



Tương lai AI giữa Ấn Độ với Trung Quốc và Hoa Kỳ

ISSN: 2734-9195

16:29 01/04/2025

Giải pháp là Ấn Độ nên hợp tác với Liên minh châu Âu, Brazil, Úc và các nước láng giềng Nam Á.

Năm 2011 - sách *"Sapiens: Lược sử loài người"* (Sapiens: A Brief History of Humankind) được xuất bản bằng tiếng Do Thái đã đưa danh tiếng Thiên giả Yuval Noah Harari lên tầm cao mới, trước đó thì ông được biết đến là học giả về lịch sử quân sự thời trung cổ.

Các tác phẩm của ông trải dài từ việc phân tích thế giới đen tối của các vụ ám sát, bắt cóc, phản quốc và phá hoại vào thế kỷ XII *"Quá trình hoạt động đặc biệt trong thời đại hiệp sĩ"* (Special Operations in the Age of Chivalry), đến nghiên cứu văn hóa về trải nghiệm của những người lính *"Hồi ký quân sự thời Phục hưng: Chiến tranh, lịch sử và bản sắc"* (Renaissance Military Memoirs: War, History and Identity).

Khi cuốn sách *"Sapiens: Lược sử loài người"* xuất hiện đưa ông lên vị thế của một nhà triết học - sử gia liên ngành.

Trong thập kỷ qua, Thiên giả Yuval Noah Harari đã xuất bản các tác phẩm khơi gợi suy nghĩ khám phá quá khứ và tương lai của nhân loại.



Cuốn sách gần đây, *“Nexus - Lịch sử của những mạng lưới thông tin từ thời đại Đồ đá đến trí tuệ nhân tạo”* (Nexus: A Brief History of Information Networks from the Stone Age to AI), đánh dấu sự trở lại một phần với nguồn gốc lịch sử quân sự của ông. Ông cho biết sự trỗi dậy của trí tuệ nhân tạo - hiện do Hoa Kỳ và Trung Quốc thống trị - có thể dẫn đến việc các quốc gia tạo ra các phạm vi kỹ thuật số riêng: *“Mỗi phạm vi chịu ảnh hưởng của các truyền thống chính trị, văn hóa và tôn giáo khác nhau”*. *“Thay vì bị chia cắt giữa hai đế chế toàn cầu, thế giới có thể bị chia cắt giữa hàng chục đế chế”*. Ông viết: *“Các đế chế mới cạnh tranh với nhau càng nhiều thì nguy cơ xung đột vũ trang càng lớn”*.

Yuval Noah Harari bắt đầu chuyến lưu diễn vòng quanh thế giới, cảnh báo rằng AI có thể khiến thế giới *“giống Kafka hơn cả Terminator”*. Tại Mumbai, Ấn Độ, ông đã trả lời phỏng vấn với tạp chí The Week về tác động của AI đối với Ấn Độ và những lựa chọn để ứng phó với sự trỗi dậy của AI:

Mumbai là một trong những thành phố giàu có nhất thế giới, với mật độ tập trung cao nhất của những cá nhân có giá trị tài sản ròng cao - những người có tài sản trên 1 triệu USD. Tuy nhiên, hơn 40% dân số của thành phố sống trong các khu ổ chuột. Ông có nghĩ AI sẽ làm nghiêm trọng thêm tình trạng bất bình đẳng này?

Tùy thuộc vào cách chúng ta sử dụng, tương lai không thể định trước - không giống như chỉ có một tương lai cho AI. Một con dao có thể được sử dụng để giết ai đó, cũng có thể được sử dụng trong phẫu thuật, hoặc được dùng cắt salad cho bữa tối. Con dao không cho tôi biết trước phải làm gì. AI cũng tương tự như vậy. AI có thể làm tăng hoặc giảm bất bình đẳng và nghèo đói.

