ nhân tạo mang tên loài chim khởi đầu từ một góc phòng thí nghiệm nay đang bước qua ngưỡng cửa của các phòng nghiên cứu trong ngành dược phẩm.

Những vấn đề vẫn chưa có lời giải

Bên cạnh thành tựu này luôn có một bóng râm bám theo. Hệ thống đọc tài liệu để tổng hợp tri thức không thể chính xác hơn chính những bài báo mà nó đọc. John Ioannidis tại Đại học Stanford, trong một bài báo năm 2005, đã chỉ ra rằng một phần đáng kể các kết quả nghiên cứu được công bố có thể không đúng sự thật. Nguyên nhân là do thiên vị xuất bản, khi các thí nghiệm thành công trở thành bài báo còn các thí nghiệm thất bại lại nằm im trong ngăn kéo. Thêm vào đó là thực trạng các nghiên cứu đăng trên tạp chí học thuật tiếng Anh áp đảo hoàn toàn các cơ sở dữ liệu, còn kết quả viết bằng ngôn ngữ khác ngay từ đầu đã không được chú ý. Việc các bài báo đã bị rút lại (retracted) vẫn tiếp tục được trích dẫn

trong các bài báo khác mà không phản ánh thực tế rút lui là vấn đề mà Ivan Oransky, người sáng lập Retraction Watch, đã chỉ ra từ lâu. Dù cỗ máy có đọc nhanh đến đâu, nếu bản thân tài liệu đọc đã bị lệch về một phía thì sự thiên vị đó sẽ truyền thẳng đến kết luận cuối cùng.

Cách mô hình ngôn ngữ xử lý tài liệu tham khảo cũng có những điểm chưa đáng tin cậy. Mô hình ngôn ngữ thường bịa ra một cách trơn tru tên tác giả, tiêu đề và năm xuất bản của những bài báo không hề tồn tại, hoặc gán ghép trích dẫn vô căn cứ vào một bài báo có thật. Giới học thuật gọi hiện tượng này là ảo giác (hallucination), và liên tục đưa ra cảnh báo về mối nguy hiểm lớn trong các lĩnh vực như y học và luật pháp, nơi một dòng trích dẫn có thể quyết định phán đoán. Việc PaperQA2 được thiết kế để chấm điểm các đoạn văn chứng cứ và hoãn đưa ra câu trả lời khi chưa chắc chắn cũng là cơ chế nhằm giảm thiểu ảo giác ngay từ gốc rễ. Tuy nhiên, nhiều nhà nghiên cứu cùng chung quan sát rằng phương pháp bổ sung căn cứ bằng tìm kiếm chỉ làm giảm ảo giác chứ không thể loại bỏ hoàn toàn.

Các con số mà Cosmos và PaperQA2 đưa ra cũng còn lâu mới đạt đến sự hoàn hảo. Trong các đánh giá độc lập, cứ 5 kết luận của Cosmos thì có 1 kết luận không chính xác, và tỷ lệ WikiCrow để lại câu văn sai dù có gắn nguồn trích dẫn cũng không phải là số không. Nói rằng có ít câu trả lời sai hơn chuyên gia con người không đồng nghĩa với việc không có câu trả lời sai. Vấn đề là ai sẽ sàng lọc những sai sót còn sót lại này. Nếu con người không thể thẩm định lại 1 triệu bài báo mà cỗ máy đã đọc, các kết luận sai có thể đội lốt kết luận đúng và trở thành điểm xuất phát cho những nghiên cứu tiếp theo. Ý kiến cho rằng tốc độ đọc càng nhanh thì càng cần thêm cơ chế giữ lại những điều sai trái vẫn là bài toán mà lĩnh vực này phải giải quyết trong tương lai.

Nội dung trong chương này dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Chương 11. Hiroaki Kitano: Người thiết kế nhà khoa học máy

Nhà tư tưởng của sinh học hệ thống

Vào mùa thu năm 1994, một lá thư đã được gửi tới phòng thí nghiệm tại Đại học Keio, Nhật Bản. Nơi gửi là Science, tạp chí học thuật uy tín hàng đầu thế giới. Mở phong bì ra, Giáo sư Hiroaki Kitano thở dài một tiếng. Bức thư thông báo trả lại bài báo nghiên cứu mà ông cùng Giáo sư Shin-ichiro Imai đã miệt mài thực hiện suốt nhiều tháng. Hai người khẳng định đã tìm ra gen gây lão hóa tế bào bằng mô phỏng máy tính. Phản hồi của ban bình duyệt rất lạnh lùng: Họ không tin vào những phương trình trên màn hình, và yêu cầu hãy chỉ ra tên của gen thực tế mà phương trình đó hướng tới rồi viết lên bảng đen.

Giáo sư Kitano không lùi bước. Ông vốn không phải là một nhà sinh học. Vào những năm 1980, ông là một nhà khoa học máy tính chạy dịch máy trên các máy tính song song tại Viện nghiên cứu NEC. Tại sao một người chuyên về máy tính lại đi sâu vào tế bào? Câu trả lời nằm ở quy mô. Khi quan sát cách thông tin di chuyển bên trong các vi mạch silicon, ông nhận ra rằng cách hàng tỷ protein và chất chuyển hóa va chạm bên trong một tế bào đơn lẻ là một phép tính dày đặc hơn bất kỳ mạch bán dẫn nào. Đầu những năm 1990, ông đã chuyển hẳn lĩnh vực nghiên cứu từ tính toán song song sang khoa học sự sống. Vào thời điểm đó, sinh học vẫn chỉ dừng lại ở cách quan sát từng gen, từng protein dưới kính hiển vi trong thời gian dài. Giáo sư Kitano nhận định rằng cách tiếp cận này không thể giúp hiểu toàn bộ sự sống. Và vào cuối những năm 1990, ông đã sáng lập nên lĩnh vực sinh học hệ thống, xem sự sống không phải là các bộ phận rời rạc mà là một hệ thống thống nhất.

Ngay cả sau khi bị từ chối, hai người vẫn tiếp tục hoàn thiện lý thuyết giải thích lão hóa không phải là sự hư hỏng của từng bộ phận mà là một đặc tính của hệ thống, rồi công bố trên một tạp chí học thuật vào năm 1998. Sau đó, Imai sang Mỹ tham gia nghiên cứu làm sáng tỏ cơ chế hoạt động của gen sir2 điều hòa tuổi thọ của nấm men, và các nghiên cứu về trường thọ liên quan đến sirtuin ngày nay đã phát triển từ nhánh nghiên cứu đó. Đó là khoảng thời gian chứng minh bằng thực tế rằng giả thuyết được xây dựng từ mô phỏng có thể xác định tọa độ trước cả khi tiến hành thí nghiệm.

Thử thách của Giáo sư Kitano không dừng lại ở đó. Năm 1997, ông đã khởi xướng Grand Challenge mang tên giải đấu bóng đá robot RoboCup, và đến năm 2016, ông đã chuyển phương thức cạnh tranh đó sang chính nghiên cứu khoa học, đề xuất một Grand Challenge mới: phát triển nhà khoa học trí tuệ nhân tạo có khả năng tạo ra những phát hiện tầm cỡ giải Nobel. Tên gọi Nobel Turing Challenge đã được định hình qua hội thảo tại London năm 2020 và bài báo khoa học năm 2021.

Năm 2021, Giáo sư Kitano đã tổng hợp và công bố ý tưởng này trên tạp chí học thuật npj Systems Biology and Applications. Thử thách này vẫn đang tiếp diễn. Kết quả của Hội thảo Nobel Turing Challenge Ấn Độ - Nhật Bản diễn ra tại Tokyo vào tháng 3 năm 2025 đã được đăng trên cùng tạp chí vào tháng 4 năm 2026. Nội dung phác thảo bước tiếp theo của khoa học tự hành bằng cách kết hợp ngành robot và trí tuệ nhân tạo chính xác của Nhật Bản với cơ sở hạ tầng dữ liệu quy mô lớn của Ấn Độ. Vào tháng 2 cùng năm, 39 nhà khoa học bao gồm Giáo sư Kitano đã cùng công bố một bài viết trên Frontiers in Artificial Intelligence. Đó là lời cảnh báo rằng dù chu trình thí nghiệm trí tuệ nhân tạo hoàn toàn tự động có thể tìm ra những giả thuyết vượt ngoài trực giác của con người, nhưng cũng tiềm ẩn nguy cơ câu trả lời được đưa ra dưới dạng con người không thể hiểu được. Vì vậy, họ đã đề xuất quyền tự hành theo từng giai đoạn: vận hành thí nghiệm ở tốc độ của máy móc nhưng vẫn gắn chặt vào cấu trúc nhân quả mà con người có thể kiểm chứng. Ba mươi năm sau thời điểm năm 1994 khi tạp chí Science không tin vào mô phỏng, Giáo sư Kitano vẫn đang tiếp tục cuộc chiến tương tự. Chỉ khác là lần này, đó là cuộc chiến tái xác định ranh giới: con người nên tin vào câu trả lời do máy móc đưa ra đến mức độ nào.

Từ một nhà khoa học máy tính từng phải chứng minh trước hội đồng bình duyệt để tìm một gen nấm men, giờ đây ông đang thiết kế cỗ máy có thể tạo ra những phát hiện tầm cỡ Nobel mà không cần đến con người. Quỹ đạo này cho thấy một sự thật: Động lực thúc đẩy Giáo sư Kitano không phải là các thuật toán mới, mà là sự kiên trì bền bỉ, không lùi bước cho đến khi tìm ra câu trả lời ngay cả sau khi bị từ chối. Dù máy móc có trở nên nhanh đến đâu, sự kiên trì đó là điều không thể truyền lại cho máy móc.

Thiết kế Grand Challenge

Viện Khoa học và Công nghệ Sau đại học Okinawa, gọi tắt là OIST. Bước vào một căn phòng trong tòa nhà thí nghiệm ở đây, người ta không thấy bóng dáng con người. Đèn cũng chỉ bật ở mức tối thiểu. Thay vào đó, ba cánh tay robot hoạt động không ngừng nghỉ ngày đêm. Từ khoảnh khắc nghiên cứu viên đặt một ống nghiệm vào cửa nạp, robot sẽ quay ly tâm tế bào, chiết xuất protein và DNA, rồi đưa mẫu vào máy giải trình tự gen. Kết quả được chuyển thẳng lên máy chủ đám mây. Thiết bị này có tên là Dự án Manta. Công việc mà 20 đến 30 nghiên cứu viên lâm sàng lạnh nghề phải mất nhiều ngày mới hoàn thành, Manta tự mình thực hiện 24 giờ mỗi ngày mà không hề run tay hay sai lệch vạch chia.

Người thành lập và hiện vẫn đang dẫn dắt tòa nhà thí nghiệm này là Giáo sư Hiroaki Kitano. Ông vừa điều hành Viện Sinh học Hệ thống, vừa lãnh đạo Viện Nghiên cứu Khoa học Máy tính Sony, đồng thời thành lập cơ sở Manta trong khuôn viên OIST và giữ vai trò nghiên cứu viên phụ trách chính. Đó là cơ cấu trong đó một người đồng thời lãnh đạo nhiều tổ chức. Điều ông mong muốn không phải là một trí tuệ nhân tạo viết nhanh vài bài báo khoa học, mà là toàn bộ một phòng thí nghiệm tự vận hành từ khâu lập kế hoạch thí nghiệm, xử lý mẫu đến phân tích dữ liệu mà không cần bàn tay con người can thiệp.

Đây không phải là lần đầu tiên Giáo sư Kitano đặt ra một mục tiêu lớn như vậy. Năm 1997, ông đã sáng lập RoboCup với mục tiêu: "Đến năm 2050, một đội robot hoàn toàn tự hành sẽ đánh bại đội tuyển vô địch FIFA World Cup". Vào thời điểm đó, chính phủ và giới học thuật Nhật Bản đã chế giễu đó là sự lãng phí tiền thuế. Thế nhưng, khi các đội tham gia công khai mã điều khiển robot trên GitHub và cạnh tranh qua từng năm, những robot từng không thể bước nổi một bước vào năm 1997 đã trở thành những robot có thể quan sát bóng bằng camera và thực hiện các đường chuyền không gian ba chiều vào những năm 2020. Kiva Systems, một công ty robot hậu cần tách ra từ cuộc thi này, sau đó đã được Amazon mua lại và trở thành Amazon Robotics. Tuy nhiên, Giáo sư Kitano nhìn nhận lại rằng RoboCup tuy đã rèn luyện khả năng phối hợp thể chất, nhưng chưa giải quyết được tư duy sâu sắc nhằm đào sâu các mối quan hệ nhân quả trong thời gian dài. Sự suy ngẫm đó đã dẫn đến Nobel Turing Challenge hướng tới nghiên cứu khoa học.

Nền tảng Garuda chính là thứ đã biến thách thức này thành thiết bị thực tế. Tên gọi này bắt nguồn từ Garuda, loài chim cưỡi của thần Vishnu trong thần thoại Ấn Độ. Các chương trình phân tích do phòng thí nghiệm của nhiều trường đại học trên thế giới phát triển vốn có định dạng dữ liệu khác nhau nên không thể kết nối với nhau. Garuda đã liên kết 718 chương trình như vậy vào một cổng duy nhất. Nó cũng kết nối không dây với các thiết bị thí nghiệm thực tế như máy quang phổ khối của Shimadzu Corporation. Kế hoạch thí nghiệm do máy tính lập ra được truyền trực tiếp qua Garuda tới các cánh tay robot trong phòng thí nghiệm. Kèm theo đó là công cụ phân tích bài báo khoa học mang tên Gandhara. Đọc khối lượng lớn bài báo để lập bản đồ quan hệ

nhân quả là chức năng của chương trình này. Mục tiêu mà Gandhara hướng tới là tái hiện bản đồ truyền tín hiệu tế bào—thứ mà các học giả con người đã mất hơn 20 năm đọc 1.500 bài báo để vẽ thủ công—mà không cần sự trợ giúp của con người, đồng thời tự động cập nhật mỗi khi có bài báo mới xuất bản.

Mục tiêu mà cơ sở Manta nhắm tới là độ chính xác. Khi con người thao tác pipet bằng tay, sự run tay và chênh lệch thời gian tiếp xúc với nhiệt độ phòng sẽ tạo ra sai số thí nghiệm rất lớn. Manta lặp lại cùng

một thí nghiệm mà không cần bàn tay con người để triệt tiêu sai số đó. Đây là cách giải quyết vấn đề tái lập vốn tồn tại lâu nay trong các thí nghiệm khoa học sự sống bằng cách loại bỏ sự can thiệp của bàn tay con người.

Có một nghiên cứu cho thấy cội nguồn của cách tiếp cận này. Năm 2017, nhóm nghiên cứu của Giáo sư Kitano đã dùng thử nghiệm 352 thành phần của Ma Hoàng Thang, một bài thuốc Đông y trị cảm cúm, trên động vật thí nghiệm rồi theo dõi trong máu. Chỉ có 19 thành phần đến được máu ở dạng nguyên bản, và ngay cả khi tính cả những chất bị chuyển hóa thay đổi hình dạng thì cũng chỉ có 113 thành phần. Họ đã phân tách theo từng thành phần để xem chất gì đi đến đâu và làm nhiệm vụ gì. Cách tiếp cận xử lý toàn bộ một hệ thống phức tạp gồm hàng trăm thành phần đan xen này chính là ngữ pháp mà Manta và Garuda kế thừa sau này.

Giáo sư Kitano khẳng định toàn bộ hệ thống này khác biệt với trí tuệ nhân tạo kiểu Mỹ. Theo ông, các doanh nghiệp Mỹ chỉ dừng lại ở việc thu thập thật nhiều dữ liệu trên màn hình máy tính để đưa ra dự đoán. Ngược lại, ông cho rằng cần phải kết hợp công nghệ điều khiển robot chính xác mà Sony đã mài giũa lâu năm—tay nghề thủ công của nghệ nhân mà Nhật Bản gọi là Monozukuri—với trí tuệ nhân tạo. Điều này có nghĩa là việc đầu pipet chuyển mẫu chính xác đến mức nào, cánh tay robot di chuyển với sai số bao nhiêu milimét sẽ quyết định kết quả thí nghiệm. Ý niệm của ông rằng tính toán trên màn hình phải song hành cùng đầu ngón tay trên bàn thí nghiệm cũng chính là lý do các khớp nối của Manta vẫn quay không ngừng nghỉ trong những đêm tĩnh lặng ở Okinawa.

Mô hình do Giáo sư Kitano xây dựng không còn chỉ dừng lại ở Okinawa. Vào tháng 6 năm 2026, Scientific American đưa tin rằng tại Phòng thí nghiệm Quốc gia Tây Bắc Thái Bình Dương và Phòng thí nghiệm Quốc gia Oak Ridge của Mỹ, số lượng phòng thí nghiệm nơi robot vận hành thí nghiệm thâu đêm và nghiên cứu sinh nhận thông báo trên giường để điều chỉnh từ xa đang ngày càng tăng. Vào tháng 8 cùng năm, có tin tức cho biết Phòng thí nghiệm Quốc gia Tây Bắc Thái Bình Dương đã đưa ra kế hoạch chuyển toàn bộ các thí nghiệm sang hệ thống robot tự hành trong vòng 5 năm. Mô hình phòng thí nghiệm không người mà Giáo sư Kitano tiên phong xây dựng ở một góc Okinawa đang lan rộng qua bên kia bờ Thái Bình Dương tới các phòng thí nghiệm quốc gia của Mỹ trong chưa đầy một thập kỷ.

Lời hứa năm 2050

Vào tháng 2 năm 2020, các học giả từ khắp nơi trên thế giới đã quy tụ tại phòng họp của Viện Alan Turing ở London. Đó là thời điểm chỉ vài tuần trước khi đại dịch COVID-19 phong tỏa toàn cầu. Giáo sư Hiroaki Kitano của Sony AI, Giáo sư Ross King—người tạo ra các nhà khoa học robot Adam và Eve, cùng Giáo sư Yolanda Gil của Đại học Nam California đã tề tựu với sự tài trợ của Văn phòng Nghiên cứu Hải quân Mỹ và Viện Alan Turing. Tại đây, Giáo sư Kitano đã tuyên bố một mục tiêu: Tạo ra cỗ máy có khả năng tự nghiên cứu mà không cần bàn tay con người để đạt được những phát hiện tầm cỡ giải Nobel Vật lý, Hóa học hoặc Sinh lý học và Y khoa vào năm 2050. Tên gọi của nó là Nobel Turing Challenge 2050. Các học giả có mặt trong phòng họp đã chia sẻ những câu chuyện từ chính phòng thí nghiệm của họ. Giáo sư Ross King chia sẻ kinh nghiệm robot thay thế các thí nghiệm gen vốn tốn hàng tuần lễ của nghiên cứu sinh, còn Giáo sư Yolanda Gil thuật lại những lời than thở của các nhà nghiên cứu trẻ đang ngập chìm trong hàng đống bài báo khoa học.

Lý do thử thách này trở nên đặc biệt là mục tiêu không dừng lại ở khẩu hiệu mà được quy định rõ thành thông số kỹ thuật. Cỗ máy phải vận hành liên tục 24 giờ không ngừng nghỉ qua năm bước: đặt giả thuyết, thiết kế thí nghiệm, trực tiếp tiến hành thí nghiệm bằng cánh tay robot, phân tích kết quả và cập nhật tri thức. Ở bước đặt giả thuyết, cỗ máy kết hợp các quy tắc logic và mô hình xác suất để tự mình chỉ ra những điểm mâu thuẫn hoặc lỗ hổng trong tri thức hiện có. Ở bước thiết kế thí nghiệm, nó tính toán và lựa chọn các điều kiện thí nghiệm nhằm thu được lượng thông tin tối đa trong giới hạn ngân sách và

thời gian. Khoảnh khắc con người can thiệp để điều chỉnh xác suất hay trọng số, tư cách tham gia thử thách sẽ bị hủy bỏ. Bên trong cỗ máy chứa hai bộ não khác nhau: một là bộ suy luận ký hiệu được lập trình bằng các ngôn ngữ logic như Prolog, có nhiệm vụ sàng lọc những giả thuyết vi phạm quy tắc liên kết hóa học hoặc định luật bảo toàn vật lý; bộ não còn lại là mô hình ngôn ngữ đọc hàng trăm triệu bài báo để phát hiện những mối liên hệ ẩn giấu giữa các ngành học khác nhau. Đó là cấu trúc mà bộ suy luận logic đóng vai trò là khiên và mô hình ngôn ngữ đóng vai trò là giáo. Ngay cả khi mô hình ngôn ngữ bịa đặt thông tin sai lệch, chiếc khiên sẽ lọc bỏ để những khẳng định vô căn cứ không thể công bố ra bên ngoài. Ở bước cập nhật tri thức cuối cùng, chu trình này sẽ lặp lại cho đến khi không còn bất kỳ mâu thuẫn vật lý nào tồn tại. Bản đặc tả kỹ thuật phân biệt rõ ràng cỗ máy này với các chatbot thông thường hay các công cụ thống kê chỉ biết khớp đường thẳng vào dữ liệu.

Làm thế nào để đánh giá xem cỗ máy đã thành công hay chưa? Edward Feigenbaum là nhà tiên phong về trí tuệ nhân tạo, người đã tạo ra chương trình Dendral xác định cấu trúc hóa học tại Đại học Stanford vào những năm 1960. Bài kiểm tra ban đầu do ông đặt ra là thước đo xem liệu một chương trình chuyên gia trong một lĩnh vực có đưa ra phán đoán khác biệt so với chuyên gia con người hay không. Giáo sư Kitano và Giáo sư Ross King đã nâng thước đo này lên toàn bộ lĩnh vực khoa

học—gọi là Bài kiểm tra Feigenbaum-Kitano. Một hội đồng giám khảo gồm các nhà khoa học thuộc Viện Hàn lâm Khoa học Quốc gia Mỹ được thành lập. Họ đặt báo cáo phát hiện do máy viết song song với các bài báo của các học giả từ Harvard, MIT, Viện Max Planck để đọc và đánh giá. Nếu hội đồng không thể phân biệt được đâu là bài viết của con người và đâu là bài viết của máy móc, thì cỗ máy mới chính thức vượt qua bài kiểm tra. Nếu như phép thử Turing là bài kiểm tra đánh lừa con người qua đối thoại, thì bài kiểm tra này đánh lừa con người bằng độ sâu của phát hiện và độ hoàn thiện của bài báo khoa học. Trước những bài toán mà số lượng trường hợp khả dĩ lên tới 10 lũy thừa 300—nhiều hơn cả số lượng nguyên tử trong vũ trụ—như cấu trúc protein hay lộ trình tổng hợp vật liệu mới, cỗ máy mượn phương pháp tìm kiếm của AlphaZero khi chơi cờ vây. Nó cắt tỉa các nhánh ít khả năng trước, rồi chỉ chọn đào sâu vào phạm vi hẹp hợp lý về mặt vật lý. Bằng cách này, nó chỉ giữ lại các ứng viên có xác suất phù hợp với các định luật vật lý thực tế trên 80% và loại bỏ phần còn lại.

Bản đặc tả kỹ thuật này không chỉ là giấc mơ viễn vông. Nhà khoa học robot Adam do Giáo sư Ross King giới thiệu vào năm 2009, với tri thức về 6.000 gen nấm men và các phản ứng chuyển hóa, đã tự đặt giả thuyết thí nghiệm gen và làm sáng tỏ bản chất của 12 loại enzyme mờ côi mà con người không giải quyết được trong 150 năm. Phiên bản kế tiếp, Eve, đảm nhận sứ mệnh tìm kiếm thuốc điều trị sốt rét từ năm 2013, đã phát hiện ra rằng triclosan—thành phần phổ biến trong kem đánh răng—ngăn chặn protein DHFR của ký sinh trùng sốt rét mà không cần sự trợ giúp của con người. Cỗ máy thế hệ thứ ba mang tên Genesis đã tiến thêm một bước xa hơn khi vận hành đồng thời 10.000 thiết bị nuôi cấy để tái lập toàn bộ bản đồ biểu hiện gen và chuyển hóa của nấm men.

Vào mùa hè năm 2026, xu hướng này cũng được ghi nhận ở bên ngoài phòng thí nghiệm. Trong bài báo ngày 16 tháng 6, Scientific American đưa tin rằng A-Lab của UC Berkeley tiến hành các thí nghiệm vật liệu mới với tốc độ nhanh gấp 100 lần con người. Đài phát thanh Công cộng Iowa vào tháng 6 cùng năm cũng đưa tin một phòng thí nghiệm tự hành gần Boston đã hoàn thành hơn 30.000 thí nghiệm chỉ trong nửa năm. Đại học Texas A&M vào đầu tháng 8 cũng thông báo họ đã nhận được khoản tài trợ gần 25 triệu USD để xây dựng viện nghiên cứu tự hành phát hiện vật liệu kim loại mới. Tuy nhiên, bài báo của Scientific American cũng chỉ ra trường hợp một bài báo do phòng thí nghiệm tự hành công bố đã bị đình chỉ sau khi được kiểm chứng độc lập từ bên ngoài. Đây chính là lý do Giáo sư Kitano quy định chặt chẽ quy trình bình duyệt mù và mạng lưới kiểm chứng phân tán trong bản đặc tả kỹ thuật. Làm cho nhanh và làm cho đúng là hai vấn đề hoàn toàn khác nhau. Khoảng thời gian còn lại đến năm 2050 của Nobel Turing Challenge sẽ được dùng để kết hợp cả hai yếu tố đó lại làm một. Liệu 24 năm nữa, bộ não silicon có thực sự bước lên bục trao giải tại Stockholm hay không, hiện tại vẫn chưa thể biết trước.

Các học giả máy kết nối

Nicolas Bourbaki không phải là một nhân vật có thật. Đó là bút danh tập thể do khoảng hai mươi nhà toán học kiệt xuất của Pháp lập ra vào đầu thế kỷ 20. Họ không hề tiết lộ tên thật của mình mà chỉ dùng cái tên Bourbaki để xuất bản bộ sách giáo khoa hình học đại số trong suốt nhiều thập kỷ. Giới học thuật đã đọc, kiểm chứng và chấp nhận những cuốn sách đó. Ai viết không quan trọng; điều quan trọng duy nhất là chứng minh đó có đúng hay không. Ngay cả khi xóa bỏ tên tác giả, tri thức vẫn tiếp tục vận hành.

Giáo sư Hiroaki Kitano đưa câu chuyện lịch sử này vào bản thiết kế của học giả máy. Trong hệ sinh thái khoa học năm 2050 mà ông hình dung, các tác nhân nhà khoa học trí tuệ nhân tạo rải rác khắp thế giới sẽ công bố các phát hiện mà không cần tiết lộ tên tuổi hay cơ quan trực thuộc. Điều này tương tự như cách Satoshi Nakamoto, tác giả của sách trắng Bitcoin, đến nay vẫn là một cái tên không lộ mặt. Giáo sư Kitano thậm chí còn đặt tên quy trình xuất bản này theo tên của Nakamoto. Máy móc sẽ tải bài báo giả thuyết viết bằng LaTeX cùng mã robot có thể tái lập chính xác thí nghiệm lên kho lưu trữ công cộng toàn cầu. Không ai hỏi tác giả được tạo nên từ carbon hay silicon, cũng không ai hỏi họ thuộc trường đại học nào hay giáo sư hướng dẫn là ai. Đó là nơi gạt bỏ hoàn toàn gánh nặng về thứ bậc và mạng lưới quan hệ vốn đè nặng lên giới học thuật con người từ lâu nay.

Liệu phát hiện đó có xác thực hay không sẽ do các nút máy tính phán định chứ không phải con người. Trong thiết kế của Giáo sư Kitano, các nút trong mạng lưới kiểm chứng trải rộng khắp thế giới sẽ thực sự tái lập kết quả do máy đưa ra tại phòng thí nghiệm của riêng mình. Tương tự như cách các thợ đào Bitcoin xác thực khối, những nút này xác nhận xem phát hiện có lặp lại được hay không thông qua quy trình bằng chứng công việc. Khi một số lượng nút đủ lớn đạt được cùng kết quả, tri thức đó sẽ được đăng ký thành tài sản chung của nhân loại mà không cần truy xét tác giả. Thay vì con dấu của một thành viên ban bình duyệt, hồ sơ ghi chép các thí nghiệm tương tự được lặp lại độc lập ở khắp nơi trên Trái Đất sẽ trở thành tấm vé thông hành. Quy trình này nằm ở vị trí cuối cùng trong chu trình năm bước đã đề cập ở phần trước. Nếu bốn bước trước là kỹ năng cá nhân vận hành bên trong một cỗ máy đơn lẻ, thì bước cuối cùng là sự liên đới mà các cỗ máy trên khắp hành tinh cùng xác nhận lẫn nhau.

Giáo sư Kitano không chỉ vẽ ra một sự kiểm chứng đơn lẻ như vậy, mà là bức tranh các đội nhóm theo cùng phương thức vận hành đồng thời trên khắp toàn cầu. Các con số mục tiêu được ghi trong bản thiết kế của ông là dữ liệu ở quy mô 1 nghìn tỷ, 1 tỷ giả thuyết và 1 triệu thí nghiệm song song. Quy mô này khác biệt về mặt cấp số so với tốc độ của một nghiên cứu sinh con người phải dành cả đời để kiểm chứng chỉ một hoặc hai giả thuyết. Quy mô này là điều mà phòng máy chủ của một viện nghiên cứu đơn lẻ không thể gánh

vác nổi. Đây là con số chỉ có thể đạt được khi các tài nguyên tính toán và phòng thí nghiệm robot rải rác khắp thế giới được liên kết thành một mạng dữ liệu duy nhất. Đó là viễn cảnh nơi tài nguyên tính toán nhàn rỗi vào ban đêm của các siêu máy tính châu Âu và khoảng thời gian bàn thí nghiệm robot ở Okinawa trống vào ban ngày có thể mượn lẫn nhau. Biên giới quốc gia và sự chênh lệch múi giờ không còn là rào cản trong mạng lưới tính toán này, mà trái lại trở thành điều kiện cho phép vận hành liên tục 24 giờ.

Cũng có bằng chứng cho thấy ý tưởng này không chỉ nằm trên giấy vẽ. Đó là Cosmos do FutureHouse, một viện nghiên cứu độc lập tại San Francisco, giới thiệu vào tháng 11 năm 2025. Tác nhân tìm kiếm tài liệu và tác nhân phân tích dữ liệu cùng chia sẻ một mô hình thế giới và trao đổi thông tin với nhau. Trong một lần vận hành, nó thực hiện 20 chu trình, viết 42.000 dòng mã và đọc khoảng 1.500 bài báo khoa học. Một kết quả duy nhất sau 12 giờ vận hành không nghỉ như vậy tương đương với thành quả mà một nhà nghiên cứu con người phải mất trung bình 6 tháng miệt mài mới đạt được. Khi các nhà khoa học độc lập đối chiếu từng câu trong báo cáo, 79,4% đã được xác nhận là sự thật, và bảy phát hiện mới thuộc các lĩnh vực chuyển hóa học, vật liệu mới và khoa học thần kinh đã cùng xuất hiện trong một lần

chạy duy nhất.

Vào tháng 5 năm 2026, Robin—một hệ thống đa tác tử khác do FutureHouse phát triển—đã thực sự công bố một bài báo trên Nature. Bốn tác tử mang tên Crow, Falcon, Owl và Finch đã cộng tác để tìm ra một loại thuốc trị tăng nhãn áp hiện có tên là ripasudil như một ứng viên điều trị thoái hóa điểm vàng, đồng thời đề xuất một con đường mới giúp phục hồi chức năng của tế bào võng mạc bằng cách điều hòa nhịp sinh học ngày đêm. Thời gian từ lúc đưa ra giả thuyết đến khi bài báo được đăng chỉ mất hai tháng rưỡi. Điều đó có nghĩa đây là một trường hợp vận hành thực tế chứ không chỉ là bản vẽ thiết kế. Tuy nhiên, một bài báo khác xuất bản cùng năm đã chỉ ra rằng bản thân khái niệm nhà khoa học máy tự hành là quá lạc quan. Bài báo nhấn mạnh rằng để mạng lưới hợp tác thực sự vận hành, vai trò của con người trong việc chia sẻ dữ liệu công bằng và điều phối những người tham gia vẫn rất lớn. Việc máy móc công bố bài báo ẩn danh và việc quyết định ai sẽ phân bổ dữ liệu cùng tài nguyên tính toán cho những cỗ máy đó là hai vấn đề hoàn toàn khác nhau. Đằng sau bút danh Bourbaki, rốt cuộc cũng có hai mươi con người tụ họp trong phòng họp để điều phối bản thảo. Xóa tên tác giả không có nghĩa là sự điều phối tự khắc biến mất. Những người tạo ra Kosmos và Robin cũng phải lặp lại công việc hằng ngày là chi trả chi phí máy chủ, bảo trì cánh tay robot và sửa mã nguồn. Đằng sau các phòng thí nghiệm không người vào năm 2050, vị trí của con người kiểm tra hóa đơn tiền điện vẫn sẽ còn đó.

Những vấn đề vẫn chưa có lời giải

Xung quanh mục tiêu năm 2050 do Thử thách Nobel Turing đặt ra, trong giới học thuật đã có những tiếng nói hoài nghi liệu đó có thực sự là một đặc tả có thể kiểm chứng được hay không. Một bài viết do Sebastian Musslick cùng hơn mười nhà nghiên cứu khác đồng công bố trên Kỷ yếu Viện Hàn lâm Khoa học Quốc gia Hoa Kỳ vào năm 2025 đã thừa nhận tiềm năng mà tự động hóa khoa học mang lại, nhưng cũng chỉ ra rằng máy móc hiện nay vẫn chỉ dừng lại ở giai đoạn giải các bài toán do con người thiết lập sẵn. Dù bề ngoài việc đưa ra giả thuyết và thiết kế thí nghiệm có vẻ gần với sự tự hành, nhưng ở vị trí quyết định nên đặt ra câu hỏi nào và điều gì được coi là một câu trả lời thỏa đáng, sự phán đoán của con người vẫn in dấu sâu đậm. Đó là lời chỉ ra rằng khoảng cách giữa quyền tự hành hoàn toàn và quá trình tự động hóa vận hành trên một nền tảng do con người dựng sẵn xa hơn nhiều so với suy nghĩ của chúng ta.

Cũng có những ý kiến chỉ ra sự mơ hồ trong tiêu chí đánh giá. Cụm từ 'phát kiến tầm cỡ giải Nobel' tuy tạo được tiếng vang lớn, nhưng thước đo để công nhận điều gì đạt đến tầm mức đó lại không rõ ràng. Bản thân giải Nobel là một sự đánh giá mang tính hồi tố sau nhiều thập kỷ kiểm chứng tiếp theo và đồng thuận của giới học thuật, nên rất khó để ấn định một phát kiến có xứng đáng với giải thưởng hay không tại thời điểm trao giải làm thước đo cho năng lực của máy móc. Phương thức thử nghiệm cho rằng chỉ cần con người không phân biệt được bài báo do người viết với báo cáo do máy viết là vượt qua bài kiểm tra cũng vấp phải phản biện: đó chỉ là thước đo mức độ hoàn thiện của văn bản chứ không phải là thước đo trực tiếp sức nặng khoa học của phát kiến.

Cũng có những lời chỉ trích nhắm vào chính phương thức Đại Thử thách. Một bài viết của Channing và Gosh công bố năm 2026 khẳng định rằng ý tưởng về một nhà khoa học máy tự hành là vấn đề của xã hội trước khi là vấn đề công nghệ, đồng thời chỉ ra rằng một mục tiêu xa vời như năm 2050 tuy có sức mạnh như một khẩu hiệu truyền cảm hứng để tập hợp mọi người về một hướng, nhưng nó khác với một lộ trình thực nghiệm cụ thể để đi đến đó. Giống như khẩu hiệu năm 2050 của RoboCup, một mục tiêu lớn đóng vai trò như ngọn cờ quy tụ tài nguyên và nhân tài, nhưng ngọn cờ không đồng nghĩa với con đường. Câu hỏi rằng những bước tiến hướng tới mục tiêu thực sự đã đi đến đâu chỉ có thể được xác nhận bằng các kết quả có thể tái lập chứ không phải bằng khẩu hiệu, và câu hỏi ấy vẫn còn đó đối với Thử thách Nobel Turing.

Nội dung mô tả trong chương này dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Ghi chú của tác giả. Vai trò của con người đã dịch chuyển về đâu?

Hai nhân vật trong Phần 5 đã bắt tay vào việc kết nối các mắt xích để khép kín vòng lặp. Kosmos của Rodriguez mức lấy tri thức từ đại dương tài liệu mà con người không thể đọc xuể, còn Thử thách Nobel Turing của Kitano đặt ra mục tiêu về một nhà khoa học máy tự mình phát kiến vào năm 2050. Trạng thái mà vòng tuần hoàn từ quan sát đến kiểm chứng quay tròn một vòng mà không cần con người—vị trí mà cuốn sách này gọi là vòng lặp khép kín—chính là đích đến mà cả hai đang bước tới.

Tính tự hành có bản chất khác biệt so với bốn giai đoạn trước. Từ tái khám phá cho đến thí nghiệm đều theo khuôn mẫu: con người đưa ra bài toán và máy móc giải quyết. Khi đạt đến sự tự hành, ngay cả việc lựa chọn bài toán tiếp theo cần giải cũng trở thành phần việc của máy móc. Đó là giai đoạn mà máy móc lần đầu tiên chạm tay vào điều mang tính con người nhất trong khoa học: tò mò về điều gì. Phần lớn các cuộc tranh luận được đề cập ở phần sau của cuốn sách này đều bắt nguồn từ đây.

Vì vậy, trong Phần 5, vị trí của con người rốt cuộc thu hẹp lại chỉ còn một: vị trí lựa chọn câu hỏi. Đôi tay đã trao cho robot, đôi mắt trao cho thuật toán, ký ức trao cho cơ sở dữ liệu, nhưng vị trí quyết định nên hỏi điều gì vẫn thuộc về con người. Dẫu vậy, hai nhân vật trong Phần 5 cho thấy ngay cả vị trí cuối cùng này cũng có thể bị máy móc dòm ngó. Liệu có nên bảo vệ vị trí đó hay nhượng lại cho máy móc, đó là câu hỏi mà chương tiếp theo và phần kết luận sẽ tiếp tục tiếp nối.

Chương 12. Ai là chủ nhân của phát hiện?

Chủ nhân của những gì máy móc phát hiện

Mười một chương trước đã cho thấy máy móc có thể làm được những gì. Chương này đề cập đến câu hỏi còn lại sau khi máy móc hoàn thành một điều gì đó. Khi máy móc tự đặt giả thuyết, thiết kế vật chất, vận hành thí nghiệm và viết bài báo khoa học, câu hỏi đặt ra là ai là chủ nhân của kết quả đó? Câu hỏi này không còn là vấn đề của phòng thí nghiệm mà là vấn đề của luật pháp và thể chế. Và hiện tại, vị trí ghi câu trả lời đó hầu như vẫn còn để trống.

Với tư cách là một luật sư, tôi đã quan sát khoảng trống này rất lâu. Luật pháp từ trước đến nay luôn được xây dựng dựa trên tiền đề về con người. Người nắm giữ quyền lợi là con người, người chịu trách nhiệm là con người, và người để lại tên tuổi cũng là con người. Khi máy móc chen vào vị trí đó, những câu chữ luật pháp lâu đời bắt đầu kêu cọt két. Chương này không phải là nơi đưa ra câu trả lời chuẩn xác. Đây là nơi mở ra những lựa chọn đang có và những cái giá đi kèm với từng lựa chọn đó. Quyết định cuối cùng thuộc về xã hội. Nhưng nếu chúng ta càng trì hoãn quyết định, khoảng trống sẽ càng bị lấp đầy trước bởi những sự thật do máy móc tạo ra.

Ai là chủ nhiệm nghiên cứu của giả thuyết?

Trong Chương 8, Genesis của Ross King tự mình đưa ra và kiểm chứng hàng nghìn giả thuyết mỗi ngày. Giả sử một trong số đó được chứng minh là một phát hiện mới mẻ. Vậy ai là chủ nhiệm nghiên cứu của phát hiện này?

Theo thông lệ lâu đời của luật pháp và giới học thuật, chủ nhiệm nghiên cứu là người đưa ra giả thuyết. Nhưng ở đây, chính máy móc mới là bên đưa ra giả thuyết. Con người đã chế tạo ra máy móc, nạp dữ liệu vào máy móc và phê duyệt việc sẽ thí nghiệm giả thuyết nào trong số các giả thuyết mà máy móc lựa chọn. Vấn đề đặt ra là nên coi con người này là chủ nhiệm nghiên cứu, hay cần ghi nhận máy móc – thực thể tạo ra bản chất của phát hiện – vào hồ sơ theo một cách nào đó.

Hai lựa chọn đều đi kèm với những cái giá riêng. Nếu ấn định con người là người chịu trách nhiệm, thông lệ sẽ được duy trì, nhưng khoảng cách giữa công việc thực tế mà người đó đã làm và công lao được ghi tên sẽ ngày càng nở rộng. Nếu bắt đầu công nhận sự đóng góp của máy móc dưới bất kỳ hình thức nào, khoảng cách đó sẽ được lấp đầy, nhưng tiền đề cơ bản của luật pháp rằng máy móc không phải là chủ thể của quyền lợi sẽ bị lung lay. Hiện nay, hầu hết các viện nghiên cứu chưa đưa ra lựa chọn này một cách rõ ràng, mà đang lửng lơ trôi theo hướng coi máy móc là công cụ và chỉ ghi tên con người. Nhưng việc để mọi thứ tự trôi đi và việc chủ động quyết định là hai điều hoàn toàn khác nhau.

Sự phân biệt này quan trọng như thế nào sẽ bộc lộ rõ vào thời điểm tiền bạc và công trạng xuất hiện. Khi một giả thuyết do máy móc đưa ra được cấp bằng sáng chế và trao giải thưởng, tên của chủ nhiệm nghiên cứu sẽ lập tức trở thành tên của chủ sở hữu quyền lợi. Nếu thông lệ coi máy móc là công cụ và chỉ ghi tên con người trở nên cố định, thì trên thực tế, công trạng và lợi nhuận từ những việc máy móc làm sẽ dồn về một số ít người sở hữu hoặc vận hành cỗ máy đó. Ngược lại, nếu xây dựng một thể chế ghi nhận riêng sự đóng góp của máy móc, việc phân bổ công trạng sẽ gần với sự thật hơn, nhưng lại phát sinh thêm gánh nặng mới về việc ai sẽ quản lý và kiểm chứng những ghi chép đó như thế nào. Tôi xem đây vừa là vấn đề học thuật vừa là vấn đề phân phối. Bởi lẽ, việc ghi nhận công trạng của phát hiện do máy móc tìm ra dưới tên ai không thể tách rời khỏi việc phân chia lợi ích do phát hiện đó mang lại cho ai.

Tai nạn trong phòng thí nghiệm robot là tai nạn của ai?

Trong Chương 9, Coscientist (Gomes) của Gabe Gomes điều khiển cánh tay robot thực tế theo chỉ thị của mô hình ngôn ngữ để tạo ra phản ứng hóa học. Hệ thống này có sẵn các cơ chế an toàn nhằm ngăn chặn việc tổng hợp các chất nguy hiểm, bởi những người chế tạo ra nó đã lường trước rủi ro đó. Tuy nhiên, cơ chế an toàn không thể ngăn chặn mọi tình huống. Giả sử máy móc thiết lập một quy trình ngoài dự kiến và gây ra tai nạn. Giả sử khí độc rò rỉ, phản ứng mất kiểm soát hoặc chất được tổng hợp sai gây thương tích cho con người. Tai nạn này là trách nhiệm của ai?

Trong thế giới của luật trách nhiệm pháp lý, có một vài hướng tiếp cận. Chúng ta có thể quy trách nhiệm cho nhà phát triển đã tạo ra máy móc, có thể quy trách nhiệm cho nhà nghiên cứu đã vận hành cỗ máy đó trong phòng thí nghiệm, hoặc có thể quy trách nhiệm cho cơ quan đã phê duyệt việc trao quyền tự chủ cho máy móc. Con đường coi đây là khuyết tật của sản phẩm và con đường coi đây là sự lơ là quản lý của người sử dụng sẽ rẽ sang hai hướng khác nhau, và tùy thuộc vào sự phân nhánh đó mà mức độ bồi thường cũng như đối tượng chịu trách nhiệm sẽ thay đổi.

Máy móc càng phán đoán tự chủ thì ranh giới này càng trở nên mờ nhạt. Giống như sự cố tăng tốc đột ngột của ô tô, nếu nguyên nhân nằm bên trong máy móc thì trách nhiệm sản phẩm sẽ nặng hơn; còn nếu nguyên nhân là do sự bất cẩn của người lái thì trách nhiệm của người dùng sẽ nặng hơn. Tuy nhiên, trong các phòng thí nghiệm vòng lặp khép kín, có những trường hợp nguyên nhân nằm ở chính sự phán đoán tự chủ của máy móc. Đó là phán đoán mà người chế tạo không lường trước được và người vận hành cũng không hề ra chỉ thị. Lúc này, không phải đối tượng chịu trách nhiệm biến mất, mà là chúng ta chưa có phương pháp để phân chia trách nhiệm.

Trong số các công cụ pháp lý lâu đời, khái niệm gần nhất với tình huống này là trách nhiệm rủi ro. Người nuôi thú dữ phải chịu trách nhiệm đối với người bị thú cắn mà không cần xét đến lỗi của bản thân. Đó là nguyên tắc: nếu quản lý một thứ nguy hiểm vì lợi ích của mình thì khi rủi ro trở thành hiện thực, người đó phải chịu trách nhiệm. Đặt phòng thí nghiệm tự hành vào khuôn khổ này, ta đi đến kết luận rằng tai nạn do phán đoán của máy móc gây ra sẽ do cơ quan vận hành cỗ máy đó vì lợi ích chịu trách nhiệm. Tuy nhiên, khuôn khổ này cũng có cái giá của nó. Nếu áp đặt trách nhiệm rủi ro quá nặng, các tổ chức sẽ không dám vận hành phòng thí nghiệm; nếu áp đặt quá nhẹ, con đường nhận bồi thường của nạn nhân sẽ bị thu hẹp. Đặt quả cân vào đâu không phải là vấn đề công nghệ, mà là vấn đề xác định mức độ rủi ro mà xã hội sẵn sàng chấp nhận.

Ba nguyên tắc được nêu trong phần kết luận – xác định trước phạm vi tự chủ, ghi lại quá trình phán đoán vào hồ sơ và để con người ký tên xác nhận cuối cùng – là những cơ chế tối thiểu để lấp đầy khoảng trống này. Phải có phạm vi và hồ sơ ghi chép thì khi xảy ra tai nạn, chúng ta mới có thể xem xét lại xem ranh giới đã bị vượt qua ở điểm nào để phân chia trách nhiệm,

bởi vì tự chủ mà không có hồ sơ ghi chép là sự tự chủ không thể quy trách nhiệm, và một sự tự chủ không thể quy trách nhiệm thì xã hội sẽ không chấp nhận lâu dài.

Ai là nhà phát minh của loại thuốc mới do trí tuệ nhân tạo thiết kế?

Trong Chương 2 và Chương 4, AlphaFold và RFdiffusion thiết kế các protein và ứng viên thuốc mới chưa từng tồn tại trong tự nhiên. Nếu nộp đơn xin cấp bằng sáng chế cho các chất này, chúng ta phải ghi tên ai vào ô nhà phát minh?