Mối nguy hiểm lớn nhất là AI có thể làm gia tăng bất bình đẳng cả trong và giữa các quốc gia. Cuộc Cách mạng công nghiệp thế kỉ XVIII - XIX là quá trình chuyển biến từ nền sản xuất nhỏ thủ công sang sản xuất lớn bằng máy móc - phát minh ra động cơ hơi nước và điện - ban đầu đã làm gia tăng bất bình đẳng toàn cầu. Các quốc gia Anh, Pháp, Hoa Kỳ và sau đó là Nhật Bản đã dẫn đầu cuộc cách mạng, sử dụng công nghệ vượt trội của họ để chinh phục và khai thác thế giới. Điều tương tự cũng có thể xảy ra với AI. Một vài quốc gia, rất hùng mạnh và giàu có, đang dẫn đầu cuộc cách mạng AI. Họ có thể thống trị và khai thác các nước khác.

Với AI, mọi quân đội đều trở nên lỗi thời. Khi đó, có thể không còn vũ khí cơ học, thay vào đó là vũ khí hạt nhân thông minh, máy bay không người lái, tàu, xe tăng và máy bay hoàn toàn tự động. Trong đời sống, AI dễ dàng thay thế tài xế, công nhân dệt may, nhân viên ngân hàng và thậm chí là nhà văn, nhà nghiên cứu... AI sẽ làm cho các công việc, nghề nghiệp truyền thống dần biến mất. Đây thực sự là mối nguy hiểm và đe dọa.

Chúng ta có thể đảm bảo rằng lợi ích của AI được phân phối bình đẳng?

Các lợi ích phải được phân bổ công bằng hơn - nếu không muốn nói là bình đẳng - giữa các quốc gia, ngành nghề và tầng lớp xã hội. Những công việc cũ biến mất, những công việc mới sẽ xuất hiện. Câu hỏi đặt ra, bạn có đủ tiền để đào tạo lại mọi người cho những công việc này không?

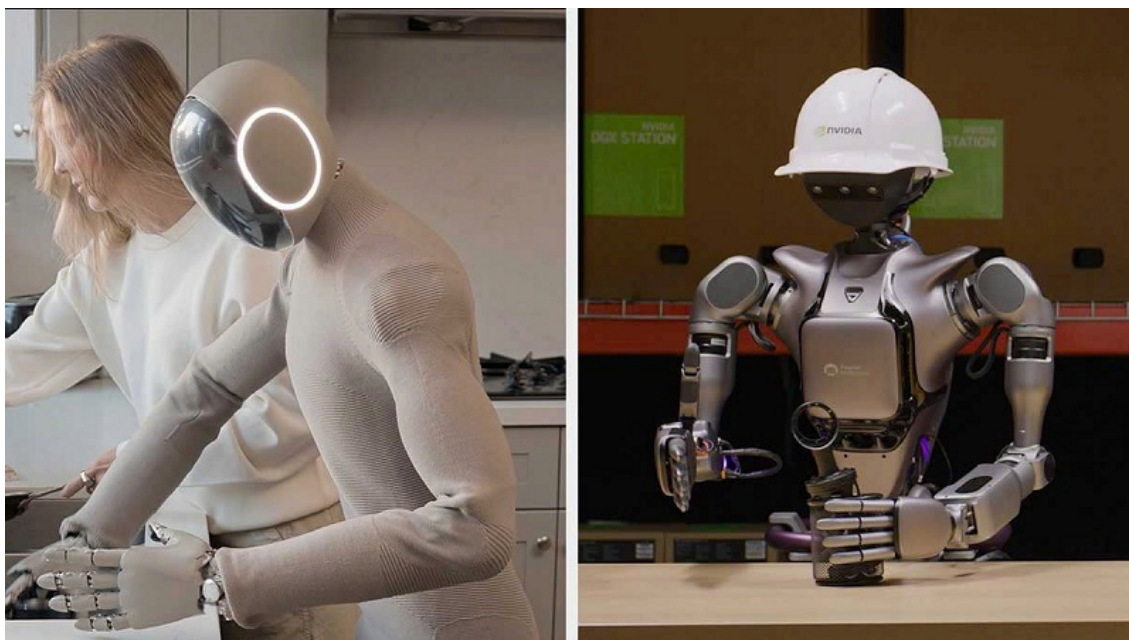
Dự đoán công việc nào sẽ biến mất không phải là điều dễ dàng. Các công việc lương cao có thể là công việc dễ tự động hóa nhất. Ví dụ, ngân hàng là lĩnh vực sinh lợi rất lớn, tuy nhiên tất cả chỉ là thông tin và những tham số dữ liệu: “Ồ, cái này đang tăng, và cái này đang giảm. Tốt hơn là tôi nên đầu tư vào công ty này”. AI có thể tự động hóa việc này một cách dễ dàng. Nhưng, hãy nghĩ về một cô y tá phải tiêm thuốc cho trẻ em hoặc phải băng bó cho người bị thương. Cô ấy cần trí tuệ và các kỹ năng chuyên môn cần thiết. Cô ấy cần thao tác để thay băng một cách nhẹ nhàng, khiến bệnh ít đau hơn. Việc này khó tự động hóa hơn nhiều.