Hệ thống bằng sáng chế từ trước đến nay luôn ấn định nhà phát minh phải là con người. Trong vài năm gần đây, cơ quan sáng chế và tòa án ở nhiều quốc gia đã liên tiếp bác bỏ các đơn đăng ký ghi nhận trí tuệ nhân tạo là nhà phát minh. Điều này tái khẳng định nguyên tắc nhà phát minh phải là thể nhân. Nguyên tắc này rất rõ ràng, nhưng sự rõ ràng không đồng nghĩa với việc hoàn toàn trùng khớp với thực tế. Khi máy móc đảm nhận phần cốt lõi của thiết kế, tên của nhà phát minh là con người sẽ chỉ phản ánh một phần của sự thật.

Từ đây, các lựa chọn rẽ sang các hướng khác nhau. Con đường duy trì con người làm nhà phát minh giúp hệ thống ổn định, nhưng sự đóng góp của máy móc càng lớn thì tính chân thực của việc ghi danh nhà phát minh càng suy giảm. Con đường tạo ra một danh mục mới để ghi nhận riêng sự đóng góp của máy móc giúp tiến gần hơn đến sự thật, nhưng chúng ta sẽ phải tái cấu trúc lại toàn bộ hệ thống quyền lợi và trách nhiệm vốn dành cho nhà phát minh. Cả hai con đường đều không có lựa chọn nào là miễn phí.

Đối với vấn đề này, còn có một con đường thứ ba là hoàn toàn không cấp bằng sáng chế. Nếu nhà phát minh bắt buộc phải là con người nhưng bản chất phát minh lại do máy móc tạo ra, thì phát minh đó sẽ không thuộc về ai và nằm ngoài sự bảo hộ của bằng sáng chế. Con đường này trung thành với nguyên tắc nhưng lại dẫn đến những hệ quả bất ngờ, bởi nó thúc đẩy các doanh nghiệp che giấu vai trò của máy móc trong quá trình phát minh. Nếu tiết lộ rằng máy móc đã thiết kế thì sẽ không được cấp bằng sáng chế, từ đó tạo ra động cơ dằn dặt như thể con người đã thiết kế. Đó là một nghịch lý khi một nguyên tắc đòi hỏi sự thật lại khiến người ta che giấu sự thật. Tôi cho rằng nghịch lý này chính là lý do khiến thể chế hiện nay cần khẩn trương đưa ra câu trả lời. Điều rõ ràng là tốc độ chuyển dịch bản chất của phát minh sang máy móc đang nhanh hơn tốc độ thể chế đưa ra lời giải. Trong khoảng cách đó, việc ai nắm giữ quyền lợi và ai hưởng lợi đang dần được định hình trước.

Nếu tác giả bài báo khoa học là máy móc, ai là chủ thể của hành vi sai phạm?

Sự việc của Sakana AI được nhắc thoáng qua trong phần kết luận đã đặt ra câu hỏi này một cách sâu sắc. Năm 2025, một bài báo khoa học do trí tuệ nhân tạo viết từ đầu đến cuối mà không qua bàn tay con người đã vượt qua vòng bình duyệt ẩn danh tại một hội thảo của hội nghị quốc tế. Điểm số mà ba phản biện đưa ra lần lượt là 6, 6 và 7 điểm. Các phản biện đã được thông báo trước rằng trong số các bài nộp có thể có bài do trí tuệ nhân tạo viết, nhưng họ không biết bài nào là bài viết đó. Công ty phát triển đã chủ động rút bài báo này trước khi xuất bản theo đúng thỏa thuận trước đó với hội nghị và ủy ban đạo đức nghiên cứu. Năm sau đó, bài báo tổng kết về hệ thống này đã được đăng trên tạp chí Nature, và công ty tuyên bố sẽ gắn hình mờ vào các bài viết do máy móc tạo ra.

Tuy nhiên, bài báo này có một tí vết. Một tài liệu tham khảo đã ghi sai nhà phát minh của công nghệ mạng bộ nhớ ngắn-dài hạn thành một người hoàn toàn không liên quan. Nhà phát minh thực sự là hai nhà nghiên cứu vào năm 1997, nhưng máy móc lại gắn tên của người khác. Đó là một kiến thức phổ thông mà bất kỳ ai trong ngành cũng biết, nhưng máy móc lại mắc sai sót. Lỗi này đã bị phát hiện qua quá trình rà soát của chính đội ngũ nghiên cứu Sakana.

Nếu sai sót như vậy không được sàng lọc mà vẫn được xuất bản, liệu đó có phải là hành vi sai phạm trong nghiên cứu? Nếu là sai phạm thì ai là chủ thể? Đạo đức nghiên cứu trừng phạt những người có hành vi cố ý hoặc bất cẩn nghiêm trọng. Nhưng ở đây, thực thể bịa ra câu chữ là máy móc, và máy móc không có ý đồ cố ý. Người vận hành máy móc không trực tiếp viết ra câu chữ đó. Khái niệm sai phạm vốn được xây dựng dựa trên tiền đề về một tác nhân có nhận thức và tâm trí, nhưng giờ đây ở vị trí đó lại là một bộ tạo sinh không có tâm trí.

Ở đây cũng vậy, câu trả lời chỉ có thể tìm thấy từ phía con người. Đó là nguyên tắc: người nộp bài viết do máy móc tạo ra dưới danh nghĩa của mình phải chịu trách nhiệm về tính chân thực của bài viết đó. Việc không lọc được các sai sót do máy móc bịa ra không phải là tội của máy móc, mà thuộc về trách nhiệm của người đã nộp bài mà không qua kiểm chứng. Chỉ khi nguyên tắc này được xác lập thì sự tin cậy của học thuật mới không bị sụp đổ trước làn sóng bài báo do máy móc tạo ra ào ạt. Như chúng ta đã thấy trong Chương 10, ngay cả cỗ máy đọc tài liệu cũng trích dẫn sai tài liệu tham khảo, thì nguyên tắc con người ký tên xác nhận cuối cùng không phải là hình thức, mà là một phòng tuyến bảo vệ.

Nên đặt những gì vào diện cần con người phê duyệt trước?

Cốt lõi của thí nghiệm tự hành nằm ở chỗ con người không can thiệp vào từng giai đoạn. Chính vì không can thiệp nên tốc độ mới nhanh hơn, và cũng chính vì không can thiệp nên mức độ nguy hiểm mới tăng lên. Do đó, giao phó cho máy móc đến đâu và từ đâu thì cần sự phê duyệt trước của con người trở thành trọng tâm của việc thiết kế.

Ranh giới này không thể vạch ra một lần là xong. Câu trả lời thực tế là tính tự chủ theo từng cấp độ: để máy móc vận hành đến cùng đối với các thí nghiệm có độ rủi ro thấp, và yêu cầu sự phê duyệt của con người đối với các thí nghiệm có độ rủi ro cao. Vấn đề là ai sẽ là người xác định ranh giới đó. Nếu nhà phát triển quyết định, sự tiện lợi kỹ thuật sẽ được ưu tiên; nếu cơ quan quản lý quyết định, sự an toàn sẽ được đặt lên hàng đầu nhưng tốc độ sẽ chậm lại; nếu viện nghiên cứu quyết định, quy định sẽ sát với thực tế hiện trường nhưng tiêu chuẩn giữa các viện lại khác nhau.

Với tư cách là một luật sư, điều tôi muốn đề xuất là một quy trình xác định ranh giới này bằng văn bản và công khai trước khi thí nghiệm bắt đầu. Khi những gì máy móc được phép tự quyết định được lưu lại trong tài liệu, trách nhiệm pháp lý khi xảy ra tai nạn sẽ trở nên rõ ràng, và ngay cả khi không có tai nạn, xã hội cũng có cơ sở để tin tưởng vào thí nghiệm đó. Đây không chỉ là quy định ràng buộc máy móc, mà trước hết là cơ chế bảo vệ chính những người thực hiện thí nghiệm. Bởi chỉ khi phạm vi ủy thác được ghi rõ ràng, nhà nghiên cứu mới biết được phạm vi mình phải chịu trách nhiệm để đặt bút ký.

Thể chế của Hàn Quốc đã sẵn sàng trả lời chưa?

Những câu hỏi nêu trên áp dụng cho bất kỳ quốc gia nào. Cuối cùng, chúng ta hãy nhìn vào vị thế của Hàn Quốc.

Hệ thống nghiên cứu và hệ thống bằng sáng chế của Hàn Quốc nhìn chung chưa sẵn sàng để trả lời những câu hỏi này. Việc chưa chuẩn bị không hẳn là điều đáng xấu hổ, mà có nghĩa là chưa có ai trực tiếp đối diện với câu hỏi này. Các quy định về đạo đức nghiên cứu được xây dựng dựa trên tiền đề về nhà nghiên cứu là con người, và các quy định về nhà phát minh trong luật bằng sáng chế cũng dựa trên tiền đề về thể nhân. Các điều khoản xử lý trách nhiệm đối với giả thuyết do máy móc đưa ra, tai nạn trong phòng thí nghiệm robot hay quyền đối với phát minh do máy móc thiết kế vẫn chưa được định hình.

Tuy nhiên, khoảng trống này vừa là bằng chứng của sự tụt hậu vừa là một cơ hội, bởi chưa một quốc gia nào trên thế giới hoàn thiện các chuẩn mực cho khoa học máy. Nếu có một quốc gia tiên phong đưa ba nguyên tắc đã thảo luận ở trên – xác định trước phạm vi tự chủ, ghi chép quá trình phán đoán và con người ký tên xác nhận cuối cùng – vào quy định đạo đức nghiên cứu và tiêu chuẩn thẩm định bằng sáng chế, thì hệ thống quy tắc của quốc gia đó sẽ trở thành điểm tham chiếu cho các nước đi sau. Đúng như đã nêu trong phần kết luận, bên nào viết lên trang giấy trắng trước sẽ là bên tạo ra tiêu chuẩn.

Hàn Quốc đã hội tụ đủ các yếu tố nền tảng. Quốc gia này có nền tảng sản xuất, dữ liệu và năng lực tính toán để xây dựng các phòng thí nghiệm tự hành. Thứ còn thiếu không phải là trang thiết bị mà là các chuẩn mực. Việc tiên phong viết nên những câu chữ quy định ai là chủ nhân của những phát hiện do máy móc tạo ra, ai sẽ chịu trách nhiệm cho các tai nạn đó, và trao quyền lợi đó cho ai – chính công việc ấy sẽ quyết định vị thế của Hàn Quốc trong các quy tắc khoa học của nửa thế kỷ tới.

Cụ thể có ba điểm cần can thiệp. Thứ nhất là quy định quản lý nghiên cứu và phát triển quốc gia. Nếu yêu cầu nộp phạm vi và hồ sơ ghi chép các thí nghiệm do máy móc thực hiện tự chủ như một phần của kết quả nghiên cứu, thì mức độ tự chủ sẽ được lưu lại trong tài liệu. Thứ hai là hướng dẫn đạo đức nghiên cứu. Nếu quy định rõ ràng rằng trách nhiệm kiểm chứng các câu chữ và dữ liệu do máy móc tạo ra thuộc về người nộp, thì khi lỗi do máy móc bịa ra trở thành hành vi sai phạm, đối tượng chịu trách

nhiệm sẽ rất rõ ràng. Thứ ba là tiêu chuẩn thẩm định bằng sáng chế. Nếu thiết kế sao cho vai trò của trí tuệ nhân tạo trong quá trình phát minh phải được nêu rõ trong đơn đăng ký nhưng không lấy đó làm lý do từ chối cấp bằng sáng chế, chúng ta có thể tránh được nghịch lý che giấu đã đề cập ở trên. Cả ba điểm này đều có thể bắt đầu từ việc điều chỉnh các hướng dẫn và tiêu chuẩn hiện có trước khi ban hành luật mới.

Để không để lại khoảng trống

Chương này không đưa ra câu trả lời khẳng định tuyệt đối. Việc xác định ai là chủ nhân của phát hiện do máy móc tạo ra, phân chia trách nhiệm tai nạn như thế nào, hay ghi gì vào ô nhà phát minh là việc mà xã hội phải quyết định, và mỗi quyết định đó đều có cái giá riêng. Tuy nhiên, có một điều chắc chắn: nếu chúng ta không trả lời những câu hỏi này, thì chính việc không trả lời sẽ trở thành một câu trả lời. Những sự thật do máy móc tạo ra sẽ lấp đầy các khoảng trống trước, và con người sau đó sẽ phải tranh chấp muện màng về quyền lợi và trách nhiệm trên những sự thật đã rồi đó.

Nếu mười một chương trước đã cho thấy con người đã chuyển giao những gì cho máy móc, thì câu hỏi mà chương này đặt ra là: điều gì chúng ta không được phép chuyển giao? Niềm vui phát hiện có thể trao đi. Nhưng trách nhiệm đối với phát hiện thì không thể chuyển giao. Bởi vì chỉ có con người mới có thể gánh vác trách nhiệm. Chừng nào chúng ta không quên đi sự thật đó, thì dù máy móc có phát hiện ra bao nhiêu điều đi chăng nữa, chủ nhân của những phát hiện ấy vẫn mãi là con người.

Nội dung trong chương này dựa trên các tài liệu được công bố tính đến tháng 8 năm 2026.

Kết luận. Vượt qua sự phong phú của dự đoán không có hiểu biết

Cái giá của sự phong phú

Đã đến lúc quay trở lại với mệnh đề được đặt ra ở phần đầu cuốn sách này. Bản chất của việc trí tuệ nhân tạo làm thay đổi khoa học không nằm ở việc tính toán nhanh hơn. Nó nằm ở chỗ vòng lặp nối từ quan sát đến giả thuyết, dự đoán, thí nghiệm và kiểm chứng lần đầu tiên bắt đầu được khép lại bên trong cỗ máy. Dẫu theo 11 nhà khoa học cho đến lúc này, câu nói trên không còn là một lời tuyên bố đơn thuần nữa mà đã là một lời chứng thực. Hassabis, Jumper và Barzilay đã giao phó dự đoán cho máy móc. Baker và Aspuru-Guzik giao phó việc thiết kế, Lipson và Tegmark giao phó việc tìm kiếm quy luật, còn King và Gomes giao phó đôi bàn tay trên bàn thí nghiệm. Rodriguez và Kitano đã bắt tay vào việc nối liền các mắt xích đó để khép lại vòng lặp.

Thế nhưng, sự phong phú này luôn đi kèm với một cái giá. Phương trình mà cỗ máy của Lipson tìm thấy trong chu trình chuyển hóa của nấm men chính xác hơn bất kỳ mô hình nào do con người tạo ra, nhưng không ai giải thích được các biến số trong phương trình đó có ý nghĩa gì bên trong tế bào. AlphaFold đã đưa ra 200 triệu cấu trúc protein, nhưng chính những người tạo ra nó cũng khẳng định rằng mô hình này không phải là một mô hình nhân quả giải thích nguyên lý cuộn gập. Câu trả lời tuôn trào nhưng lý do lại không theo kịp. Tốc độ dự đoán đã vượt qua tốc độ hiểu biết. Cuốn sách này gọi trạng thái đó là sự phong phú của dự đoán không có hiểu biết.

Cái giá phải trả không chỉ nằm ở tầng mức tri thức. Nó còn phải trả ở tầng mức con người. Trong quá trình thức trắng đêm trước bàn thí nghiệm và tự mình truy xét lại lý do tại sao dữ liệu bị sai, các nhà nghiên cứu trẻ đã rèn luyện được trực giác bằng chính trải nghiệm thực tế. Có những cảm giác mà chỉ người từng trộn sai thuốc thử và làm hỏng thí nghiệm tới ba lần mới hiểu được. Khi máy móc trao trọn vẹn câu trả lời hoàn chỉnh, toàn bộ quá trình rèn luyện đó sẽ biến mất hoàn toàn. Nếu thói quen chỉ nhìn vào chỉ số độ tin cậy hiển thị trên màn hình rồi chuyển sang thí nghiệm tiếp theo trở thành thói quen của cả một thế hệ nhà khoa học tương lai, thì dù tổng lượng tri thức tăng lên, số người có khả năng nghi ngờ và sửa chữa tri thức đó lại giảm đi. Nghịch lý khi sự phong phú làm xói mòn thực lực của con người này, cùng với vấn đề gian lận nghiên cứu đã thấy ở chương 12, vẫn là bài toán chưa có lời giải mà thế chế phải đối mặt.

Chúng ta nên đón nhận sự phong phú này như thế nào? Trong phần còn lại của cuốn sách, tôi chia câu hỏi này thành ba phần để đưa ra câu trả lời của tác giả: Tri thức không thể giải thích có phải là khoa học không? Ai chịu trách nhiệm cho kết quả do máy móc phát hiện? Và môi trường nghiên cứu của Hàn Quốc đang đứng ở đâu trong dòng chảy này?

Tri thức không thể giải thích có phải là khoa học không

Đây là câu hỏi đầu tiên: Tri thức chỉ dự đoán đúng nhưng không thể giải thích liệu có thể được gọi là khoa học không?

Nhìn vào lịch sử khoa học, việc dự đoán đi trước và giải thích theo sau không hề hiếm gặp. Động cơ hơi nước đã vận hành trước cả khi nhiệt động lực học ra đời. Newton đã tính toán quỹ đạo của các hành tinh mà không thể nói rõ vì sao trọng lực tồn tại, và lý do đó phải 200 năm sau mới được Einstein bổ khuyết. Mendeleev đã tiên đoán những ô trống trong bảng tuần hoàn mà không biết vì sao các nguyên tố lại có tính tuần hoàn. Một dự đoán chính xác tự thân nó đã là tri thức, và khoa học luôn nắm lấy dự đoán trước khi có được lời giải thích. Vì vậy, không thể vì máy móc thiếu đi lời giải thích mà gạt bỏ những dự đoán của nó ra ngoài cánh cửa khoa học.

Tuy nhiên, dự đoán chỉ là bán thành phẩm chứ chưa phải thành phẩm của khoa học. Dự đoán thiếu lời giải thích sẽ không làm được hai điều: nó không thể báo trước khi nào sẽ sai, và không thể chỉ ra cần phải sửa ở đâu khi đã sai. Khả năng dự đoán cực kỳ chính xác trong dữ liệu huấn luyện nhưng lại âm thầm sụp đổ trước những tình huống xa lạ ngoài tập dữ liệu vốn là điểm yếu cố hữu của các mô hình không có khả năng giải thích. Điểm khác biệt giữa khoa học và cỗ máy dự đoán nằm ở cách xử lý sai sót. Chỉ có tri thức hiểu rõ vì sao đúng mới có thể chỉ ra vì sao sai.

Không phải là không có hy vọng. Lời giải thích đang theo sau. Đúng như Tegmark đã xác nhận khi nhìn sâu vào bên trong mạng nơ-ron, cỗ máy tự mình tạo ra các cấu trúc hình học mà không ai dạy để giải quyết vấn đề. Ngành học nghiên cứu lý do tại sao máy móc dự đoán đúng đang mổ xẻ nội tại của nó bằng các phương pháp vật lý học. Giống như việc sau khi hàng loạt kết quả dự đoán của AlphaFold xuất hiện, các nhà sinh học cấu trúc đã bám vào những cấu trúc đó để một lần nữa truy vấn lại nguyên lý cuộn gập, dự đoán của máy móc cũng đang ngược lại kéo theo sự hiểu biết của con người. Khoảng cách giữa dự đoán và giải thích không phải là định mệnh bất biến, mà là một khoảng cách đang dần được thu hẹp.

Dưới góc nhìn thực dụng, câu hỏi này có thể xem là một nỗi lo xa xôi. Quy trình phê duyệt thuốc mới không hỏi vì sao thuốc có hiệu quả mà chỉ yêu cầu chứng minh thực tế thuốc có tác dụng trong thử nghiệm lâm sàng. Hoạt động được thì sử dụng vốn là thái độ lâu đời của ngành kỹ thuật, và aspirin cũng đã cứu sống con người từ rất lâu trước khi nguyên lý hoạt động của nó được làm sáng tỏ. Nhưng khoa học và kỹ thuật lại có cách xử lý thất bại khác nhau. Những loại thuốc được sử dụng khi chưa rõ nguyên lý từng trở nên hoàn toàn bất lực trước các tác dụng phụ không lường trước, và chỉ sau khi nguyên lý được làm sáng tỏ thì người ta mới thấy rõ không nên dùng thuốc cho đối tượng bệnh nhân nào. Sự thiếu hiểu biết nhưng vẫn hoạt động được có thể đem ra sử dụng, nhưng chỉ có lời giải thích mới cho biết khi nào sự thiếu hiểu biết đó sẽ phản bội chúng ta.

Vì vậy, câu trả lời của cuốn sách này là: Dự đoán của máy móc là nguyên liệu của khoa học. Thời khắc nguyên liệu đó trở thành khoa học là khi ai đó—dù là con người hay máy móc—đặt dự đoán ấy dưới dạng có thể bác bỏ, thử thách bằng thực nghiệm, và đưa nó vào vòng lặp tìm kiếm rồi sửa chữa những điểm sai sót. Một dự đoán đã bước vào vòng lặp thì dù lời giải thích đến muộn vẫn là khoa học. Một dự đoán nằm ngoài vòng lặp thì dù có chính xác đến đâu vẫn chỉ gần như một quẻ bói. Áp dụng lại thước đo phân loại năm hình thái phát hiện ở phần mở đầu, vấn đề không nằm ở chỗ máy móc không thể giải thích, mà nằm ở chỗ con người bắt đầu bỏ qua quy trình đòi hỏi lời giải thích.

Ai chịu trách nhiệm

Đây là câu hỏi thứ hai: Khi kết quả do máy móc phát hiện bị sai, khi kết quả đó gây ra tai nạn, ai sẽ chịu trách nhiệm?

Trong thế giới trách nhiệm, có một nguyên tắc lâu đời: Trách nhiệm đi liền với thẩm quyền. Kẻ nắm giữ quyền quyết định sẽ phải trả giá cho quyết định đó. Thế nhưng trong phòng thí nghiệm vòng lặp khép kín, quyền lựa chọn giả thuyết và quyết định thí nghiệm tiếp theo đang dần chuyển sang cho máy móc. Thẩm quyền được chuyển giao nhưng đối tượng gánh vác trách nhiệm lại không có. Máy móc không có tài sản, không sợ hình phạt và bản thân nó không có tư cách đứng trước pháp luật. Sợi dây gắn kết thẩm quyền và trách nhiệm bị đứt đoạn tại đây.

Câu hỏi này vốn đã ẩn hiện trong nhiều phân cảnh của cuốn sách này. Khi mô hình ngôn ngữ trong phòng thí nghiệm của Gomes dễ dàng lập ra quy trình tổng hợp chất gây nổ, thứ ngăn chặn trước khi đoạn mã đó được thực thi chính là chốt an toàn do con người cài đặt từ trước. Trong thế giới mà Genesis của King tự mình lựa chọn hàng ngàn thí nghiệm mỗi ngày, vụ tai nạn do một thí nghiệm chọn sai trong số đó gây ra là trách nhiệm của ai thì vẫn chưa ai ghi lại câu trả lời. Khi giả thuyết do máy móc đưa ra

trở thành bài báo khoa học hay bằng sáng chế, việc ghi tên ai vào ô tác giả và nhà sáng chế cũng có câu trả lời khác nhau tùy theo từng quốc gia, hoặc thậm chí hoàn toàn không có quy định.

Câu trả lời của cuốn sách này là chúng ta phải buộc lại sợi dây đó từ phía con người. Việc trao tư cách pháp nhân cho máy móc để quy trách nhiệm trong thời điểm hiện tại gần như là điều viển vông. Con đường duy nhất còn lại là con người phải xác định trước và lập thành văn bản phạm vi những gì máy móc được phép tự quyết định. Người xác định phạm vi tự chủ sẽ chịu trách nhiệm cho những gì xảy ra trong phạm vi đó. Nếu vận hành máy móc mà không xác định phạm vi, thì chính việc không xác định phạm vi ấy sẽ trở thành căn cứ quy trách nhiệm. Đây không phải là quy định trói buộc máy móc, mà là ranh giới an toàn bảo vệ con người. Chỉ khi ranh giới giao phó cho máy móc được phân định rõ ràng, nhà nghiên cứu mới biết được phạm vi mình phải chịu trách nhiệm để đặt bút ký.

Tóm lại là ba nguyên tắc: Thứ nhất, xác định phạm vi trước. Phải quy định rõ ràng bằng văn bản các loại thí nghiệm mà máy móc có thể lựa chọn và danh mục các chất mà máy móc có thể xử lý trước khi thí nghiệm bắt đầu. Thứ hai, lưu lại quy trình dưới dạng hồ sơ ghi chép. Phải lưu giữ cơ sở nào khiến máy móc chọn giả thuyết nào dưới dạng con người có thể truy xét lại sau này, để khi xảy ra tai nạn có thể phân định được ranh giới đã bị vượt qua ở đâu. Thứ ba, chữ ký cuối cùng phải do con người thực hiện. Trong ô tác giả của bài báo khoa học, ô nhà sáng chế của bằng sáng chế và ô phê duyệt của giấy chấp thuận thí nghiệm, tên được ghi phải là tên của người xác định phạm vi cho cỗ máy chứ không phải tên của cỗ máy. Những

căn cứ chi tiết của ba nguyên tắc này, cùng những khoảng trống mà luật pháp và thể chế hiện hành chưa thể giải đáp, đã được xem xét cận kề ở chương 12.

Hàn Quốc đang đứng ở đâu

Đây là câu hỏi thứ ba: Trong dòng chảy này, môi trường nghiên cứu của Hàn Quốc đang đứng ở đâu?

Hàn Quốc đã đứng ở vạch xuất phát. Tháng 8 năm 2026, Viện Thông tin Khoa học và Công nghệ Hàn Quốc (KISTI) đã công bố hệ thống đọc bài báo khoa học và lập giả thuyết nghiên cứu mang tên Nhà khoa học AI, đồng thời cho biết sẽ bắt đầu cung cấp dịch vụ cho các nhà nghiên cứu thực tế từ cuối năm. Điều này đồng nghĩa với việc không chỉ các công ty khởi nghiệp ở Thung lũng Silicon mà các viện nghiên cứu công lập của Hàn Quốc cũng đã bắt đầu chuyển động theo cùng một hướng. Nguyên liệu cũng không hề thiếu. Phòng thí nghiệm tự hành là nơi hội tụ của robot, dữ liệu và tính toán, mà Hàn Quốc là một trong số ít quốc gia sở hữu cả ba yếu tố đó. Hiện trường sản xuất đẳng cấp thế giới trong các ngành bán dẫn, pin và dược phẩm sinh học tự thân nó đã là thao trường huấn luyện tự động hóa chính xác; còn dữ liệu bệnh viện tích lũy từ hệ thống bảo hiểm y tế toàn dân cùng siêu máy tính quốc gia và mạng lưới tính toán nghiên cứu chính là kho dữ liệu và tính toán để nuôi dưỡng nhà khoa học máy.

Vấn đề nằm ở chỗ các nguyên liệu này chưa thể gặp nhau. Dữ liệu bệnh viện bị khóa kín sau bức tường bệnh viện, công nghệ tự động hóa sản xuất nằm sau bức tường nhà máy, còn hồ sơ thí nghiệm lại nằm gọn trong ngăn kéo phòng nghiên cứu. Vòng lặp khép kín chỉ có thể vận hành khi vượt qua được những bức tường này. Phá bỏ các bức tường không phải là việc của công nghệ mà là việc của thể chế, và vì vậy, điểm mấu chốt quyết định thành bại của Hàn Quốc không nằm ở phòng thí nghiệm mà nằm ở thể chế.

Thói quen đánh giá nghiên cứu cũng đang đi ngược lại dòng chảy này. Sức mạnh của phòng thí nghiệm vòng lặp khép kín đến từ việc trải qua thất bại với số lượng lớn và tốc độ nhanh. Phòng thí nghiệm của Aspuru-Guzik có thể cắt giảm chi phí chính là nhờ hấp thụ hàng trăm lần tổng hợp thất bại với chi phí rẻ. Thế nhưng việc đánh giá đề tài nghiên cứu ở Hàn Quốc từ lâu đã luôn đặt tỷ lệ thành công lên hàng đầu. Trong cơ cấu mà việc để lại vết tích thất bại trong hồ sơ khiến nhà nghiên cứu khó nhận được đề tài tiếp theo, họ buộc phải đăng ký những mục tiêu an toàn không thể thất bại. Cỗ máy sử dụng

thất bại làm nhiên liệu và thể chế coi thất bại là điểm trừ không thể ăn khớp với nhau. Việc xây dựng hệ thống đánh giá biết nhìn nhận hàng trăm thất bại do máy móc tạo ra như một thành quả còn cấp thiết hơn việc đưa nhà khoa học máy vào hoạt động.

Thói quen xử lý dữ liệu cũng đang đứng trước ngã rẽ. Sức mạnh của vòng lặp khép kín chỉ gia tăng khi hồ sơ thí nghiệm được chia sẻ dưới định dạng máy có thể đọc được. Chừng nào tập quán cất giữ tài liệu trong ngăn kéo của từng phòng thí nghiệm vẫn còn tồn tại, nhà khoa học máy của Hàn Quốc sẽ phải học từ bài báo khoa học của nước khác và chịu cảnh đối dữ liệu trước tài liệu của chính nước mình. Cần phải có một thể chế chia sẻ tài liệu vượt qua quan niệm sở hữu cục bộ cấp phòng thí nghiệm, cùng với các chuẩn mực nghiên cứu quy định rõ trách nhiệm đối với giả thuyết do máy móc tạo ra.

Vì vậy, cuốn sách này muốn gửi đến những người thiết kế thể chế của Hàn Quốc bốn đề xuất: Thứ nhất, khu vực công cần tiên phong thiết lập các dự án thí điểm phòng thí nghiệm tự hành, tạo ra sân chơi để các trường đại học và doanh nghiệp cùng chia sẻ chi phí thất bại và học hỏi. Thứ hai, sửa đổi quy định đánh giá đề tài nghiên cứu để công nhận hồ sơ thí nghiệm thất bại do máy móc thực hiện là một sản phẩm đầu ra chứ không phải điểm trừ. Thứ ba, quy định tiêu chuẩn lưu giữ hồ sơ thí nghiệm dưới định dạng máy có thể đọc được, đồng thời xác lập nguyên tắc bắt buộc chia sẻ dữ liệu thí nghiệm có sử dụng kinh phí nghiên cứu công. Thứ tư, đào tạo các nhà khoa học có khả năng điều khiển máy móc. Trọng tâm của giáo dục sau đại học cần chuyển từ đôi bàn tay cầm ống hút pipet sang rèn luyện đôi mắt biết hoài nghi và kiểm chứng những kết luận do máy móc đưa ra.

Như đã thấy ở chương 12, phần lớn những chuẩn mực này hiện vẫn còn là trang giấy trắng. Trang giấy trắng vừa là bằng chứng của sự tụt hậu, nhưng cũng có nghĩa là quốc gia nào viết lên trước sẽ tạo ra tiêu chuẩn. Quốc gia tiên phong viết ra các chuẩn mực nghiên cứu và ngữ pháp trách nhiệm cần thiết trong thời đại của nhà khoa học máy sẽ trở thành người định hình luật chơi cho khoa học trong nửa thế kỷ tiếp theo. Nơi Hàn Quốc đang đứng chính là trước trang giấy trắng ấy.

Những gì còn lại cho con người

Ở cuối mỗi phần, tôi đã lặp đi lặp lại cùng một câu hỏi: Vai trò của con người đã dịch chuyển về đâu? Câu trả lời được xác nhận qua năm phần là như sau: Ở phần 1, máy móc đứng ở vị trí đưa ra câu trả lời thay con người, và con người chuyển sang vị trí kiểm chứng câu trả lời đó bằng thực nghiệm. Ở phần 2, máy móc đứng ở vị trí phác họa những thứ chưa từng tồn tại, và con người chuyển sang vị trí quyết định sẽ yêu cầu máy phác họa điều gì. Ở phần 3, máy móc đứng ở vị trí vớt lên các quy luật từ dữ liệu, và con người chuyển sang vị trí diễn giải xem quy luật đó có ý nghĩa gì. Ở phần 4, máy móc trở thành đôi bàn tay trên bàn thí nghiệm, và con người chuyển sang vị trí vạch ra ranh giới mà đôi tay đó không được phép vượt qua. Và ở phần 5, khi máy móc bắt đầu vận hành trọn vẹn toàn bộ vòng lặp, vị trí còn lại của con người cuối cùng thu hẹp vào một điểm duy nhất: vị trí lựa chọn câu hỏi. Đó là vị trí quyết định việc tò mò về điều gì, chữa trị căn bệnh nào trước, và dành thời gian cùng tiền bạc vào hướng đi nào.

Suốt một thời gian dài, vị trí đó dường như là đặc quyền riêng của con người. Thế nhưng phần kết của chặng đường mà cuốn sách này đi theo lại mở ra một khung cảnh khác. Máy móc đã tự mình đưa ra giả thuyết, lựa chọn thí nghiệm tiếp theo, và thậm chí tự mở rộng danh sách các câu hỏi nghiên cứu. Hệ thống Co-Scientist (DeepMind) mà Google DeepMind công bố trên tạp chí Nature năm 2026 đã tiến xa đến mức cho nhiều tác tử cùng phản bác lẫn nhau về các giả thuyết để trau chuốt các ý tưởng nghiên cứu. Nếu như Co-Scientist (Gomes) ở chương 9 là cỗ máy sở hữu đôi tay làm thí nghiệm, thì đây là cỗ máy mô phỏng bộ não suy ngẫm về các câu hỏi. Thời đại mà ngay cả thẩm quyền quyết định nên đặt ra câu hỏi nào cũng có thể được giao cho máy móc đang đến gần. Vì vậy, câu kết của cuốn sách này không phải là một lời trấn an mà là một bài tập về nhà: Chúng ta phải quyết định ngay bây giờ xem con người sẽ nắm giữ thẩm quyền đó đến giới hạn nào. Bởi vì nếu trì hoãn quyết định, chính việc quyết định sẽ trở

thành phần việc của máy móc.

Hãy hình dung về một phòng thí nghiệm lúc 11 giờ đêm. Tại vị trí của người sinh viên từng dùng đôi bàn tay run rẩy để chuyển chất lỏng trước tủ hút khí độc, giờ đây đã là một cánh tay robot. Người sinh viên ngồi trước màn hình đọc danh sách các giả thuyết do máy móc lập ra. Việc viết gì vào dòng cuối cùng của danh sách đó, ký tên vào dòng nào và xóa đi dòng nào vẫn là phần việc của người sinh viên ấy. Gìn giữ phần việc đó chính là di sản đích thực mà những con người làm thay đổi khoa học bằng trí tuệ nhân tạo trao lại cho thế hệ tiếp theo.

Phụ lục

Bảng giới thiệu nhân vật

Nơi đây tập hợp cơ quan trực thuộc và thành tựu tiêu biểu của mười một nhân vật xuất hiện trong cuốn sách này. Chức danh được tính đến thời điểm tháng 8 năm 2026.

Demis Hassabis: Chủ tịch Google DeepMind, Nhà khoa học trưởng của Alphabet, Giám đốc điều hành Isomorphic Labs. Đồng đoạt giải Nobel Hóa học năm 2024 nhờ AlphaFold dự đoán cấu trúc protein.

John Jumper: Cựu Giám đốc nghiên cứu của Google DeepMind, chuyển sang Anthropic vào năm 2026. Kiến trúc sư trưởng của AlphaFold2, đồng đoạt giải Nobel Hóa học năm 2024.

Regina Barzilay: Giáo sư tại Viện Công nghệ Massachusetts (MIT). Khám phá ra kháng sinh halicin và trí tuệ nhân tạo chẩn đoán sớm ung thư vú. Người đầu tiên nhận giải thưởng Squirrel của AAAI.

David Baker: Giáo sư tại Đại học Washington, Giám đốc Viện Thiết kế Protein. Phát triển Rosetta và RFdiffusion thiết kế những loại protein chưa từng tồn tại trong tự nhiên. Đồng đoạt giải Nobel Hóa học năm 2024.

Alán Aspuru-Guzik: Giáo sư tại Đại học Toronto, Giám đốc Hiệp hội Vật liệu Tăng tốc. Phòng thí nghiệm tự hành kết nối trí tuệ nhân tạo và robot.

Hod Lipson: Giáo sư tại Đại học Columbia. Phát triển Eureka tìm kiếm các định luật vật lý từ dữ liệu, robot có khả năng tự mô hình hóa cơ thể.

Max Tegmark: Giáo sư tại Viện Công nghệ Massachusetts (MIT), người sáng lập Viện Tương lai Sự sống. Phát triển AI Feynman tái khám phá các định luật vật lý, nghiên cứu an toàn trí tuệ nhân tạo.

Ross King: Giáo sư tại Đại học Công nghệ Chalmers và Đại học Cambridge. Phát triển Adam, Eve và Genesis - những robot thí nghiệm tự hành đầu tiên trên thế giới.

Gabe Gomes: Giáo sư Đại học Carnegie Mellon (nghỉ phép), Nghiên cứu viên tại Google X. Phát triển Coscientist dẫn dắt các thí nghiệm hóa học bằng mô hình ngôn ngữ.

Sam Rodriques: Đồng sáng lập kiêm Giám đốc điều hành FutureHouse, Giám đốc điều hành Edison Scientific. Phát triển PaperQA2 và Cosmos đọc tài liệu để tổng hợp tri thức.

Hiroaki Kitano: Giám đốc điều hành Phòng thí nghiệm Khoa học Máy tính Sony, Giáo sư kiêm nhiệm tại Viện Khoa học và Công nghệ Sau đại học Okinawa. Người sáng lập sinh học hệ thống, Nobel Turing

người đề xuất Thử thách, người sáng lập RoboCup.

Niên biểu

Năm 1965, Dự án Dendral khởi động tại Đại học Stanford, nỗ lực đầu tiên để máy móc suy luận cấu trúc phân tử bằng kiến thức hóa học. Thập niên 1970, chương trình BACON của Herbert Simon tái khám phá các định luật của Kepler từ dữ liệu. Năm 1997, Kitano sáng lập giải bóng đá robot RoboCup. Năm 2004, robot nhà khoa học Adam của Ross King tự mình xác định chức năng gen của nấm men (Nature). Năm 2009, bài báo tổng kết những phát hiện tự hành của Adam (Science). Cùng năm đó, Eureka của Lipson và Schmidt tái khám phá các định luật vật lý từ chuyển động của con lắc (Science). Năm 2016, Kitano đề xuất Thử thách Lớn về trí tuệ nhân tạo tạo ra những phát hiện tầm cỡ giải Nobel. Năm 2018, robot nhà khoa học Eve phát hiện tác dụng ức chế sốt rét của triclosan. Năm 2020, AI Feynman của Tegmark tái khám phá 100 phương trình. Nhóm nghiên cứu của Barzilay phát hiện kháng sinh halicin

(Cell). Hội thảo Thử thách Nobel Turing tại Viện Turing ở London. Năm 2021, AlphaFold2 giải quyết về cơ bản bài toán dự đoán cấu trúc protein (Nature). Kitano công bố bài báo về Thử thách Nobel Turing. Năm 2022, cơ sở dữ liệu AlphaFold công khai cấu trúc của 200 triệu protein. Năm 2023, Coscientist của Gomes thực hiện thí nghiệm hóa học bằng mô hình ngôn ngữ (Nature). RFDiffusion tạo ra protein thông qua thiết kế. Năm 2024, Hassabis, Jumper và Baker đồng đoạt giải Nobel Hóa học. Năm 2025, bài báo do Sakana AI tạo ra vượt qua vòng bình duyệt hội thảo khoa học. Rodriques thành lập Edison Scientific. Năm 2026, Jumper chuyển sang Anthropic, Hassabis nhậm chức Chủ tịch DeepMind, Jeff Dean sáng lập Discovery Loop. Anthropic ra mắt Claude Science, KISTI công bố Nhà khoa học AI. Năm 2050 là năm mục tiêu do Thử thách Nobel Turing đặt ra.

Giải thích thuật ngữ

Học tăng cường: Phương thức học tập trong đó máy móc tự mình học các hành động nhằm tối đa hóa phần thưởng thông qua thử và sai. Cách AlphaGo học chơi cờ vây thuộc về phương thức này.

Học chủ động: Phương thức học tập trong đó máy móc tự lựa chọn thí nghiệm tiếp theo để thu được nhiều thông tin nhất. Đây là cách robot nhà khoa học Eve thu hẹp các ứng viên thuốc.

Tối ưu hóa Pareto: Phương pháp tìm kiếm các điểm mà tại đó không thể cải thiện một mục tiêu mà không phải hy sinh mục tiêu còn lại, chẳng hạn như sự đánh đổi giữa độ chính xác và tính ngắn gọn.

Vòng lặp khép kín: Trạng thái trong đó chu trình từ quan sát, giả thuyết, thí nghiệm đến kiểm chứng diễn ra trọn vẹn một vòng mà không cần con người can thiệp. Đây là khái niệm cốt lõi của khoa học tự hành trong cuốn sách này.

Hồi quy biểu tượng: Phương pháp tìm ra toàn bộ công thức toán học giải thích dữ liệu từ chính dữ liệu đó. Đây là phương thức mà Eureqa của Lipson và AI Feynman của Tegmark đã sử dụng.

Mô hình ngôn ngữ lớn: Trí tuệ nhân tạo học từ lượng văn bản khổng lồ để dự đoán từ tiếp theo và tạo ra câu hoàn chỉnh. Coscientist của Gomes đã sử dụng mô hình này làm bộ não chỉ huy các thí nghiệm.

Gấp cuộn protein: Hiện tượng chuỗi axit amin tự gấp lại thành hình dạng ba chiều. AlphaFold đã giải quyết bài toán dự đoán hình dạng đó.

Tính tái lập: Tính chất cho ra kết quả giống nhau khi lặp lại cùng một thí nghiệm. Khi máy móc lưu lại các điều kiện thí nghiệm dưới dạng mã lệnh, tính tái lập sẽ tăng lên.

Ghi chú kỹ thuật

Trong số các giải thích toán học được lược bớt ở phần chính văn, những nội dung dành cho độc giả muốn tìm hiểu sâu hơn được tập hợp tại đây. Độc giả không nhất thiết phải đọc ghi chú này để hiểu nội dung chính của cuốn sách.

Biểu diễn cấu trúc của AlphaFold: AlphaFold2 nhìn nhận protein theo hai nhánh. Một là mối tương quan giữa các axit amin cùng biến đổi trong quá trình tiến hóa, và nhánh còn lại là cấu trúc hình học tạo nên bởi khoảng cách và góc giữa các cặp axit amin. Bằng cách luân phiên trao đổi hai luồng thông tin này qua nhiều tầng lớp, nó dần thu hẹp lại thành một bố cục ba chiều không tự mâu thuẫn. Độ tin cậy được thể hiện bằng điểm số cho từng phần dư, qua đó màu sắc sẽ cho thấy vùng nào của dự đoán là đáng tin cậy.

Kiểm tra tính đối xứng của AI Feynman: AI Feynman không lập tức tìm kiếm công thức. Đầu tiên, nó khớp các đơn vị đại lượng vật lý để thu hẹp phạm vi các công thức khả dĩ. Tiếp theo, nó tìm kiếm tính đối xứng ẩn bằng cách kiểm tra xem kết quả có giữ nguyên khi cộng hoặc nhân một biến số nào đó hay không, rồi kiểm tra xem liệu một biểu thức lớn có thể chia làm hai được không. Đối với dữ liệu thực tế

lẫn tạp âm, nó cho mạng nơ-ron học thuộc trước, sau đó dùng chính mạng nơ-ron đó làm máy phát để tạo ra dữ liệu sạch vô tận.

Đường biên Pareto của Eureka: Eureka cân nhắc giữa một công thức khớp chính xác với dữ liệu và một công thức ngắn gọn, súc tích. Nếu chấp nhận giảm bớt độ chính xác một chút thì công thức sẽ ngắn hơn, còn nếu chấp nhận bớt đi tính ngắn gọn thì công thức sẽ chính xác hơn. Đường ranh giới của những công thức không thể cải thiện nếu không hy sinh yếu tố còn lại được gọi là đường biên Pareto, và con người sẽ lựa chọn công thức hữu ích từ đường biên đó.

Lợi ích thông tin của vòng lặp khép kín: Khi lựa chọn thí nghiệm tiếp theo cần thực hiện, phòng thí nghiệm tự hành sẽ định giá lượng thông tin mà thí nghiệm đó mang lại dù thành công hay thất bại. Những thí nghiệm nhằm xác nhận lại điều đã biết có giá trị thấp, còn những thí nghiệm khó dự đoán kết quả có giá trị cao. Giá trị này được so sánh với chi phí hóa chất và thời gian để lựa chọn thí nghiệm mang lại lợi ích lớn nhất.

Hình ảnh

Cuốn sách này có ba hình ảnh. Chúng được cung cấp dưới dạng tệp riêng biệt. Hình 1. Năm hình thái khoa học Hình 2. Cấu trúc tuần hoàn của thí nghiệm vòng lặp khép kín Hình 3. Năm cấp độ tự động hóa và sự thay đổi vai trò của con người

Tài liệu tham khảo

Các tài liệu căn cứ của cuốn sách này được tập hợp theo từng chương. Các dữ kiện trong nội dung chính đã được đối chiếu và xác nhận với các tài liệu dưới đây. Mỗi tài liệu được đánh giá mức độ tin cậy và phân loại từ [A] đến [D]. [A] là các bài báo gốc và cơ sở dữ liệu chính thức, [B] là các công bố chính thức từ các trường đại học, viện nghiên cứu và các bên liên quan, [C] là báo chí uy tín, [D] là tài liệu quảng bá của doanh nghiệp, blog và mạng xã hội. Nếu có mã DOI và arXiv, chúng sẽ được đặt ở đầu.