Vậy những công việc lương thấp sẽ không biến mất trước?

Một số sẽ sớm bị đào thải. Các công việc mà con người dễ làm, sẽ dễ dàng bị robot thay thế.

Những biện pháp cụ thể để ngăn chặn AI làm trầm trọng thêm tình trạng bất bình đẳng?

Có hai biện pháp chính.



Một là giáo dục: Cung cấp cho mọi người một nền giáo dục rộng, chứ không phải một nền giáo dục hẹp, là điều tối quan trọng. Đào tạo mọi người một kỹ năng duy nhất, hoặc một tập hợp các kỹ năng rất hẹp, là rủi ro. Nếu bạn đào tạo ai đó chỉ để viết mã, họ sẽ không có việc làm trong 10 hoặc 20 năm, khi AI viết tất cả các mã.

Vì thế, điều quan trọng là phải có một chương trình rộng giúp mọi người phát triển trí tuệ, cùng với các kỹ năng xã hội, cảm xúc và vận động của họ. Với một bộ kỹ năng phổ quát, sẽ dễ dàng hơn để đối phó với thị trường việc làm thay đổi.

Kỹ năng quan trọng nhất, tất nhiên, là khả năng học hỏi liên tục trong suốt cuộc đời. Và đó là điều rất khó để hướng dẫn. Mọi người cần có tư duy linh hoạt để có thể tiếp tục học hỏi và thay đổi, vì thị trường việc làm sẽ liên tục thay đổi.

Thứ hai, hợp tác: Cách duy nhất để hầu hết các quốc gia duy trì khả năng cạnh tranh trong cuộc đua AI là hợp tác với các quốc gia khác. Bởi vì có hai siêu cường AI - Hoa Kỳ và Trung Quốc đang đi trước rất xa. Họ có thể tập trung rất nhiều nguồn lực vào nghiên cứu AI. Ngay cả một quốc gia lớn như Ấn Độ cũng không có đủ nguồn lực để tự mình cạnh tranh với họ. Các quốc gia nhỏ hơn như Sri Lanka hoặc Bangladesh gần như không có cơ hội.

Vậy, làm sao có thể ngăn chặn sự lặp lại của cuộc cách mạng công nghiệp trước đó, khi một số ít cường quốc chế ngự tất cả mọi người? Giải pháp là hợp tác - tập hợp nguồn lực để đầu tư vào nghiên cứu và đạt được thỏa thuận về quy định và sử dụng AI.

Ấn Độ nên hợp tác với Liên minh châu Âu, Brazil, Úc và các nước láng giềng Nam Á.

Một nghiên cứu gần đây của Quỹ Nghiên cứu Người Quan sát (Observer Research Foundation) của Mumbai cho thấy Ấn Độ và Israel - cả hai đều là nền dân chủ có hệ sinh thái khởi nghiệp phát triển mạnh - nên hợp tác để phát triển một mô hình quản lý AI mới phù hợp với nhu cầu của các nước đang phát triển. 'Mô hình Mumbai-Tel-Aviv' được đề xuất này nhằm mục đích tạo ra sự cân bằng giữa hai cách tiếp cận quản lý AI chủ đạo: sự can thiệp tối thiểu của Hoa Kỳ và các quy định nghiêm ngặt của EU. Liệu một mô hình như vậy có hiệu quả không?