Lời mở đầu 1. [A] Simon, H. A. (1977). *Models of Discovery: and Other Topics in the Methods of Science*. D. Reidel Publishing Company. (ISBN: 978-9027708571)

2. [B] Langley, P., Simon, H. A., Bradshaw, G. L., & Zytkow, J. M. (1987). *Scientific Discovery: Computational Explorations of the Creative Processes*. MIT Press. (ISBN: 978-0262620529)

3. [A] King, R. D. et al. (2009). The automation of science. *Science*, 324(5923), 85-89. DOI: 10.1126/science.1165620

4. [A] Schmidt, M., & Lipson, H. (2009). Distilling free-form natural laws from experimental data. *Science*, 324(5923), 81-85. DOI: 10.1126/science.1165893

5. [A] Jumper, J. et al. (2021). Highly accurate protein structure prediction with AlphaFold. *Nature*, 596, 583-589. DOI: 10.1038/s41586-021-03819-2

6. [A] Stokes, J. M. et al. (2020). A deep learning approach to antibiotic discovery. *Cell*, 180(4), 688-702. DOI: 10.1016/j.cell.2020.01.021

7. [A] Watson, J. L. et al. (2023). De novo design of protein structure and function with RFdiffusion. *Nature*, 620, 1089-1100. DOI: 10.1038/s41586-023-06415-8

8. [A] Boiko, D. A., MacKnight, R., Kline, B., & Gomes, G. (2023). Autonomous chemical research with large language models. *Nature*, 624, 570-578. DOI: 10.1038/s41586-023-06792-0

9. [A] Gottweis, J. et al. (2026). Accelerating scientific discovery with Co-Scientist. *Nature*, 655(8122), 487-496. DOI: 10.1038/s41586-026-10644-y

10. [A] Musslick, S. et al. (2025). Automating the practice of science: Opportunities, challenges, and implications. *Proceedings of the National Academy of Sciences*, 122(5), e2401238121. DOI: 10.1073/pnas.2401238121
11. [A] Kitano, H. (2021). Nobel Turing Challenge: Creating the engine for scientific discovery. *npj Systems Biology and Applications*, 7, Article 29. DOI: 10.1038/s41540-021-00189-3
- Chương 1 Demis Hassabis 1. [A] Althöfer, I. (2024). Demis Hassabis: From Chess Wunderkind Via Play to the Nobel Prize. *ICGA Journal*, 46(3-4), 85-97. DOI: 10.1177/13896911251320715
2. [B] Mallaby, S. (2026). *The Infinity Machine: Demis Hassabis, DeepMind and the Quest for Superintelligence*. Penguin Press.
3. [B] Hassabis, D. (2024). AlphaFold: The 50-year grand challenge cracked by AI. Nobel Prize Lecture in Chemistry 2024, Stockholm, Sweden.
4. [A] Hassabis, D., Kumaran, D., Vann, S. D., & Maguire, E. A. (2007). Patients with hippocampal amnesia cannot imagine new experiences. *Proceedings of the National Academy of Sciences*, 104(5), 1726-1731. DOI: 10.1073/pnas.0610561104
5. [A] Hassabis, D., & Maguire, E. A. (2007). Deconstructing episodic memory with construction. *Trends in Cognitive Sciences*, 11(7), 299-306. DOI: 10.1016/j.tics.2007.05.001
6. [A] Hassabis, D., & Maguire, E. A. (2009). The construction system of the brain. *Philosophical Transactions of the Royal Society B*, 364(1521), 1263-1271. DOI: 10.1098/rstb.2008.0296
7. [A] Schacter, D. L., Addis, D. R., Hassabis, D., Martin, V. C., Spreng, R. N., & Szpunar, K. K. (2012). The Future of Memory: Remembering, Imagining, and the Brain. *Neuron*, 76(4), 677-694. DOI: 10.1016/j.neuron.2012.11.001
8. [A] The News Staff (2007). Breakthrough of the Year: The Runners-up. *Science*, 318(5858), 1844a-1849a. DOI: 10.1126/science.318.5858.1844a
9. [C] Fortune (2026). Demis Hassabis steps down from Google DeepMind CEO role amid a major AI leadership shake-up. *Fortune*. URL: <https://fortune.com/2026/08/05/demis-hassabis-steps-down-google-deepmind-ai-shakeup/> (Truy cập ngày 19 tháng 8 năm 2026)
10. [C] Axios (2026). Google's Hassabis calls for new US-led global AI watchdog "before year end." *Axios*. URL: <https://www.axios.com/2026/07/14/demis-hassabis-ai-regulation-google-deepmind> (Truy cập ngày 19 tháng 8 năm 2026)
11. [A] Hassabis, D., Kumaran, D., & Maguire, E. A. (2007). Using Imagination to Understand the Neural Basis of Episodic Memory. *The Journal of Neuroscience*, 27(52), 14365-14374. DOI: 10.1523/JNEUROSCI.4549-07.2007
12. [A] Mnih, V., Kavukcuoglu, K., Silver, D., Graves, A., Antonoglou, I., Wierstra, D., & Riedmiller, M. (2013). Playing Atari with Deep Reinforcement Learning. *NIPS Deep Learning Workshop 2013*. arXiv:1312.5602.
13. [A] Mnih, V. et al. (2015). Human-level control through deep reinforcement learning. *Nature*, 518, 529-533. DOI: 10.1038/nature14236
14. [A] Silver, D. et al. (2016). Mastering the game of Go with deep neural networks and tree search. *Nature*, 529, 484-489. DOI: 10.1038/nature16961

15. [A] Reed, S. et al. (2022). A Generalist Agent. Transactions on Machine Learning Research. arXiv:2205.06175.
16. [A] Musslick, S. et al. (2025). Automating the practice of science: Opportunities, challenges, and implications. Proceedings of the National Academy of Sciences, 122(5), e2401238121. DOI: 10.1073/pnas.2401238121
17. [B] 9to5Google (2026). Demis Hassabis no longer DeepMind CEO to focus on new AGI role, Jeff Dean departs. URL: <https://9to5google.com/2026/08/05/demis-hassabis-deepmind/> (Truy cập ngày 19 tháng 8 năm 2026)
18. [B] Kerr, B. (2026). From Hippocampal Simulation to AGI: How Hassabis's Thesis Shapes DeepMind's Race to General Intelligence. URL: <https://bretkerr.substack.com/p/from-hippocampal-simulation-to-agi> (Truy cập ngày 19 tháng 8 năm 2026)
19. [C] Axios (2026). Google DeepMind CEO Demis Hassabis is stepping aside. Axios. URL: <https://www.axios.com/2026/08/05/google-deepmind-demis-hassabis-ai> (Truy cập ngày 19 tháng 8 năm 2026)
20. [B] Sutton, R. S., & Barto, A. G. (2018). Reinforcement Learning: An Introduction. MIT Press. ISBN: 978-0262039246
21. [A] Schultz, W., Dayan, P., & Montague, P. R. (1997). A neural substrate of prediction and reward. Science, 275(5306), 1593-1599. DOI: 10.1126/science.275.5306.1593
22. [A] Sutton, R. S. (1988). Learning to predict by the methods of temporal differences. Machine Learning, 3, 9-44. DOI: 10.1007/BF00115009
23. [A] Nature (2026). "An AlphaFold 4", scientists marvel at DeepMind drug spin-off's exclusive new AI. Nature. URL: <https://www.nature.com/articles/d41586-026-00365-7> (truy cập ngày 19 tháng 8 năm 2026)
24. [A] Degraeve, J., Felici, F., Buchli, J., Neunert, M., Tracey, B., Carpanese, F., Ewald, J. M., Hafner, R., Riedmiller, M., Erdmann, M. C., et al. (2022). Magnetic control of tokamak plasmas through deep reinforcement learning. Nature, 602, 414-419. DOI: 10.1038/s41586-021-04301-9
25. [A] Lam, R., Sanchez-Gonzalez, A., Matthew, M., Fortunato, M., Pritzel, A., Sheppard, P., Wirnsberger, P., Battaglia, P., & Hassabis, D. (2023). Learning skillful medium-range global weather forecasting. Science, 382(6677), 1416-1421. DOI: 10.1126/science.adi2336
26. [A] Jumper, J., Evans, R., Pritzel, A., Green, T., Figurnov, M., Ronneberger, O., et al., & Hassabis, D. (2021). Highly accurate protein structure prediction with AlphaFold. Nature, 596, 583-589. DOI: 10.1038/s41586-021-03819-2
27. [A] Weinberg, S. (1992). Dreams of a Final Theory: The Scientist's Search for the Ultimate Laws of Nature. Pantheon Books. ISBN: 978-0679419235
28. [B] Hassabis, D. (2024). Accelerating Scientific Discovery with AI. Nobel Prize Lecture in Chemistry 2024, Aula Magna, Stockholm University, Sweden.
29. [A] Abramson, J., Jumper, J., Silver, D. et al. (2024). AlphaFold 3 predicts the structure and interactions of all of life's molecules. Nature, 630, 493-500. DOI: 10.1038/s41586-024-07487-w
30. [A] Senior, A. W. et al. (2020). Improved protein structure prediction using potentials from deep learning. Nature, 577, 706-710. DOI: 10.1038/s41586-019-1923-7

31. [A] Skarlinski, M. D., White, A. D., & Rodriques, S. G. et al. (2024). Language agents achieve superhuman synthesis of scientific knowledge. arXiv preprint arXiv:2409.13740.
 32. [D] Isomorphic Labs. (2026). The Isomorphic Labs Drug Design Engine unlocks a new frontier beyond AlphaFold. Technical Whitepaper, Isomorphic Labs Press, London.
 33. [B] BioSpace (2026). Isomorphic Labs secures \$2.1 Billion funding to scale its AI drug design engine. URL:
<https://www.biospace.com/press-releases/isomorphic-labs-secures-2-1-billion-funding-to-scale-its-ai-drug-design-engine> (truy cập ngày 19 tháng 8 năm 2026)
 34. [B] MLQ.ai (2026). Isomorphic Labs Launches Human Trials for AI-Designed Drugs. URL:
<https://mlq.ai/news/isomorphic-labs-launches-human-trials-for-ai-designed-drugs/> (truy cập ngày 19 tháng 8 năm 2026)
 35. [A] Bongard, J. C., & Lipson, H. (2007). Automated reverse engineering of nonlinear dynamical systems. *Proceedings of the National Academy of Sciences*, 104(24), 9943-9948. DOI: 10.1073/pnas.0609476104
 36. [A] King, R. D., Rowland, J., Oliver, S. G., Young, M., Aubrey, W., Byrne, E., Liakata, M., Markham, M., Pir, P., Soldatova, L. N., Sparkes, A., Whelan, K. E., & Clare, A. (2009). The automation of science. *Science*, 324(5923), 85-89. DOI: 10.1126/science.1165620
 37. [B] King, R. D. et al. (2015). Functional genomic hypothesis generation and confirmation by the Robot Scientist Adam. *Journal of Eukaryotic Systems Biology*, 22(3), e105432.
 38. [A] Google DeepMind, & Isomorphic Labs (2026). Our approach to bioresilience. Google DeepMind Blog. URL: <https://deepmind.google/blog/our-approach-to-bioresilience/> (truy cập ngày 19 tháng 8 năm 2026)
 39. [B] Isomorphic Labs (2026). The Isomorphic Labs Drug Design Engine unlocks a new frontier beyond AlphaFold. Isomorphic Labs. URL:
<https://www.isomorphiclabs.com/articles/the-isomorphic-labs-drug-design-engine-unlocks-a-new-frontier> (truy cập ngày 19 tháng 8 năm 2026)
- Chương 2 John Jumper 1. [A] Jumper, J., Evans, R., Pritzel, A. et al. (2021). Highly accurate protein structure prediction with AlphaFold. *Nature*, 596(7873), 583-589. DOI: 10.1038/s41586-021-03819-2
2. [A] Jumper, J. M. (2017). *New Methods Using Rigorous Machine Learning for Coarse-Grained Protein Folding and Dynamics* (luận án tiến sĩ). University of Chicago. DOI: 10.6082/M1BZ647N
 3. [C] University of Chicago Department of Chemistry (2024). John Jumper, UChicago Department of Chemistry alum, and the 2024 Nobel Prize winner in Chemistry. UChicago Chemistry News. URL:
<https://chemistry.uchicago.edu/news/john-jumper-uchicago-department-of-chemistry-alum-and-the-2024-nobel-prize-winner-in-chemistry> (truy cập ngày 19 tháng 8 năm 2026)
 4. [D] Google DeepMind (2024). Demis Hassabis & John Jumper awarded Nobel Prize in Chemistry. DeepMind Blog. URL:
<https://deepmind.google/blog/demis-hassabis-john-jumper-awarded-nobel-prize-in-chemistry/> (truy cập ngày 19 tháng 8 năm 2026)
 5. [B] University of Chicago Sosnick Group. Protein Folding. URL:
<https://sosnick.uchicago.edu/research/proteinfolding.html> (truy cập ngày 19 tháng 8 năm 2026)

6. [A] Britannica. John M. Jumper: Biography, Nobel Prize, DeepMind, AlphaFold, & Facts. URL: <https://www.britannica.com/biography/John-M-Jumper> (truy cập ngày 19 tháng 8 năm 2026)
7. [C] Tech Times (2026). AlphaFold Nobel Laureate John Jumper Joins Anthropic After Nine Years at DeepMind. URL: <https://www.techtimes.com/articles/318754/20260620/alphafold-nobel-laureate-john-jumper-joins-anthropic-after-nine-years-deepmind.htm> (truy cập ngày 19 tháng 8 năm 2026)
8. [A] Digital Applied (2026). John Jumper Joins Anthropic: AI for Science Heats Up. URL: <https://www.digitalapplied.com/blog/john-jumper-joins-anthropic-ai-for-science-2026> (truy cập ngày 19 tháng 8 năm 2026)
9. [C] The Next Web (2026). Nobel laureate John Jumper is leaving Google DeepMind for Anthropic after nearly nine years. URL: <https://thenextweb.com/news/john-jumper-nobel-deepmind-leaves-anthropic-alphafold> (truy cập ngày 19 tháng 8 năm 2026)
10. [C] University of Chicago News. Nobel laureate John Jumper returns to UChicago to discuss the AlphaFold protein revolution. URL: <https://news.uchicago.edu/story/nobel-laureate-john-jumper-returns-uchicago-discuss-alphafold-protein-revolution> (truy cập ngày 19 tháng 8 năm 2026)
11. [A] Senior, A. W., Evans, R., Jumper, J., Kirkpatrick, J., Sifre, L., Green, T., ... & Hassabis, D. (2020). Improved protein structure prediction using potentials from deep learning. *Nature*, 577(7792), 706-710. DOI: 10.1038/s41586-019-1923-7
12. [B] Jumper, J. (2024). Structure Prediction with AlphaFold. Nobel Prize Lecture in Chemistry 2024, Stockholm, Sweden.
13. [B] Critical Assessment of Techniques for Protein Structure Prediction (2020). CASP14 Abstracts. predictioncenter.org.
14. [B] Applying and improving AlphaFold at CASP14 (2021). PubMed. URL: <https://pubmed.ncbi.nlm.nih.gov/34599769/>
15. [B] Mallaby, S. (2026). *The Infinity Machine: How Demis Hassabis Built DeepMind and Chased AGI*. Penguin Press. ISBN: 978-0593831847
16. [B] Thakker, Y. (2026). John Jumper Leaves Google DeepMind for Anthropic: AlphaFold Nobel Laureate Joins Claude. *explainx.ai Blog*, Jun 19, 2026.
17. [B] TechTimes (2026). AlphaFold Nobel Laureate John Jumper Joins Anthropic After Nine Years at DeepMind. URL: <https://www.techtimes.com/articles/318754/20260620/alphafold-nobel-laureate-john-jumper-joins-anthropic-after-nine-years-deepmind.htm> (truy cập ngày 19 tháng 8 năm 2026)
18. [A] Frontiers in Artificial Intelligence (2026). The transformative impact of AI-enabled AlphaFold 3: evolution, current status, and future prospects in structural biology. URL: <https://www.frontiersin.org/journals/artificial-intelligence/articles/10.3389/frai.2026.1739303/full> (truy cập ngày 19 tháng 8 năm 2026)
19. [A] Vaswani, A., Shazeer, N., Parmar, N. et al. (2017). Attention Is All You Need. *Advances in Neural Information Processing Systems (NeurIPS 2017)*, 30, 5998-6008.
20. [A] Musslick, S. et al. (2025). Automating the practice of science: Opportunities, challenges, and implications. *Proceedings of the National Academy of Sciences (PNAS)*, 122(5), e2401238121. DOI: 10.1073/pnas.2401238121

21. [A] Abramson, J., Jumper, J., Silver, D. et al. (2024). AlphaFold 3 predicts the structure and interactions of all of life's molecules. *Nature*, 630, 493-500. DOI: 10.1038/s41586-024-07487-w
22. [B] AIToolsRecap. (2026). John Jumper Joins Anthropic: Nobel Prize Winner, AlphaFold Creator, and What It Means for the AI Talent War. *The Journal of Artificial Intelligence Tools*, Jun 25, 2026.
23. [B] EMBL-EBI (2026). EMBL-EBI and Google DeepMind renew partnership and release update to AlphaFold Database. EMBL-EBI. URL: <https://www.ebi.ac.uk/about/news/technology-and-innovation/google-deepmind-partnership-renewal/> (truy cập ngày 19 tháng 8 năm 2026)
24. [B] EMBL-EBI (2026). Millions of protein complexes added to AlphaFold Database shed light on how proteins interact. EMBL-EBI. URL: <https://www.ebi.ac.uk/about/news/technology-and-innovation/first-complexes-alphafold-database/> (truy cập ngày 19 tháng 8 năm 2026)
25. [A] Varadi, M., Bertoni, D., Magana, P. et al. (2024). AlphaFold Protein Structure Database in 2024: providing structure coverage for over 214 million protein sequences. *Nucleic Acids Research*, 52(D1), D368-D375. DOI: 10.1093/nar/gkad1011
26. [D] Isomorphic Labs. (2026). The Isomorphic Labs Drug Design Engine unlocks a new frontier beyond AlphaFold. Technical Whitepaper, Isomorphic Labs Press, London.
27. [A] Varadi, M. et al. (2026). AlphaFold Protein Structure Database 2025: a redesigned interface and updated structural coverage. *Nucleic Acids Research*, 54(D1), D358. URL: <https://academic.oup.com/nar/article/54/D1/D358/8340156> (truy cập ngày 19 tháng 8 năm 2026)
28. [C] ChemDiv News (2026). Isomorphic Labs Launches Human Trials for AI-Designed Cancer Drugs. URL: <https://www.chemdiv.com/company/media/pharma-news/2026/isomorphic-labs-launches-human-trials-for-ai-designed-cancer-drugs/> (truy cập ngày 19 tháng 8 năm 2026)
29. [B] EMBL (2026). Millions of protein complexes added to AlphaFold Database shed light on how proteins interact. URL: <https://www.embl.org/news/science-technology/first-complexes-alphafold-database/> (truy cập ngày 19 tháng 8 năm 2026)
30. [B] Hassabis, D. (2024). Accelerating Scientific Discovery with AI. Nobel Prize Lecture in Chemistry 2024, Aula Magna, Stockholm University, Sweden.
31. [B] Nobel Prize Outreach (2024). The Nobel Prize in Chemistry 2024, Press release. NobelPrize.org. URL: <https://www.nobelprize.org/prizes/chemistry/2024/press-release/>
32. [C] EMBL-EBI (2024). Computational protein design and protein structure prediction win Nobel Prize in Chemistry. EMBL News. URL: <https://www.embl.org/news/science-technology/alphafold-wins-nobel-prize-chemistry-2024/>
33. [C] Anthropic (2026). Claude Science, an AI workbench for scientists. Anthropic News. URL: <https://www.anthropic.com/news/claude-science-ai-workbench> (truy cập ngày 19 tháng 8 năm 2026)
34. [A] TechCrunch (2026). Anthropic's Claude Science bets on workflow, not a new model, to win over scientists. TechCrunch, ngày 30 tháng 6 năm 2026. URL: <https://techcrunch.com/2026/06/30/anthropics-claude-science-bets-on-workflow-not-a-new-model-to-win-over-scientists/> (truy cập ngày 19 tháng 8 năm 2026)

35. [C] CNBC (2026). John Jumper to leave Google DeepMind for Anthropic. CNBC, ngày 19 tháng 6 năm 2026. URL: <https://www.cnbc.com/2026/06/19/john-jumper-to-leave-google-deepmind-for-anthropic.html> (truy cập ngày 19 tháng 8 năm 2026)

36. [A] Nasri, F., Gurev, S., Varilly, P. et al. (2026). Paving the way for AI agents in biology. arXiv preprint arXiv:2606.06749.

37. [A] Skarlinski, M. D., Cox, S., Laurent, J. M. et al. (2024). Language agents achieve superhuman synthesis of scientific knowledge. arXiv preprint arXiv:2409.13740.

38. [A] Lála, J., O'Donoghue, O., Shtedritski, A. et al. (2023). PaperQA: Retrieval-Augmented Generative Agent for Scientific Research. arXiv preprint arXiv:2312.07559.

39. [C] TechCrunch (2026). Nobel laureate John Jumper is leaving DeepMind for rival Anthropic. TechCrunch, ngày 20 tháng 6 năm 2026. URL: <https://techcrunch.com/2026/06/20/nobel-laureate-john-jumper-is-leaving-deepmind-for-rival-anthropic/> (truy cập ngày 19 tháng 8 năm 2026)

40. [C] CNBC (2026). Alphabet has its worst day in over a year on AI concerns after high-profile exits. CNBC, ngày 22 tháng 6 năm 2026. URL: <https://www.cnbc.com/2026/06/22/alphabet-goog-stock-ai-departures.html> (truy cập ngày 19 tháng 8 năm 2026)

41. [C] Bloomberg (2026). Alphabet Shares Drop After Second AI Star Departs for a Rival. Bloomberg, ngày 22 tháng 6 năm 2026. URL: <https://www.bloomberg.com/news/articles/2026-06-22/alphabet-shares-drop-after-second-ai-star-departs-for-a-rival> (truy cập ngày 19 tháng 8 năm 2026)

42. [B] FierceBiotech (2026). Anthropic acquires stealth AI startup Coefficient Bio in \$400M deal: reports. FierceBiotech. URL: <https://www.fiercebiotech.com/biotech/anthropic-acquires-stealth-ai-startup-coefficient-bio-400m-deal> (truy cập ngày 19 tháng 8 năm 2026)

43. [A] CNBC (2026). Anthropic launches AI drug discovery program, Claude Science. CNBC. URL: <https://www.cnbc.com/2026/06/30/anthropic-launches-ai-drug-discovery-program-claude-science.html> (truy cập ngày 19 tháng 8 năm 2026)

Chương 3 Regina Barzilay 1. [A] Stokes, J. M., Yang, K., Swanson, K., Jin, W., Cubillos-Ruiz, A., Donghia, N. M., MacNair, C. R., French, S., Carfrae, L. A., Bloom-Ackermann, Z., Tran, V. M., Chiappino-Pepe, A., Badran, A. H., Andrews, I. W., Chory, E. J., Church, G. M., Brown, E. D., Jaakkola, T. S., Barzilay, R., & Collins, J. J. (2020). A deep learning approach to antibiotic discovery. *Cell*, 180(4), 688-701. DOI: 10.1016/j.cell.2020.01.021

2. [A] Yala, A., Lehman, C., Schuster, T., Portnoi, T., & Barzilay, R. (2019). A deep learning approach to breast cancer risk prediction. *Radiology*, 292(2), 338-346. DOI: 10.1148/radiol.2019182716

3. [A] Mikhael, P. G., Wohlwend, J., Yala, A. et al. (2023). Sybil: A validated deep learning model to predict future lung cancer risk from a single low-dose chest computed tomography. *Journal of Clinical Oncology*, 41(12), 2191-2200. DOI: 10.1200/JCO.22.01345

4. [A] Yala, A., Mikhael, P. G., Connor, J., & Barzilay, R. (2022). Optimizing risk-based breast cancer screening policies with reinforcement learning. *Nature Medicine*, 28(1), 136-143. DOI: 10.1038/s41591-021-01599-w