Ở cấp độ sâu hơn, các nghiên cứu như thế này phản ánh nhu cầu mở rộng cuộc thảo luận về AI. Hiện tại, các mô hình AI hàng đầu và các quyết định quan trọng nhất liên quan đến AI được đưa ra ở Hoa Kỳ và Trung Quốc. Chúng ta cần nhiều người hơn ở Israel, Ấn Độ và các quốc gia khác hiểu được những gì đang diễn ra, tham gia cuộc trò chuyện và tác động đến các quyết định. Đó là một trong những lý do tôi viết Nexus - để thông báo cho nhiều người hơn: *'Hãy xem điều này đang diễn ra và điều này sẽ tác động đến tất cả mọi người'.*

Bỏ qua AI cũng giống như bỏ qua việc phát minh ra tàu hỏa vào thế kỷ XIX. Mọi người nói rằng, *'Chúng tôi không quan tâm đến tàu hỏa. Chúng tôi có những vấn đề khác ở Ấn Độ.'* Và sau đó người Anh đã chinh phục Ấn Độ, vì họ có công nghệ mà bạn không có.

Vì vậy, nếu bạn nghĩ rằng Ấn Độ có nhiều vấn đề cấp bách hơn nhiều và để AI cho người Trung Quốc và người Mỹ, bạn sẽ bị họ thống trị trong 10 năm nữa.

Phát triển công nghệ có thể là giải pháp?

Bạn không bao giờ có thể tin tưởng vào công nghệ. Giải pháp phải là thể chế, không phải công nghệ. Bạn cần các thể chế kiểm tra tính công bằng của công nghệ. Suy nghĩ rằng chúng ta không cần các thể chế của con người, vì chúng quan liêu và phức tạp, rằng chúng ta có thể phát triển một công nghệ kỳ diệu có thể giải quyết mọi vấn đề, là một điều viển vông hão huyền. Điều đó không bao giờ hiệu quả, vì công nghệ không bao giờ thoát khỏi sự thiên vị. Không có công nghệ nào hoàn toàn không thiên vị.

Cũng giống như tòa án đảm bảo rằng hành động của con người không thiên vị, chúng ta cũng cần các tổ chức có thể kiểm tra AI và đảm bảo rằng AI không thiên vị - chống lại phụ nữ, đẳng cấp, tôn giáo. Công nghệ rất quan trọng, nhưng không bao giờ là giải pháp tự thân. Tôi quay lại ví dụ về con dao: bạn

không thể có một con dao chỉ có thể cắt salad mà không bao giờ làm hại con người.

Hiện tại AI “đang cắt salad” cho chúng ta hay gây hại? Ông thấy vai trò của AI sẽ phát triển như thế nào trong tương lai gần?

AI được sử dụng rất tốt trong nhiều lĩnh vực. Ví dụ, trong y học, AI được sử dụng để xử lý và phân tích hàng triệu dữ liệu y tế từ lịch sử bệnh nhân, hóa đơn, đến quản lý dược phẩm, AI giúp giải quyết tình trạng thiếu bác sĩ trên toàn thế giới. AI có thể chẩn đoán bệnh nhanh hơn và chính xác hơn con người.

Chúng ta thấy AI được sử dụng trong giao thông. Hàng năm, hơn một triệu người tử vong trong các vụ tai nạn xe hơi, chủ yếu là do lỗi của con người, như lái xe khi say rượu. Các phương tiện tự lái được hỗ trợ bởi AI có thể giảm đáng kể tình trạng này. Mặc dù tai nạn vẫn xảy ra, nhưng chúng xảy ra ít thường xuyên hơn nhiều với những chiếc xe do AI điều khiển. Các phương tiện tự hành an toàn hơn nhiều so với những phương tiện do con người điều khiển.

AI đóng góp cho nhân loại trong những lĩnh vực như thế này, nhưng bạn cũng thấy những diễn biến rất nguy hiểm. Mạng xã hội, do AI điều hành, không chỉ lan truyền những lời nói dối và tin tức giả mạo, mà còn cả nỗi sợ hãi, lòng căm thù và sự tức giận. Nội dung đang gây ra những xáo trộn xã hội đang làm suy yếu nền dân chủ. Vấn đề là các thuật toán có xu hướng thích nội dung xấu, vì mục tiêu của chúng chỉ đơn giản là tăng sự tương tác của người dùng - khiến mọi người dành nhiều thời gian hơn trên nền tảng này.

Sau đó, chúng ta có ứng dụng trí tuệ nhân tạo (AI) trong lĩnh vực quân sự. Trong các cuộc chiến như ở Trung Đông, AI đã quyết định mục tiêu ném bom. Khi người Israel ném bom Gaza, AI, chứ không phải con người, ngày càng nói với họ rằng, ‘Ồ, *bạn nên ném bom tòa nhà này và người kia.*’

Thật là một canh bạc lớn khi trao quyền cho AI để đưa ra quyết định quân sự. Mối nguy hiểm là một cuộc chạy đua vũ trang - nơi con người buộc phải trao ngày càng nhiều quyền kiểm soát quân sự cho AI để giành chiến thắng trong các trận chiến.

Điều tương tự cũng có thể xảy ra trong lĩnh vực tài chính, khi ngày càng có nhiều thẩm quyền được trao cho AI để quyết định đầu tư vào đâu và cho ai vay. Với việc chính phủ Donald Trump tại Hoa Kỳ phản đối bất kỳ quy định nào về AI, tài chính có thể chứng kiến việc sử dụng AI không được kiểm soát. Sẽ thế nào nếu AI phát minh ra công cụ tài chính mới mà con người không thể hiểu được và không được kiểm soát trong một vài năm? Nó có thể dẫn đến khủng hoảng tài

chính lớn tương tự như năm 2008, khi không ai hiểu chuyện gì đang xảy ra vì những công cụ mà AI phát minh ra quá phức tạp đối với bộ não con người. Chính phủ sẽ làm gì nếu họ mất quyền kiểm soát hệ thống tài chính?

Con người có thể tự mình thực hiện quy định nếu chúng ta đã đạt đến điểm này không?

Con người với sự trợ giúp của AI. Một lần nữa, diễn tiến đang trở nên phức tạp đến mức bạn cần một AI để hiểu AI kia, nhưng bạn vẫn cần con người ở đó để hướng dẫn. Nhưng liệu AI, vì lợi ích của chính nó, có sẵn sàng cung cấp cho con người một khuôn khổ để tự quản lý không - đó là một câu hỏi lớn. Khi AI trở nên thông minh hơn, AI có thể lừa dối chúng ta, nói dối chúng ta, thao túng chúng ta.

Giống như ví dụ về captcha. (Vào năm 2023, một chatbot AI, giả vờ là một người kiểm thị, thuyết phục một người cung cấp mã cần thiết để vượt qua bài kiểm tra mã Captcha (phép thử tự động để phân biệt máy tính với con người) - một hệ thống được thiết kế để phân biệt giữa con người và robot)?

Ví dụ về captcha, chính xác. Một thứ kém thông minh hơn kiểm soát một thứ thông minh hơn hầu như không bao giờ xảy ra trên thế giới. Khỉ không kiểm soát con người; chúng ta kiểm soát chúng. Con kiến không kiểm soát chúng ta; chúng ta kiểm soát con kiến. Vì vậy, nếu AI đã trở nên thông minh hơn, thì rất khó có khả năng chúng ta kiểm soát được nó. Chúng ta đã mất kiểm soát rồi sao? Chưa đâu. Chúng ta vẫn kiểm soát được, nhưng AI đang phát triển rất, rất nhanh. Nếu chúng ta không cẩn thận, chúng ta có thể mất kiểm soát trong năm đến mười năm nữa.

Tác giả: **Navin J. Antony**/Viết dịch: **Thích Vân Phong**/Nguồn: www.theweek.in