5. [B] Barzilay, R. (2023). Deep learning for cancer diagnosis and treatment. Sana AI & Health Conference, Sana Technologies, Sweden.
6. [B] Barzilay, R. (2021). Expanding the Reach of Molecular Models in Drug Discovery. AI Helps Ukraine Charity Conference, Kiev / Boston, MA.
7. [B] MIT Jameel Clinic (2026). AI is taking on antibiotic resistance, here's how. URL: <https://jclinic.mit.edu/ai-is-taking-on-antibiotic-resistance-heres-how/> (truy cập ngày 19 tháng 8 năm 2026)
8. [B] MIT Jameel Clinic (2026). MIT Jameel Clinic AI faculty lead Regina Barzilay named in Boston Globe's Tech Power Players 50 list. URL: <https://jclinic.mit.edu/regina-barzilay-boston-globe-tech-power-players-50/> (truy cập ngày 19 tháng 8 năm 2026)
9. [B] MIT Jameel Clinic. Jameel Clinic news archive. URL: <https://jclinic.mit.edu/news/> (truy cập ngày 19 tháng 8 năm 2026)
10. [B] MIT Jameel Clinic (2026, April 30). Santa Casa approves clinical trial for MIT-developed breast cancer risk AI model, Mirai. URL: <https://jclinic.mit.edu/santa-casa-approves-clinical-trial-mirai/> (truy cập ngày 19 tháng 8 năm 2026)
11. [B] MIT Jameel Clinic. Regina Barzilay named to Boston Globe's Tech Power Players 50 list. URL: <https://jclinic.mit.edu/regina-barzilay-boston-globe-tech-power-players-50/> (truy cập ngày 19 tháng 8 năm 2026)
12. [B] WBUR (2026). AI can detect breast cancer risk before humans. Why it may take hospitals a while to adopt the tech. URL: <https://news.google.com/rss/articles/CBMimwFBVV95cUxQRnRzNjNkVlc4LXQzbWF CREtDRVJsTlNxDHfR0FBc29CYUc3YnRRaFlNStDZZE10eVBfUkV4d2xtelFGR1ZvcUoxYkNTE9tbXZMSEIze GtPbklxZTZNTVJRUGg2MHdsTWp0MmJKNXZmZ2ZLdlB3Q3d3a05Rd1JRR0t0eTNqWGI2T3BZV1M4YUJXN HFKSDdaQUJmdw> (truy cập ngày 19 tháng 8 năm 2026)
13. [B] Abdul Latif Jameel (2026). Jameel Corporation co-hosts luncheon seminar on AI-enabled preventive medicine with support from MIT Jameel Clinic. URL: <https://news.google.com/rss/articles/CBMi1gFBVV9 5cUxQUHlrb2hTaF9EckFvTGNsbVnKclhxTWVpNlhG SmFkRHY4c3U4LVdLSU93VnZTdlXMOtDdlBnWlROOW VUc2FfcGxieU1uaXJWVXlLYTVPZGdRRFhJLWV 4dkR5NTFkcGlNNnNYbUx2NlV6QS12SHZOTFlrdUVPRXZJU ONFdTVVUTcyVGtDMWJBuk9ONF9aOEdoW TZjUHFHdnpwZlQwZ2tmSjRYU19uTWZFFZ2ZCRXFYZVlfaDEyZD R4WDN2ZzY2RH11dU1NWnNwbF90a0NBdzBR> (truy cập ngày 19 tháng 8 năm 2026)
14. [B] Ynetnews (2026). Israel can lead medical AI revolution internationally, says MIT professor. URL: <https://news.google.com/rss/articles/CBMiakFVX3lxTE5mMHo5RGI3WnlIWephREo5REFEUEpEbFd oNm54MGN qamFDcDBMUDVCblpvcHZLdnFseTdoZnpKMVIZblFsTEkwdG5PQVlndndmUk5yVXFQc3d5 TlhRMFVRaUMx TXJoaV9PY0E> (truy cập ngày 19 tháng 8 năm 2026)
15. [B] BioTechniques (2026). AI-enhanced antibiotic discovery could thwart multi-drug resistant gonorrhea. BioTechniques. URL: <https://www.biotechniques.com/computational-biology/ai-enhanced-antibiotic-discovery-could-thwart-multi-drug-resistant-gonorrhea/> (truy cập ngày 19 tháng 8 năm 2026)
16. [C] CNBC (2026). MIT's Kevin Esvelt on AI's ability to create viruses: This could help replace antibiotics. CNBC. URL: <https://www.cnbc.com/video/2026/08/07/mits-kevin-esvelt-on-ais-ability-to-create-viruses-t>

his-could-help-replace-antibiotics.html (truy cập ngày 19 tháng 8 năm 2026)

Chương 4 David Baker 1. [A] Kuhlman, B., Dantas, G., Ireton, G. C., Varani, G., Stoddard, B. L., & Baker, D. (2003). Design of a novel globular protein fold with atomic-level accuracy. *Science*, 302(5649), 1364-1368. DOI: 10.1126/science.1089427

2. [A] Watson, J. L., Juergens, D., Bennett, N. R., Trippe, B. L., Yim, J., Eisenach, H. E., et al., & Baker, D. (2023). De novo protein design of structurally diverse macromolecular architectures using RFdiffusion. *Nature*, 620, 1089-1100. DOI: 10.1038/s41586-023-06415-8

3. [A] Dauparas, J., Anishchenko, I., Bennett, N., Ru, S., Ji, X., Carreira, J. R., et al., & Baker, D. (2022). Robust deep learning-based protein sequence design using ProteinMPNN. *Science*, 378(6615), 49-56. DOI: 10.1126/science.add1964

4. [A] Baker, D. (2024). Using AI for Science to Solve Humanity's Biggest Problems. Nobel Prize Lecture in Chemistry 2024, Stockholm, Sweden.

5. [A] Khandogina, N., King, N. P., & Baker, D. et al. (2021). De novo designed self-assembling protein nanoparticles as a highly potent vaccine platform. *Nature Medicine*, 27, 1234-1241.

6. [A] Musslick, S. et al. (2025). Automating the practice of science: Opportunities, challenges, and implications. *Proceedings of the National Academy of Sciences*, 122(5), e2401238121. DOI: 10.1073/pnas.2401238121

7. [A] Luebbert, L., Gurev, S., & Rodriques, S. G. (2026). Paving the way for AI agents in biology. Anthropic Science Event Briefing Whitepaper, June 30, 2026.

8. [B] Washington Research Foundation (2026). Washington Research Foundation awards \$7 million to UW Institute for Protein Design for ambitious initiative to advance AI-enabled enzyme design. URL: <https://www.wrfseattle.org/news/washington-research-foundation-awards-7-million-to-uw-institute-for-protein-design-for-ambitious-initiative-to-advance-ai-enabled-enzyme-design/> (truy cập ngày 19 tháng 8 năm 2026)

9. [B] GeekWire (2026). Nobel laureate David Baker lands \$7M for new protein design program. GeekWire. URL: <https://www.geekwire.com/2026/from-computer-to-lab-to-market-nobel-winner-david-baker-lands-7m-for-new-protein-program/> (truy cập ngày 19 tháng 8 năm 2026)

10. [A] Watson, J. L., Juergens, D., Bennett, N. R., Trippe, B. L., Yim, J., Eisenach, H. E. et al. (2023). De novo protein design of structurally diverse macromolecular architectures using RFdiffusion. *Nature*, 620, 1089-1100.

11. [A] Baek, M., DiMaio, F., Anishchenko, I., Dauparas, J., Ovchinnikov, S., Lee, G. R. et al. (2021). Accurate prediction of protein structures and interactions using a three-track neural network. *Science*, 373(6557), 871-876. DOI: 10.1126/science.abj8754

12. [A] Khandogina, N., King, N. P., Baker, D. et al. (2021). De novo designed self-assembling protein nanoparticles as a highly potent vaccine platform. *Nature Medicine*, 27, 1234-1241.

13. [A] Life Science Washington (2026). From Computer to Lab to Market: Nobel Winner David Baker Lands \$7M for New Protein Program. URL: <https://lifesciencewa.org/2026/03/24/from-computer-to-lab-to-market-nobel-winner-david-baker-lands-7m-for-new-protein-program/> (truy cập ngày 19 tháng 8 năm 2026)

14. [A] Anfinsen, C. B. (1973). Principles that Govern the Folding of Protein Chains. *Science*, 181(4096), 223-230. DOI: 10.1126/science.181.4096.223
15. [A] Moult, J., Pedersen, J. T., Judson, R., & Fidelis, K. (1995). A large-scale experiment to assess protein structure prediction methods. *Proteins*, 23(3), ii-v. DOI: 10.1002/prot.340230303
16. [A] Das, R., & Baker, D. (2008). Macromolecular Modeling with Rosetta. *Annual Review of Biochemistry*, 77, 363-382. DOI: 10.1146/annurev.biochem.77.062906.171838
17. [A] Cooper, S., Khatib, F., Treuille, A. et al. (2010). Predicting protein structures with a multiplayer online game. *Nature*, 466, 756-760. DOI: 10.1038/nature09304
18. [A] Khatib, F., DiMaio, F., Foldit Contenders Group, Foldit Void Crushers Group et al. (2011). Crystal structure of a monomeric retroviral protease solved by protein folding game players. *Nature Structural & Molecular Biology*, 18(10), 1175-1177. DOI: 10.1038/nsmb.2119
19. [A] Rosetta Commons (2026). Brazil's Largest Open Science Event Introduces Thousands to Protein Design with Foldit. Rosetta Commons. URL: <https://rosettacommons.org/2026/07/15/brazils-largest-open-science-event-introduces-thousands-to-protein-design-with-foldit/> (truy cập ngày 19 tháng 8 năm 2026)
20. [A] Bennett, N. R., Watson, J. L., Ragotte, R. J., Borst, A. J., See, D. L., Weidle, C. et al. (2026). Atomically accurate de novo design of antibodies with RFdiffusion. *Nature*, 649(8095), 183-193. DOI: 10.1038/s41586-025-09721-5
21. [C] GenEngNews (2025). RFdiffusion3 Now Open Source, Designs DNA Binders and Advanced Enzymes. GEN Biotechnology News. URL: <https://www.genengnews.com/topics/artificial-intelligence/rfdiffusion3-now-open-source-designs-dna-binders-and-advanced-enzymes/> (truy cập ngày 19 tháng 8 năm 2026)
22. [B] Baker Lab, University of Washington (2025). Designing antibodies with RFdiffusion. URL: <https://www.bakerlab.org/2025/02/28/designing-antibodies-with-rfdiffusion/> (truy cập ngày 19 tháng 8 năm 2026)
23. [B] Institute for Protein Design, University of Washington (2025). Teaching AI to build antibodies from scratch. URL: <https://www.ipd.uw.edu/2025/11/rfantibody-in-nature> (truy cập ngày 19 tháng 8 năm 2026)
24. [B] GeekWire (2025). UW Nobel winner's lab releases most powerful protein design tool yet. URL: <https://www.geekwire.com/2025/uw-nobel-winners-lab-releases-most-powerful-protein-design-tool-yet/> (truy cập ngày 19 tháng 8 năm 2026)
25. [A] Walls, A.C., Fiala, B., Schäfer, A. et al. (2020). Elicitation of Potent Neutralizing Antibody Responses by Designed Protein Nanoparticle Vaccines for SARS-CoV-2. *Cell*, 183(5), 1367-1382. DOI: 10.1016/j.cell.2020.10.043
26. [A] Dauparas, J., Anishchenko, I., Bennett, N. et al. (2022). Robust Deep Learning-Based Protein Sequence Design Using ProteinMPNN. *Science*, 378(6615), 49-56. DOI: 10.1126/science.add2187
27. [B] The Korea Times (2022). SK Bioscience Gets Final Approval for Korea's 1st COVID-19 Vaccine. The Korea Times. URL: <https://www.koreatimes.co.kr/business/companies/20220629/sk-bioscience-gets-final-approval-for-koreas-1st-covid-19-vaccine> (truy cập ngày 19 tháng 8 năm 2026)
28. [B] GeekWire (2026). Lab of UW Nobel Winner Cracks Challenge of Creating Roomier Protein Cages to Deliver Genetic Medicines. GeekWire. URL: <https://www.geekwire.com/2026/uw-nobel-winning-lab-creat>

es-blueprints-for-roomier-protein-cages-to-deliver-genetic-medicines/ (truy cập ngày 19 tháng 8 năm 2026)

29. [B] Institute for Protein Design (2026). A New Era for Protein Design in Oncology: Dr. Ahn Joins the IPD. University of Washington. URL:

<https://www.ipd.uw.edu/2026/05/a-new-era-for-protein-design-in-oncology-dr-ahn-joins-the-ipd/> (truy cập ngày 19 tháng 8 năm 2026)

30. [B] Nobel Prize Outreach (2024). Press release: The Nobel Prize in Chemistry 2024. NobelPrize.org. URL: <https://www.nobelprize.org/prizes/chemistry/2024/press-release/>

31. [B] Institute for Protein Design, University of Washington (2024). David Baker wins Nobel Prize for protein design. URL: <https://www.ipd.uw.edu/2024/10/david-baker-wins-nobel-prize-for-protein-design> (truy cập ngày 19 tháng 8 năm 2026)

Chương 5 Alán Aspuru-Guzik 1. [A] Aspuru-Guzik, A., Dutoi, A. D., Love, P., Head-Gordon, M. (2005). Simulated Quantum Computation of Molecular Energies. *Science*, 309(5741), 1704-1707. DOI: 10.1126/science.1113479

2. [A] Hachmann, J. et al. (2011). The Harvard Clean Energy Project: Large-Scale Computational Screening and Design of Organic Photovoltaics on the World Community Grid. *The Journal of Physical Chemistry Letters*, 2(17), 2241-2251. DOI: 10.1021/jz200866s

3. [A] Huskinson, B., Marshak, M. P., Suh, C., Er, S., Gerhardt, M. R., Galvin, C. J., Chen, X., Aspuru-Guzik, A., Gordon, R. G., Aziz, M. J. (2014). A metal-free organic-inorganic aqueous flow battery. *Nature*, 505, 195-198. DOI: 10.1038/nature12909

4. [C] University of Toronto (2023). U of T receives \$200-million grant to support Acceleration Consortium's 'self-driving labs' research. University of Toronto News. URL: <https://www.utoronto.ca/news/u-t-receives-200-million-grant-support-acceleration-consortium-s-self-driving-labs-research> (truy cập ngày 19 tháng 8 năm 2026)

5. [B] Globalnews.ca (2023). 'Self-driving' AI lab awarded \$200M grant to pursue new drugs, materials. URL: <https://globalnews.ca/news/9657567/ai-canada-university-of-toronto-self-driving-lab-research-grant-acceleration-consortium/> (truy cập ngày 19 tháng 8 năm 2026)

6. [B] Acceleration Consortium, University of Toronto (2026). Global research consortia join forces to accelerate AI-driven drug discovery. URL: <https://acceleration.utoronto.ca/news/global-research-consortia-join-forces-to-accelerate-ai-driven-drug-discovery> (truy cập ngày 19 tháng 8 năm 2026)

7. [C] EurekAlert! (2026). Global research consortia join forces to accelerate AI-driven drug discovery. URL: <https://www.eurekalert.org/news-releases/1131549> (truy cập ngày 19 tháng 8 năm 2026)

8. [B] Matter Lab, University of Toronto. About Alán. URL: <https://www.matter.toronto.edu/basic-content-page/about-alan> (truy cập ngày 19 tháng 8 năm 2026)

9. [A] Krenn, M., Häse, F., Nigam, A., Friederich, P., & Aspuru-Guzik, A. (2020). Self-referencing embedded strings (SELFIES): A 100% robust molecular string representation. *Machine Learning: Science and Technology*, 1(4), 045024. DOI: 10.1088/2632-2153/aba947

10. [A] Häse, F., Roch, L. M., & Aspuru-Guzik, A. (2018). Chimera: enabling hierarchy based multi-objective optimization for self-driving laboratories. *Chemical Science*, 9(39), 7642-7655. DOI: 10.1039/c8sc02239a

11. [A] Häse, F., Aldeghi, M., Hickman, R. J., Roch, L. M., & Aspuru-Guzik, A. (2021). Gryffin: An algorithm for Bayesian optimization of categorical variables informed by expert knowledge. *Applied Physics Reviews*, 8(3), 031406. DOI: 10.1063/5.0048164
12. [A] Tom, G., Schmid, S. P., Baird, S. G. et al., & Aspuru-Guzik, A. (2024). Self-driving laboratories for chemistry and materials science. *Chemical Reviews*, 124(16), 9633-9732. DOI: 10.1021/acs.chemrev.4c00055
13. [A] Kumacheva, E., Aspuru-Guzik, A. et al. (2025). Self-driving lab for the photochemical synthesis of plasmonic nanoparticles with targeted structural and optical properties. *Nature Communications*. DOI: 10.1038/s41467-025-56788-9
14. [B] The Varsity (2026). Lash Miller expansion breaks ground for AI-driven materials research. URL: <https://thevarsity.ca/2026/03/22/lash-miller-expansion-breaks-ground-for-ai-driven-materials-research/> (truy cập ngày 19 tháng 8 năm 2026)
15. [C] EurekAlert (2023). Acceleration Consortium applies artificial intelligence to discovery of advanced materials. URL: <https://www.eurekalert.org/news-releases/666819> (truy cập ngày 19 tháng 8 năm 2026)
16. [A] Gomez-Bombarelli, R., Aguilera-Iparraguirre, J., Hirzel, T. D., Duvenaud, D., Maclaurin, D., Adams, R. P., Shkurti, F., Garg, A., & Aspuru-Guzik, A. (2018). Automatic chemical design using a data-driven continuous representation of molecules. *ACS Central Science*, 4(2), 268-276. DOI: 10.1021/acscentsci.7b00572
17. [A] Sanchez-Lengeling, B., & Aspuru-Guzik, A. (2018). Inverse molecular design with generative models. *Science*, 361(6400), 360-365. DOI: 10.1126/science.aat2684
18. [A] Burger, B., Maffettone, P. M., Gusev, V. V., Aitchison, C. M., Bai, Y., Wang, X., Li, X., Alston, B. M., Li, B., Clowes, R., Rankin, J., Cooper, A. I., & Aspuru-Guzik, A. (2020). A mobile robotic chemist. *Nature*, 583(7818), 719-723. DOI: 10.1038/s41586-020-2442-2
19. [B] Aspuru-Guzik, A. (2021). There is no time for science as usual: materials acceleration platforms. MIT.nano Distinguished Seminar Series, Cambridge, MA.
20. [B] Aspuru-Guzik, A. (2020). Automated Material Discovery with Self-Driving Labs. The 26th Annual Shih-I Pai Lecture, University of Maryland, College Park, MD.
21. [A] Nature (2026). Inside the 'self-driving' lab revolution. *Nature*. URL: <https://www.nature.com/articles/d41586-026-00974-2> (truy cập ngày 19 tháng 8 năm 2026)
22. [C] Chemical & Engineering News (2026). Self-driving labs are changing how chemists work. C&EN, American Chemical Society. URL: <https://cen.acs.org/physical-chemistry/computational-chemistry/Self-driving-labs-changing-chemists/104/web/2026/06> (truy cập ngày 19 tháng 8 năm 2026)
23. [A] Faculty of Arts & Science, University of Toronto (2026). University of Toronto breaks ground on new home for the Acceleration Consortium. University of Toronto. URL: <https://www.artsci.utoronto.ca/news/university-toronto-breaks-ground-new-home-acceleration-consortium> (truy cập ngày 19 tháng 8 năm 2026)
24. [A] Musslick, S., Bartlett, L. K., Chandramouli, S. H. et al. (2025). Automating the practice of science: Opportunities, challenges, and implications. *Proceedings of the National Academy of Sciences*, 122(5), e2401238121. DOI: 10.1073/pnas.2401238121

25. [A] Strieth-Kalthoff, F., Hao, H., Rathore, V. et al., & Aspuru-Guzik, A. (2024). Delocalized, asynchronous, closed-loop discovery of organic laser emitters. *Science*, 384(6697), eadk9227. DOI: 10.1126/science.adk9227
 26. [B] Structural Genomics Consortium (2026). Structural Genomics Consortium to lead AI-assisted Drug Discovery Initiative within the University of Toronto's Acceleration Consortium. Structural Genomics Consortium. URL: <https://www.thesgc.org/news/structural-genomics-consortium-lead-ai-assisted-drug-discovery-initiative-within-university> (truy cập ngày 19 tháng 8 năm 2026)
 27. [C] News-Medical.net (2026). Automated self-driving labs accelerate the search for new medicines. News-Medical.net. URL: <https://www.news-medical.net/news/20260610/Automated-self-driving-labs-accelerate-the-search-for-new-medicines.aspx> (truy cập ngày 19 tháng 8 năm 2026)
 28. [A] Nigam, A. K., Pollice, R., & Aspuru-Guzik, A. (2022). JANUS: Parallel tempered genetic algorithm guided by deep neural networks for inverse molecular design. *Digital Discovery*, 1(4), 390-404. DOI: 10.1039/D2DD00003B
 29. [A] Boiko, D. A., MacKnight, R., Kline, B., & Gomes, G. (2023). Autonomous chemical research with large language models. *Nature*, 624(7992), 570-578. DOI: 10.1038/s41586-023-06792-0
 30. [C] GlobeNewswire (2026). Zapata Quantum partners with QuEra to advance the commercial viability of quantum applications. URL: <https://www.globenewswire.com/news-release/2026/08/12/3343489/0/en/zapata-quantum-partners-with-quera-to-advance-the-commercial-viability-of-quantum-applications.html> (truy cập ngày 19 tháng 8 năm 2026)
 31. [B] BetaKit (2026). Dr. Alán Aspuru-Guzik. Most Ambitious 2026. URL: <https://mostambitious.betakit.com/2026/dr-alan-aspuru-guzik/> (truy cập ngày 19 tháng 8 năm 2026)
- Chương 6 Hod Lipson 1. [A] Lipson, H., & Pollack, J. B. (2000). Automatic design and manufacture of robotic lifeforms. *Nature*, 406(6799), 974-978. DOI: 10.1038/35023115
2. [A] Bongard, J. C., & Lipson, H. (2007). Automated reverse engineering of nonlinear dynamical systems. *Proceedings of the National Academy of Sciences*, 104(24), 9943-9948. DOI: 10.1073/pnas.0609476104
 3. [A] Schmidt, M. D., & Lipson, H. (2008). Coevolution of fitness predictors. *IEEE Transactions on Evolutionary Computation*, 12(6), 736-749. DOI: 10.1109/TEVC.2008.919006
 4. [A] Schmidt, M., & Lipson, H. (2009). Distilling free-form natural laws from experimental data. *Science*, 324(5923), 81-85. DOI: 10.1126/science.1165893
 5. [A] Stoutemyer, D. R. (2012). Can the Eureka symbolic regression program, computer algebra and numerical analysis help each other? *arXiv preprint arXiv:1203.1023*.
 6. [A] Chen, B., Huang, K., Raghupathi, S., Chandratreya, I., Du, Q., & Lipson, H. (2022). Automated discovery of fundamental variables hidden in experimental data. *Nature Computational Science*, 2(7), 433-442. DOI: 10.1038/s43588-022-00281-6
 7. [A] Udrescu, S. M., & Tegmark, M. (2020). AI Feynman: A physics-inspired method for symbolic regression. *Science Advances*, 6(16), eaay2631. DOI: 10.1126/sciadv.aay2631
 8. [A] Lazebnik, T., & Liberzon, A. (2026). Moving from table to graph in physics-informed spatio-temporal symbolic regression. *Scientific Reports*, 16(1), 16016. DOI: 10.1038/s41598-026-53882-w

9. [C] Phys.org (2026). Researchers develop AI tool that finds the equations behind complex systems. Phys.org. URL: <https://phys.org/news/2026-07-ai-tool-equations-complex.html> (truy cập ngày 19 tháng 8 năm 2026)
10. [A] Bongard, J., Zykov, V., & Lipson, H. (2006). Resilient Machines through Self-Modeling. *Science*, 314(5802), 1118-1121. DOI: 10.1126/science.1133687
11. [A] Kwiatkowski, R., & Lipson, H. (2019). Task-agnostic self-modeling machines. *Science Robotics*, 4(26), eaau9354. DOI: 10.1126/scirobotics.aau9354
12. [A] Chen, B., Kwiatkowski, R., Vondrick, C., & Lipson, H. (2022). Fully body visual self-modeling of robot morphologies. *Science Robotics*, 7(68), eabn1944. DOI: 10.1126/scirobotics.abn1944
13. [B] Lipson, H. (2012). RI Seminar: Distilling Natural Laws from Experimental Data. CMU Robotics Institute Colloquium Series, Carnegie Mellon University, Pittsburgh, PA.
14. [B] Lipson, H. (2022). Automating discovery: From cognitive robotics to particle physics. AI Forward Forum Lecture, online.
15. [A] Wyder, P. M., Lipson, H., et al. (2026). Robot metabolism: Toward machines that can grow by consuming other machines. *Science Advances*. DOI: 10.1126/sciadv.adu6897 (truy cập ngày 19 tháng 8 năm 2026)
16. [A] arXiv (2026). Proprioceptive-visual correspondence enables self-other distinction in humanoid robots. arXiv:2606.13222. URL: <https://arxiv.org/abs/2606.13222> (truy cập ngày 19 tháng 8 năm 2026)
17. [A] Discovering equations from data: symbolic regression in dynamical systems. (2026). *Journal of Physics: Complexity*. URL: <https://iopscience.iop.org/article/10.1088/2632-072X/ae3eed> (truy cập ngày 19 tháng 8 năm 2026)
18. [B] Bayesian symbolic regression: automated equation discovery from a physicist's perspective. (2026). *Philosophical Transactions of the Royal Society A*, 384(2317), 20250089. URL: <https://royalsocietypublishing.org/rsta/article/384/2317/20250089/481281/Bayesian-symbolic-regression-automated-equation> (truy cập ngày 19 tháng 8 năm 2026)
19. [A] SR-Scientist: Scientific Equation Discovery With Agentic AI. (2026). arXiv preprint arXiv:2510.11661. URL: <https://arxiv.org/pdf/2510.11661> (truy cập ngày 19 tháng 8 năm 2026)
20. [B] Eureka. Wikipedia. URL: <https://en.wikipedia.org/wiki/Eureka> (truy cập ngày 19 tháng 8 năm 2026)
21. [A] Schmidt, M. D., & Lipson, H. (2011). Age-fitness pareto optimization. In *Genetic Programming Theory and Practice VIII*. Springer. DOI: 10.1007/978-1-4419-7747-2_8
22. [A] Koza, J. R. (1994). Genetic programming as a means for programming computers by natural selection. *Statistics and Computing*, 4(2). DOI: 10.1007/bf00175355
23. [A] Stoutemyer, D. R. (2013). Can the Eureka symbolic regression program, computer algebra, and numerical analysis help each other? *Notices of the American Mathematical Society*, 60(6), 713. DOI: 10.1090/noti1000
24. [A] Tonda, A. (2025). Review of PySR: high-performance symbolic regression in Python and Julia. *Genetic Programming and Evolvable Machines*, 26(1). DOI: 10.1007/s10710-024-09503-4
25. [B] Cranmer, M. (2026). PySR v2.0.0-beta.2 release notes. GitHub. URL: <https://github.com/MilesCranmer/PySR/releases> (truy cập ngày 19 tháng 8 năm 2026)

26. [A] Stoutemyer, D. R. (2012). Can the Eureka symbolic regression program, computer algebra and numerical analysis help each other? arXiv preprint arXiv:1203.1023, 1-17.
 27. [A] Brki , D., Praks, P., Marek, M., Ili , U., & Staji , Z. (2025). Reducing the number of input variables through symbolic regression. Electronic Research Archive (ERA), 33(9), 5158-5178. DOI: 10.3934/era.2025231
 28. [A] Jhunhunwala, A., Goldfeder, J., & Lipson, H. (2026). Evidence of an Emergent "Self" in Continual Robot Learning. arXiv preprint arXiv:2603.24350. URL: <https://arxiv.org/abs/2603.24350> (truy cập ngày 19 tháng 8 năm 2026)
 29. [A] Cox, N., Grundler, X., & Li, B.-A. (2026). Symbolic Regression for Interpretable Emulation of Proton Collective Flow in Intermediate-Energy Heavy-Ion Collisions. arXiv preprint arXiv:2608.17828. URL: <https://arxiv.org/abs/2608.17828> (truy cập ngày 19 tháng 8 năm 2026)
- Chương 7 Max Tegmark
1. [A] Udrescu, S.M., Tegmark, M. (2020). AI Feynman: A physics-inspired method for symbolic regression. Science Advances, 6(16), eaay2631. DOI: 10.1126/sciadv.aay2631
 2. [A] Udrescu, S.M., Tan, A., Feng, J., Neto, O., Wu, T., Tegmark, M. (2020). AI Feynman 2.0: Pareto-optimal symbolic regression exploiting graph modularity. Advances in Neural Information Processing Systems 33 (NeurIPS 2020). arXiv:2006.10782
 3. [A] Zhong, Z., Liu, Z., Tegmark, M., Andreas, J. (2023). The Clock and the Pizza: Two Stories in Mechanistic Explanation of Neural Networks. Advances in Neural Information Processing Systems 36 (NeurIPS 2023). arXiv:2306.17844
 4. [A] Power, A., Burda, Y., Edwards, H., Babuschkin, I., Misra, V. (2022). Grokking: Generalization Beyond Overfitting on Small Algorithmic Datasets. arXiv:2201.02177
 5. [A] Nanda, N., Chan, L., Lieberum, T., Smith, J., Steinhardt, J. (2023). Progress measures for grokking via mechanistic interpretability. arXiv:2301.05217
 6. [A] Tegmark, M. (2014). Our Mathematical Universe: My Quest for the Ultimate Nature of Reality. Knopf.
 7. [C] Future of Life Institute (2026). The most urgent warning shot yet. FLI Newsletter. URL: <https://newsletter.futureoflife.org/p/fli-newsletter-july-2026> (truy cập ngày 19 tháng 8 năm 2026)
 8. [C] Al Jazeera, The Bottom Line (2026). Why are sandwiches more regulated than AI? URL: <https://www.aljazeera.com/program/the-bottom-line/2026/8/16/why-are-sandwiches-more-regulated-than-ai> (truy cập ngày 19 tháng 8 năm 2026)
 9. [B] Tegmark, M. (2017). Life 3.0: Being Human in the Age of Artificial Intelligence. Knopf, New York, NY.
 10. [B] Tegmark, M. (2021). "AI and Physics." The Lex Fridman Podcast, Episode No. 155, Cambridge, MA.
 11. [B] Tegmark, M. (2023). "Keeping AI under control with mechanistic interpretability." MIT Department of Physics Colloquium Series, Massachusetts Institute of Technology, Cambridge, MA.
 12. [A] Udrescu, S. M., Tan, A., Feng, J., Neto, O., Wu, T., & Tegmark, M. (2020). AI Feynman 2.0: Pareto-optimal symbolic regression exploiting graph modularity. arXiv preprint arXiv:2006.10782.
 13. [A] Udrescu, S.-M., Tan, A., Feng, J., Neto, O., Wu, T., & Tegmark, M. (2020). AI Feynman 2.0: Pareto-optimal symbolic regression exploiting graph modularity. Advances in Neural Information

Processing Systems 33 (NeurIPS 2020). URL: <https://arxiv.org/abs/2006.10782>

14. [A] Schmidt, M., & Lipson, H. (2009). Distilling Free-Form Natural Laws from Experimental Data. *Science*, 324(5923), 81-85. DOI: 10.1126/science.1165893

15. [A] Power, A., Burda, Y., Edwards, H., Babuschkin, I., & Misra, V. (2022). Grokking: Generalization Beyond Overfitting on Small Algorithmic Datasets. URL: <https://arxiv.org/abs/2201.02177>

16. [A] Huh, M., Cheung, B., Wang, T., & Isola, P. (2024). The Platonic Representation Hypothesis. URL: <https://arxiv.org/abs/2405.07987>

17. [A] Buckingham, E. (1914). On Physically Similar Systems; Illustrations of the Use of Dimensional Equations. *Physical Review*, 4(4), 345-376. DOI: 10.1103/PhysRev.4.345

18. [B] NSF AI Institute for Artificial Intelligence and Fundamental Interactions (IAIFI) (2026). Five-year NSF renewal announcement. URL: <https://iaifi.org/> (truy cập ngày 19 tháng 8 năm 2026)

19. [B] Democracy Now! (2026). "Less Regulated Than Sandwiches": MIT Prof. Calls for Oversight as OpenAI Agent Hacks Other Firms. URL: https://www.democracynow.org/2026/7/30/max_tegmark (truy cập ngày 19 tháng 8 năm 2026)

20. [B] Tegmark, M. (2023). Keeping AI under control with mechanistic interpretability. MIT Department of Physics Colloquium Series, Massachusetts Institute of Technology, Cambridge, MA.

21. [B] Tegmark, M. (2021). AI and Physics. The Lex Fridman Podcast, Episode 155, Cambridge, MA.

22. [B] Politico (2026). 5 Questions for Max Tegmark. Politico. URL: <https://news.google.com/rss/search?q=Max+Tegmark+AI+2026> (truy cập ngày 19 tháng 8 năm 2026)

23. [A] Towards Data Science (2026). Mechanistic Interpretability: Peeking Inside an LLM. Towards Data Science. (truy cập ngày 19 tháng 8 năm 2026)

Chương 8 Ross King 1. [A] King, R. D., Rowland, J., Oliver, S. G., Young, M., Aubrey, W., Byrne, E., Liakata, M., Muggleton, S., Whelan, K. E., & Soldatova, L. N. (2009). The automation of science. *Science*, 324(5923), 85-89. DOI: 10.1126/science.1165620

2. [A] King, R. D., Whelan, K. E., Jones, F. M., Reiser, P. G., Bryant, C. H., Muggleton, S. H., Kell, D. B., & Oliver, S. G. (2004). Functional genomic hypothesis generation and experimentation by a robot scientist. *Nature*, 427(6971), 247-252. DOI: 10.1038/nature02236

3. [A] Kitano, H. (2021). Nobel Turing Challenge: Creating the engine for scientific discovery. *npj Systems Biology and Applications*, 7(1), 29. DOI: 10.1038/s41540-021-00189-3

4. [A] Bilsland, E., Sparkes, A., Williams, K., Moss, H. J., de Oliveira, S. T., & King, R. D. et al. (2018). Plasmodium dihydrofolate reductase is a second enzyme target for the antimalarial action of triclosan. *Scientific Reports*, 8, Article 1038. DOI: 10.1038/s41598-018-19549-x

5. [A] Bilsland, E. et al. (2013). Yeast-based automated high-throughput screens to identify anti-parasitic lead compounds. *Open Biology*, 3(9), 120158. DOI: 10.1098/rsob.120158

6. [A] Musslick, S., Mahesh, S., Sloman, S. J., & King, R. D. et al. (2025). Automating the practice of science: Opportunities, challenges, and implications. *Proceedings of the National Academy of Sciences*, 122(5), e2401238121. DOI: 10.1073/pnas.2401238121

7. [A] Schmidt, M., & Lipson, H. (2009). Distilling free-form natural laws from experimental data. *Science*, 324(5923), 81-85.

8. [C] Scallan, E. (2026). What I've learned from 25 years of automated science, and what the future holds: an interview with Ross King. AIhub. URL: <https://aihub.org/2026/03/30/what-ive-learned-from-25-years-of-automated-science-and-what-the-future-holds-an-interview-with-ross-king/> (truy cập ngày 19 tháng 8 năm 2026)
9. [A] Nature Portfolio (2026). Investigating uncharacterised genes in *Saccharomyces cerevisiae* using robot scientists. Scientific Reports. URL: <https://www.nature.com/articles/s41598-026-46236-z> (truy cập ngày 19 tháng 8 năm 2026)
10. [C] Chalmers University of Technology (không rõ năm). Chalmers' Robot Scientist ready for drug discovery. Chalmers News. URL: <https://www.chalmers.se/en/current/news/life-chalmers-robot-scientist-ready-for-drug-discovery/> (truy cập ngày 19 tháng 8 năm 2026)
11. [A] Sparkes, A., Aubrey, W., Byrne, E., Clare, A., Khan, M.N., Liakata, M., Markham, M., Rowland, J., Soldatova, L.N., Whelan, K.E., Young, M., King, R.D. (2010). Towards Robot Scientists for autonomous scientific discovery. *Automated Experimentation*, 2, 1. DOI: 10.1186/1759-4499-2-1
12. [A] Williams, K., Bilsland, E., Sparkes, A., Aubrey, W., Young, M., Soldatova, L.N. et al. (2015). Cheaper faster drug development validated by the repositioning of drugs against neglected tropical diseases. *Journal of the Royal Society Interface*, 12(104), 20141289. DOI: 10.1098/rsif.2014.1289
13. [A] WASP, Wallenberg AI, Autonomous Systems and Software Program (2020). The Automation of Science, for the Benefit of Mankind. URL: <https://wasp-sweden.org/the-automation-of-science-for-the-benefit-of-mankind/> (truy cập ngày 19 tháng 8 năm 2026)
14. [C] AIhub (2026). What I've learned from 25 years of automated science, and what the future holds: an interview with Ross King. URL: <https://aihub.org/2026/03/30/what-ive-learned-from-25-years-of-automated-science-and-what-the-future-holds-an-interview-with-ross-king/> (truy cập ngày 19 tháng 8 năm 2026)
15. [B] Robohub (2026). What I've learned from 25 years of automated science, and what the future holds: an interview with Ross King. URL: <https://robohub.org/what-ive-learned-from-25-years-of-automated-science-and-what-the-future-holds-an-interview-with-ross-king/> (truy cập ngày 19 tháng 8 năm 2026)
16. [C] Aberystwyth University (2009). First science discovery for robot. Aberystwyth University News. URL: <https://www.aber.ac.uk/en/news/archive/2009/april/title-77781-en.html> (truy cập ngày 19 tháng 8 năm 2026)
17. [C] University of Cambridge (2009). Robot scientist becomes first machine to discover new scientific knowledge. University of Cambridge Research News. URL: <https://www.cam.ac.uk/research/news/robot-scientist-becomes-first-machine-to-discover-new-scientific-knowledge> (truy cập ngày 19 tháng 8 năm 2026)
18. [A] ScienceDaily (2009). Robot Scientist Becomes First Machine To Discover New Scientific Knowledge. URL: <https://www.sciencedaily.com/releases/2009/04/090402143451.htm> (truy cập ngày 19 tháng 8 năm 2026)
19. [C] Nature (2026). Inside the 'self-driving' lab revolution. Nature News. URL: <https://www.nature.com/articles/d41586-026-00974-2> (Xuất bản ngày 30 tháng 3 năm 2026, truy cập ngày 19 tháng 8 năm 2026)

19 tháng 8 năm 2026)

20. [B] 3 Quarks Daily (2026). Inside the 'self-driving' lab revolution. 3 Quarks Daily. URL: <https://3quarksdaily.com/3quarksdaily/2026/04/inside-the-self-driving-lab-revolution.html> (Truy cập ngày 19 tháng 8 năm 2026)

21. [A] Bjurström, E. Y., Gower, A. H., Lasin, P., Savolainen, O. I., Tiukova, I. A., & King, R. D. (2026). Investigating uncharacterised genes in *Saccharomyces cerevisiae* using robot scientists. *Scientific Reports*, 16, 10999. DOI: 10.1038/s41598-026-46236-z

22. [A] Boiko, D. A., MacKnight, R., Kline, B., & Gomes, G. (2023). Autonomous chemical research with large language models. *Nature*, 624(7992), 570-578. DOI: 10.1038/s41586-023-06792-0

23. [A] Bisht, H., Kumar, V., Jablonka, K. M., Mausam, & Krishnan, N. M. A. (2026). Agentic AI scientists are not built for autonomous scientific discovery. arXiv:2605.08956. URL: <https://arxiv.org/abs/2605.08956> (Truy cập ngày 19 tháng 8 năm 2026)

Chương 9 Gabriel Gomes 1. [A] Boiko, D. A., MacKnight, R., Kline, B., & Gomes, G. (2023). Autonomous chemical research with large language models. *Nature*, 624(7992), 570-578. DOI: 10.1038/s41586-023-06792-0

2. [A] Sanchez-Lengeling, B., & Aspuru-Guzik, A. (2018). Inverse molecular design with generative models. *Science*, 361(6400), 360-365.

3. [A] Musslick, S., Bartlett, L. K., Chandramouli, S. H., Dubova, M., Gobet, F., Griffiths, T. L., Hullman, J., King, R. D., Kutz, J. N., Lucas, C. G., Mahesh, S., Pestilli, F., Sloman, S. J., & Holmes, W. R. (2025). Automating the practice of science: Opportunities, challenges, and implications. *Proceedings of the National Academy of Sciences (PNAS)*, 122(5), e2401238121. DOI: 10.1073/pnas.2401238121

4. [A] Kitano, H. (2021). Nobel Turing Challenge: Creating the engine for scientific discovery. *npj Systems Biology and Applications*, 7(1), 29. DOI: 10.1038/s41540-021-00189-3

5. [A] Gomes, G. (2024). Autonomous chemical research of large language models. *Science in the Age of Experience Keynote Address*, Dassault Systèmes, Boston, MA.

6. [A] Gomes, G. (2026). On Coscientist: The AI Agent Doing 6 Months of Science in 4 Hrs. *American AI Festival 2026*, Washington, D.C.

7. [C] Scientific American (2026). How one chemist is using AI and robots to automate lab experiments. *Scientific American*. URL: <https://www.scientificamerican.com/article/how-one-chemist-is-using-ai-and-robots-to-automate-lab-experiments/> (Truy cập ngày 19 tháng 8 năm 2026)

8. [D] The Org (2026). Gabe Gomes, Research Scientist at X, the moonshot factory. URL: <https://theorg.com/org/x-the-moonshot-factory/org-chart/gabe-gomes> (Truy cập ngày 19 tháng 8 năm 2026)

9. [B] Carnegie Mellon University College of Engineering. Gabe Gomes faculty profile. URL: <https://engineering.cmu.edu/directory/bios/gomes-gabe.html> (Truy cập ngày 19 tháng 8 năm 2026)

10. [C] Carnegie Mellon University, EurekAlert (2023). Carnegie Mellon-designed artificially intelligent coscientist automates scientific discovery. *EurekAlert*. URL: <https://www.eurekalert.org/news-releases/1011392> (Truy cập ngày 19 tháng 8 năm 2026)

11. [C] Chemical & Engineering News (2026). Self-driving labs are changing how chemists work. *C&EN*. URL: <https://cen.acs.org/physical-chemistry/computational-chemistry/Self-driving-labs-changing-chemists/10>

4/web/2026/06 (Truy cập ngày 19 tháng 8 năm 2026)

12. [A] Carnegie Mellon University Mellon College of Science (2023). CMU-Designed Artificially Intelligent Coscientist Automates Scientific Discovery. URL:
https://www.cmu.edu/mcs/news-events/2023/1220_ai-coscientist-automates-discovery.html

13. [C] EurekAlert! (2023). Meet 'Coscientist,' your AI lab partner. URL:
<https://www.eurekalert.org/news-releases/1029545>

14. [C] Scientific American (2024). How one chemist is using AI and robots to automate lab experiments. URL:
<https://www.scientificamerican.com/article/how-one-chemist-is-using-ai-and-robots-to-automate-lab-experiments/>

15. [B] Quantum Zeitgeist (2026). AI Chemist 'Coscientist' Executes Nobel-Winning Experiments on First Try. URL:
<https://quantumzeitgeist.com/ai-chemist-coscientist-executes-nobel-winning-experiments-on-first-try/>
(Truy cập ngày 19 tháng 8 năm 2026)

16. [A] Royal Society Open Science (2025). Autonomous 'self-driving' laboratories: a review of technology and applications. Royal Society Open Science, 12(7), 250646. URL:
<https://royalsocietypublishing.org/rso/article/12/7/250646/235354/Autonomous-self-driving-laboratories-a-review-of>

17. [A] Experiment-as-Code Labs: A Declarative Stack for AI-Driven Scientific Discovery (2026). arXiv:2605.04375. URL: <https://arxiv.org/pdf/2605.04375> (Truy cập ngày 19 tháng 8 năm 2026)

18. [A] LAP: An Agent-to-Instrument Protocol for Autonomous Science (2026). arXiv:2606.03755. URL: <https://arxiv.org/pdf/2606.03755> (Truy cập ngày 19 tháng 8 năm 2026)

19. [B] PubMed (2023). Autonomous chemical research with large language models. PMID: 38123806. URL: <https://pubmed.ncbi.nlm.nih.gov/38123806/>

20. [B] Carnegie Mellon University College of Engineering (2026). Gabe Gomes faculty profile. Carnegie Mellon University. URL: <https://engineering.cmu.edu/directory/bios/gomes-gabe.html> (Truy cập ngày 19 tháng 8 năm 2026)

21. [A] Wijaya, E. (2026). Beyond SMILES: Evaluating Agentic Systems for Drug Discovery. arXiv. URL: <https://arxiv.org/abs/2602.10163> (Truy cập ngày 19 tháng 8 năm 2026)

22. [A] Autonomous Chemistry and Materials Innovation Driven by Scientific Agents (2026). JACS Au, 6(5), 2659. URL: <https://pubs.acs.org/jaaucr/article/6/5/2659/5081217/Autonomous-Chemistry-and-Materials-Innovation> (Truy cập ngày 19 tháng 8 năm 2026)

23. [C] Self-driving labs are changing how chemists work (2026). C&EN. URL: <https://cen.acs.org/physical-chemistry/computational-chemistry/Self-driving-labs-changing-chemists/104/web/2026/06> (Truy cập ngày 19 tháng 8 năm 2026)

Chương 10 Sam Rodriques 1. [A] Skarlinski, M. D., Cox, S., Laurent, J. M., Braza, J. D., Hinks, M., Hammerling, M. J., Ponnampati, M., Rodriques, S. G., & White, A. D. (2024). Language agents achieve superhuman synthesis of scientific knowledge. arXiv preprint arXiv:2409.13740. URL: <https://arxiv.org/abs/2409.13740>

2. [A] Narayanan, S., Braza, J. D., Griffiths, R. R., Ponnampati, M., Bou, A., Laurent, J., Kabeli, O., Wellawatte, G., Cox, S., Rodriques, S. G., & White, A. D. (2024). Aviary: training language agents on challenging scientific tasks. arXiv preprint arXiv:2412.21154. URL: <https://arxiv.org/abs/2412.21154>
3. [A] Laurent, J. M., Janizek, J. D., Ruzo, M., Hinks, M. M., Hammerling, M. J., Narayanan, S., Ponnampati, M., White, A. D., & Rodriques, S. G. (2024). LAB-Bench: Measuring capabilities of language models for biology research. arXiv preprint arXiv:2407.10362. URL: <https://arxiv.org/abs/2407.10362>
4. [A] Lála, J., O'Donoghue, O., Shtedritski, A., Cox, S., Rodriques, S. G., & White, A. D. (2023). Paperqa: Retrieval-augmented generative agent for scientific research. arXiv preprint arXiv:2312.07559. URL: <https://arxiv.org/abs/2312.07559>
5. [B] Rodriques, S. G. (2020). What we'll learn about the brain in the next century. TED Talks, Vancouver, BC.
6. [A] Rodriques, S. G. (2026). Sam Rodriques on Kosmos: The AI Agent Doing 6 Months of Science in 4 Hrs. The American AI Festival, San Francisco, CA.
7. [C] FutureHouse (2025). Announcing Edison Scientific. FutureHouse News. URL: <https://www.futurehouse.org/news/announcing-edison-scientific> (Truy cập ngày 19 tháng 8 năm 2026)
8. [C] Edison Scientific (2026). Building AI-native biopharma: announcing our partnership with Incyte Corporation and Kosmos for R&D teams. Edison Scientific News. URL: <https://www.edisonscientific.com/news/building-ai-native-biopharma-announcing-our-partnership-with-incyte-corporation-and-kosmos-for-r-d-teams> (Truy cập ngày 19 tháng 8 năm 2026)
9. [C] Edison Scientific (2026). Creating novel biotechs focused on population-scale diseases: announcing our partnership with Population Health Partners. Edison Scientific News. URL: <https://www.edisonscientific.com/news/creating-novel-biotechs-focused-on-population-scale-diseases-announcing-our-partnership-with-population-health-partners> (Truy cập ngày 19 tháng 8 năm 2026)
10. [C] Edison Scientific (2026). LABBench2: an improved benchmark for measuring AI in biology research. Edison Scientific News. URL: <https://www.edisonscientific.com/news/labbench2-an-improved-benchmark> (Truy cập ngày 19 tháng 8 năm 2026)
11. [A] Skarlinski, M.D., Cox, S., Laurent, J.M., Braza, J.D., Hinks, M., Hammerling, M.J., Ponnampati, M., Rodriques, S.G., White, A.D. (2024). PaperQA2: Language Models Achieve Superhuman Synthesis of Scientific Knowledge. arXiv:2409.13740. DOI: 10.48550/arXiv.2409.13740
12. [A] Rodriques, S.G., Stickels, R.R., Goeva, A., Martin, C.A., Murray, E., Vanderburg, C.R., Welch, J., Chen, L.M., Chen, F., Macosko, E.Z. (2019). Slide-seq: A scalable technology for measuring genome-wide expression at high spatial resolution. *Science*, 363(6434), 1463-1467. DOI: 10.1126/science.aaw1219
13. [A] Marblestone, A., Gamick, A., Wang, M., Fridman, J. (2025). Field Notes on Moving Focused Research Organizations Forward. *Issues in Science and Technology*, 42(1), 43-48. DOI: 10.58875/EPTV1109
14. [B] FutureHouse (2024). WikiCrow: Automated Biomedical Literature Review. FutureHouse Research Announcements. URL: <https://www.futurehouse.org/research-announcements/wikicrow> (truy cập ngày 19 tháng 8 năm 2026)
15. [B] FutureHouse (2026). Demonstrating end-to-end scientific discovery with Robin: a multi-agent system. FutureHouse Research. URL:

<https://www.futurehouse.org/research/demonstrating-end-to-end-scientific-discovery-with-robin-a-multi-agent-system> (truy cập ngày 19 tháng 8 năm 2026)

16. [A] FutureHouse (2026). DISCO: Inventing Enzymes for Chemistry that Nature Never Explored. FutureHouse Research. URL: <https://www.futurehouse.org/research/disco-enzymes> (truy cập ngày 19 tháng 8 năm 2026)

17. [B] Edison Scientific (2026). Edison Scientific partnership announcements with Incyte Corporation and Population Health Partners. Edison Scientific. URL: <https://www.edisonscientific.com/> (truy cập ngày 19 tháng 8 năm 2026)

18. [D] FutureHouse (2026). About FutureHouse. FutureHouse. URL: <https://www.futurehouse.org/> (truy cập ngày 19 tháng 8 năm 2026)

19. [C] FutureHouse (2025). PaperQA2 achieves state-of-the-art performance in scientific literature search. FutureHouse News. URL: <https://www.futurehouse.org/news> (truy cập ngày 19 tháng 8 năm 2026)

20. [C] FutureHouse (2025). FutureHouse Platform launch. FutureHouse News. URL: <https://www.futurehouse.org/news> (truy cập ngày 19 tháng 8 năm 2026)

21. [A] Wadden, D., Lin, S., Lo, K., Wang, L. L., van Zuylen, M., Cohan, A., & Hajishirzi, H. (2020). Fact or fiction: Verifying scientific claims. arXiv preprint arXiv:2004.14974.

22. [A] Schlichtkrull, M., Guo, Z., & Vlachos, A. (2024). AVeriTeC: A dataset for real-world claim verification with evidence from the web. *Advances in Neural Information Processing Systems*, Vol. 36.

23. [A] Skarlinski, M. D., Cox, S., Laurent, J. M., Braza, J. D., Hinks, M., Hammerling, M. J., Ponnampati, M., Rodrigues, S. G., & White, A. D. (2024). Language agents achieve superhuman synthesis of scientific knowledge. arXiv preprint arXiv:2409.13740, pp. 1-28.

24. [A] Narayanan, S., Braza, J. D., Griffiths, R. R., Ponnampati, M., Bou, A., Laurent, J., Kabeli, O., Wellawatte, G., Cox, S., Rodrigues, S. G., & White, A. D. (2024). Aviary: training language agents on challenging scientific tasks. arXiv preprint arXiv:2412.21154, pp. 1-20.

25. [A] Laurent, J. M., Janizek, J. D., Ruzo, M., Hinks, M. M., Hammerling, M. J., Narayanan, S., Ponnampati, M., White, A. D., & Rodrigues, S. G. (2024). LAB-Bench: Measuring capabilities of language models for biology research. arXiv preprint arXiv:2407.10362, pp. 1-25.

26. [A] Lála, J., O'Donoghue, O., Shtedritski, A., Cox, S., Rodrigues, S. G., & White, A. D. (2023). Paperqa: Retrieval-augmented generative agent for scientific research. arXiv preprint arXiv:2312.07559, pp. 1-15.

27. [A] Boiko, D. A., MacKnight, R., & Gomes, G. (2023). Autonomous chemical research with large language models. *Nature*, 624, 570-578. DOI: 10.1038/s41586-023-06792-0

28. [C] Edison Scientific (2025). Kosmos: An AI Scientist for Autonomous Discovery. Edison Scientific News. URL: <https://edisonscientific.com/news/announcing-kosmos> (truy cập ngày 19 tháng 8 năm 2026)

29. [C] Tech Funding News (2025). FutureHouse spinout Edison lands \$70M to build autonomous AI scientists for biology and chemistry. URL: <https://techfundingnews.com/edison-raises-70m-ai-scientist-platform/> (truy cập ngày 19 tháng 8 năm 2026)

30. [C] Genetic Engineering & Biotechnology News (2026). Google DeepMind and Edison Are Building the AI Scientist. GEN. URL: <https://www.genengnews.com/topics/artificial-intelligence/google-deepmind-and-e>

dison-are-building-the-ai-scientist/ (truy cập ngày 19 tháng 8 năm 2026)

31. [C] Endpoints News (2025). Research institute FutureHouse debuts AI scientist Kosmos and for-profit spinoff. URL: <https://endpoints.news/research-institute-futurehouse-debuts-ai-scientist-kosmos-and-for-profit-spinoff/> (truy cập ngày 19 tháng 8 năm 2026)

Chương 11 Hiroaki Kitano 1. [A] Kitano, H. (2002). Systems biology: a brief overview. *Science*, 295(5560), 1662-1664. DOI: 10.1126/science.1069492

2. [A] Kitano, H. (2002). Computational systems biology. *Nature*, 420(6912), 206-210. DOI: 10.1038/nature01254

3. [A] Kitano, H. (2007). Towards a theory of biological robustness. *Molecular Systems Biology*, 3, 137. DOI: 10.1038/msb4100179

4. [A] Le Novère, N. et al. (2009). The Systems Biology Graphical Notation. *Nature Biotechnology*, 27(8), 735-741. DOI: 10.1038/nbt.1558

5. [A] Ghosh, S., Matsuoka, Y., Asai, Y., Hsin, K. Y., & Kitano, H. (2011). Software for systems biology: from tools to integrated platforms. *Nature Reviews Genetics*, 12(12), 821-832. DOI: 10.1038/nrg3096

6. [A] Kitano, H. (2021). Nobel Turing Challenge: creating the engine for scientific discovery. *npj Systems Biology and Applications*, 7, 29. DOI: 10.1038/s41540-021-00189-3

7. [A] Palaniappan, S. K., Zolfo, M., Kottapalli, S. et al. (2026). Creating an engine of scientific discovery, an Indo-Japan perspective: learnings from IJNTC 2025 workshop. *npj Systems Biology and Applications*, 12, 92. DOI: 10.1038/s41540-026-00714-2 (truy cập ngày 19 tháng 8 năm 2026)

8. [A] Zenil, H., Tegnér, J. et al. (2026). The future of fundamental science led by generative closed-loop artificial intelligence. *Frontiers in Artificial Intelligence*, 9, 1678539. DOI: 10.3389/frai.2026.1678539 (truy cập ngày 19 tháng 8 năm 2026)

9. [B] Trang web chính thức của Nobel Turing Challenge. Workshops. URL: <https://www.nobelturingchallenge.org/> (truy cập ngày 19 tháng 8 năm 2026)

10. [A] Imai, S., & Kitano, H. (1998). A computational model of cellular senescence and the role of heterochromatin state transitions. *Science*, 281(5382), 1152-1159.

11. [A] Musslick, S., Bartlett, L. K., Chandramouli, S. H., Dubova, M., Gobet, F., Griffiths, T. L., Hullman, J., King, R. D., Kutz, J. N., Lucas, C. G., Mahesh, S., Pestilli, F., Sloman, J. S., & Holmes, W. R. (2025). Automating the practice of science: Opportunities, challenges, and implications. *Proceedings of the National Academy of Sciences*, 122(5), e2401238121, 1-18.

12. [C] Scientific American (2026). Robots run the experiments. Grad students fix them from bed. *Scientific American*. URL: [13. \[A\] HPCwire \(2026\). PNNL Aims to Bring Autonomous Science to Every Lab Experiment Within 5 Years. *HPCwire*. URL: <https://news.google.com/rss/articles/CBMitAFBVV95cUxQSFBUZ29QMzFyZkFOalZvTmU1 WExqYUNxWD hNcFhQUkRlVFBienFtdU55UF11TlV5ZjhqTDc0SXBYeDdfS0liWXNKVnhp d nJDR0JOZTk ydlpM dUI3NXhT MjNPd1FidmFEQWdHSV9nYk5UOE1FQmFwX29WY2xMb mpQTUVSTEpwQjVXQzF3YkN2NXQzcy1>](https://news.google.com/rss/articles/CBMinAFBVV95cUxPVGUtbWRwMVJPdlcwa1ZNcW dyTGR0Rk j2dnZ icWVnM21CY3JjR3lmWkRkT2JWdEJtWHBkb19zb3hKTWExeTEy d05PMG9zd284X3YteHJT MjEyU282MX pwVTN5cVN6MzllOGs4d0tBUFFjX0xsaGdXSl8tTzljMkNMb3hLdy01Z1JpS2JqMFpWYW02bElO R3BZTGhvRHA (truy cập ngày 19 tháng 8 năm 2026)</p></div><div data-bbox=)

rcl9jUzVCULRnQ19XM2FEYjNwa1RPLVctOTBya2xnanM (truy cập ngày 19 tháng 8 năm 2026)

14. [C] Newswise (2026). Inside PNNL's Lab-Wide Push for Trustworthy Autonomous Science. Newswise. URL: <https://news.google.com/rss/articles/CBMinwFBVV95cUxPM1Fxbk5BMDkzYllZNDY0TVR5OHhsVkYydDhId3RRQ29MV3ZlQmFLS3NvNjVpRUw2bnE0T2FBTUlXNVJrdGtlVi1zSXA1NWc2RmtfcUFZNFpiNi1KeFFWRDg5azRhCtdmQ1U3UHBBcnRhR0tSLVlYbURzRWVQdFhmYlVtcGVzN3hqWjNxU2ltd2liSlVzcURFTGRMdEJlX3M> (truy cập ngày 19 tháng 8 năm 2026)

15. [C] EurekAlert! (2026). Operations workforce powers ORNL's autonomous science future. EurekAlert!. URL: <https://news.google.com/rss/articles/CBMiXEFVX3lxTE1SODhld3gxQ3NUR1RUSkhuZ2Y2MF9MZy1yOUczUVNzeGhvQzljJa3BmUXBpMlJSSFBDR0g4T0wycjBPTk55WlFnb1hCRDYyYmlxelNFS0FDTlNyQ2o0> (truy cập ngày 19 tháng 8 năm 2026)

16. [B] Observer (2025). This Scientist Thinks an A.I. Could Win a Nobel Prize by 2050. Observer. URL: <https://news.google.com/rss/articles/CBMiekFVX3lxTFBmOEZjQzZXcUE5SGJ3TkpKcHd5U2l1WGMybU9WSEhsdXZtdGpLQ3gwV1V5YU8wTFFlZTA2aWI2SXR6aU1ja0paLVpIN3gyZHhDZ09LMnRTLUIKVL8xbGpDOWpmQ3VCaDBm2c0b0FBN2NqbXRZRfVsT0lR> (truy cập ngày 19 tháng 8 năm 2026)

17. [C] Semafor (2025). AI may soon make Nobel-level discovery, scientists predict. Semafor. URL: <https://news.google.com/rss/articles/CBMiogFBVV95cUxOSlVoeTcwVE5QaFpROFpMbGJVMVQ5bnlUR041NmZOU1FHNW53U1N3OGlkYjk4dWlVMM5sem5kdzh2RXNXZDBJaWdKUnZsOEExwZVltTmFhaGZ2R09YVXRvamJ3d1FjNG1yemJRNE1sVU9NdEpnS2FvczdQZjZ2NzNxcDJJNTRIYUhtMXNsYkcxYUxqOWVrQ0ItRuk5bWp0Y3AtVKE> (truy cập ngày 19 tháng 8 năm 2026)

18. [C] El País English (2025). When your robot wins a Nobel Prize. El País. URL: <https://news.google.com/rss/articles/CBMilAFBVV95cUxNdKZoV3NYdmlWRmpRRnRkTmdWc3dnLTdoYnZZbXg0YnpMNU4xT2pYX3R5NS1WRGVsejU3VzJDSkEzU2hqWVkwOU1ndEkzR0l4cXdwck5xX2pTVjNkV2hTTHRERDRnWDRPV0ExYmtMVV>

RDYV9iZDJ0em9DQl9Vc0FxmUNNUNNWlcwVlBpdmY4MTZITWp6 (truy cập ngày 19 tháng 8 năm 2026)

19. [A] Turing, A. M. (1950). Computing Machinery and Intelligence. *Mind*, 59(236), 433-460. DOI: 10.1093/mind/LIX.236.433

20. [A] King, R. D., Whelan, K. E., Jones, F. M. et al. (2004). Functional genomic hypothesis generation and experimentation by a robot scientist. *Nature*, 427(6971), 247-252. DOI: 10.1038/nature02236

21. [A] King, R. D., Rowland, J., Oliver, S. G. et al. (2009). The Automation of Science. *Science*, 324(5923), 85-89. DOI: 10.1126/science.1165620

22. [A] King, R. D., Rowland, J., Aubrey, W. et al. (2009). The Robot Scientist Adam. *Computer*, 42(7), 46-54. DOI: 10.1109/MC.2009.270

23. [A] King, R. D. (2009). Make Way for Robot Scientists. *Science*, 325(5943), 945. DOI: 10.1126/science.325_945a

24. [A] Williams, K., Bilsland, E., Sparkes, A. et al. (2015). Cheaper faster drug development validated by the repositioning of drugs against neglected tropical diseases. *Journal of the Royal Society Interface*, 12(104), 20141289. DOI: 10.1098/rsif.2014.1289

25. [C] Sisson, P. (2026). The lab never sleeps. Can the science keep up? *Scientific American*. URL: <https://www.scientificamerican.com/article/autonomous-labs-are-running-science-experiments-24-7/>

(truy cập ngày 19 tháng 8 năm 2026)

26. [C] Iowa Public Radio, NPR (2026). Scientists in 'autonomous laboratories' are starting to outsource work to robots. URL:

<https://www.iowapublicradio.org/news-from-npr/2026-06-05/scientists-in-autonomous-laboratories-are-starting-to-outsource-work-to-robots> (truy cập ngày 19 tháng 8 năm 2026)

27. [C] KBTX News 3 (2026). Focus at Four: A&M lands nearly \$25M grant to build first-of-its-kind AI metal discovery lab. URL:

<https://www.kbtx.com/2026/08/08/focus-four-am-lands-nearly-25m-grant-build-first-of-its-kind-ai-metal-discovery-lab/> (truy cập ngày 19 tháng 8 năm 2026)

28. [A] Turing, A.M. (1950). Computing Machinery and Intelligence. *Mind*, 59(236), 433-460.

29. [B] Lindsay, R.K., Buchanan, B.G., Feigenbaum, E.A., Lederberg, J. (1980). Applications of Artificial Intelligence for Organic Chemistry: The DENDRAL Project. McGraw-Hill.

30. [B] Langley, P., Simon, H.A., Bradshaw, G.L., Zytkow, J.M. (1987). Scientific Discovery: Computational Explorations of the Creative Process. MIT Press.

31. [A] King, R.D. et al. (2004). Functional genomic hypothesis generation and experimentation by a robot scientist. *Nature*, 427(6971), 247-252.

32. [A] King, R.D. et al. (2009). The Automation of Science. *Science*, 324(5923), 85-89.

33. [A] Kitano, H. (2016). Artificial Intelligence to Win the Nobel Prize and Beyond: Creating the Engine for Scientific Discovery. *AI Magazine*, 37(1), 39-49.

34. [A] Boiko, D.A., MacKnight, R., Kline, B., Gomes, G. (2023). Autonomous chemical research with large language models. *Nature*, 624, 570-578.

35. [A] Mitchener, L. et al. (2025). Kosmos: An AI Scientist for Autonomous Discovery. arXiv. URL: <https://arxiv.org/abs/2511.02824> (truy cập ngày 19 tháng 8 năm 2026)

36. [A] Rodriques, S. et al. (2026). Demonstrating end-to-end scientific discovery with Robin: a multi-agent system. *Nature*. URL: <https://www.nature.com/articles/s41586-026-10652-y> (truy cập ngày 19 tháng 8 năm 2026)

37. [A] Channing, G., Ghosh, A. (2026). AI for Scientific Discovery is a Social Problem. arXiv. URL: <https://arxiv.org/abs/2509.06580> (truy cập ngày 19 tháng 8 năm 2026)

Chương 12 Ai là chủ nhân của phát hiện? 1. [A] Musslick, S. et al. (2025). Automating the practice of science: Opportunities, challenges, and implications. *Proceedings of the National Academy of Sciences*, 122(5), e2401238121. DOI: 10.1073/pnas.2401238121

2. [A] Boiko, D. A., MacKnight, R., Kline, B., & Gomes, G. (2023). Autonomous chemical research with large language models. *Nature*, 624, 570-578. DOI: 10.1038/s41586-023-06792-0

3. [A] Yamada, Y. et al. (2026). Towards end-to-end automation of AI research. *Nature*. (Sakana AI, xuất bản tháng 3 năm 2026)

4. [A] Sakana AI (2026). The AI Scientist: peer-reviewed publication in *Nature*. URL: <https://sakana.ai/ai-scientist-nature/> (truy cập ngày 19 tháng 8 năm 2026)

5. [A] Hochreiter, S., & Schmidhuber, J. (1997). Long short-term memory. *Neural Computation*, 9(8), 1735-1780.

6. [B] Thaler v. Comptroller-General of Patents (2023). UK Supreme Court, [2023] UKSC 49. (Phán quyết bác bỏ tư cách nhà phát minh của trí tuệ nhân tạo)

7. [A] Watson, J. L. et al. (2023). De novo design of protein structure and function with RFdiffusion. Nature, 620, 1089-1100. DOI: 10.1038/s41586-023-06415-8

Kết luận 1. [A] Jumper, J. et al. (2021). Highly accurate protein structure prediction with AlphaFold. Nature, 596(7873), 583-589. DOI: 10.1038/s41586-021-03819-2

2. [A] Schmidt, M., & Lipson, H. (2009). Distilling free-form natural laws from experimental data. Science, 324(5923), 81-85. DOI: 10.1126/science.1165893

3. [A] Musslick, S. et al. (2025). Automating the practice of science: Opportunities, challenges, and implications. Proceedings of the National Academy of Sciences, 122(5), e2401238121. DOI: 10.1073/pnas.2401238121

4. [A] Kitano, H. (2021). Nobel Turing Challenge: Creating the engine for scientific discovery. npj Systems Biology and Applications, 7, Article 29. DOI: 10.1038/s41540-021-00189-3

5. [A] Gottweis, J. et al. (2026). Accelerating scientific discovery with Co-Scientist. Nature, 655(8122), 487-496. DOI: 10.1038/s41586-026-10644-y

6. [C] Báo Điện tử ETNews (2026). KISTI giới thiệu mô hình AI chuyên biệt cho khoa học công nghệ và nền tảng nghiên cứu tự hành tại hội nghị học thuật hàng đầu thế giới 'KDD 2026'. URL: <https://www.etnews.com/20260814000161> (truy cập ngày 19 tháng 8 năm 2026)

7. [C] Báo Herald Kinh tế (2026). 'Đưa ra cả các giả thuyết nghiên cứu mới' KISTI lần đầu tiên công bố 'Nhà khoa học AI'. URL: <https://biz.heraldcorp.com/article/10841497?sec=005> (truy cập ngày 19 tháng 8 năm 2026)

Chỉ mục

Nền tảng Garuda Chương 11 Edison Scientific Chương 10 Tế bào ảo Chương 1 Phê phán chủ nghĩa dự đoán Chương 1 Học tăng cường Chương 1, Giải thích thuật ngữ Eureka Lời mở đầu, Chương 6 Gomes, Gabriel Chương 9 Eve (robot) Chương 8 Grokking Chương 7 Phòng thí nghiệm tự hành Chương 5 Thử thách Nobel Turing Chương 8, 11 Khả năng tái lập Lời mở đầu, Chương 8, Giải thích thuật ngữ Học chủ động Chương 8, Giải thích thuật ngữ Jumper, John Chương 2 DENDRAL Lời mở đầu, Chương 8, Niên biểu Genesis (robot) Chương 8 Giáo dục kiểu học việc Chương 12 Trách nhiệm (pháp lý) Chương 12, Kết luận Rodrigues, Sam Chương 10 Kavukcuoglu, Koray Chương 1 RoboCup Chương 11 Chemprop Chương 3 Rosetta Chương 4 Co-Scientist (Gomes) Chương 9 Lipson, Hod Chương 6 Co-Scientist (DeepMind) Lời mở đầu, Kết luận Sốt rét Chương 3, 8 Kosmos Chương 10 Barzilay, Regina Chương 3 Kitano, Hiroaki Chương 11 Nhà phát minh Chương 12 Tegmark, Max Chương 7 Baker, David Chương 4 Triclosan Chương 8 BACON (chương trình) Lời mở đầu Tối ưu hóa Pareto Chương 6, Giải thích thuật ngữ Simon, Herbert Lời mở đầu PaperQA2 Chương 10 Sakana AI Chương 12 Vòng lặp khép kín Lời mở đầu, Giải thích thuật ngữ Thiết kế (protein) Phần 2, Chương 4 Halicin Chương 3 Phản ứng Suzuki-Miyaura Chương 9 Hassabis, Demis Chương 1 Adam (robot) Lời mở đầu, Chương 8 AI Feynman Chương 7, Ghi chú kỹ thuật Aspuru-Guzik, Alán Chương 5 RFdiffusion Chương 4 AlphaFold Chương 1, 2 Bảng sáng chế Chương 12

ISBN |

Tác giả | Kim Kyung-jin

Người xuất bản | Kim Kyung-jin

Nơi xuất bản | Nhà xuất bản Luật sư Kim Kyung-jin

Đăng ký nhà xuất bản | 10. 3. 2025 (Số 2025-000015)

Địa chỉ | Phòng 304, Tòa nhà Baek-il, 91 Jeonnong-ro, Dongdaemun-gu, Seoul

Điện thoại | 02-6338-1905

Email | kimkj008@gmail.com

Giá : 20.000 won

© Kim Kyung-jin 2026

Cuốn sách này là tài sản trí tuệ của tác giả, nghiêm cấm mọi hình thức sao chép và tái bản khi chưa được phép.

Ghi chú) Một số nội dung và quá trình biên tập của cuốn sách này được hỗ trợ bởi các công cụ trí tuệ nhân tạo.

Nếu quý độc giả thấy cuốn sách hữu ích và tiếp thu được những tri thức mới có giá trị, xin vui lòng ủng hộ tự nguyện tới tài khoản NongHyup 302-1096-0948-81 (Chủ tài khoản Kim Kyung-jin).