

"ĐẠI CHIẾN PHÁP LÝ" GIỮA *THE NEW YORK TIMES* VÀ *OPENAI* TRONG KỶ NGUYÊN TRÍ TUỆ NHÂN TẠO

Luật sư Lê Quang Vinh

Công ty sở hữu trí tuệ Bross & Cộng sự (Bross & Partners)

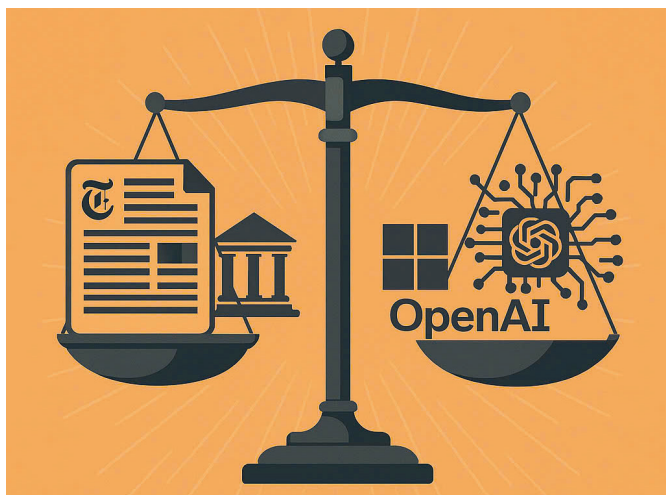
“

Tháng 12/2023, Thời báo New York (The New York Times - NYT) đã đệ đơn kiện OpenAI và Microsoft ra Tòa án Khu vực phía Nam New York, Hoa Kỳ, cáo buộc họ vi phạm quyền tác giả khi sử dụng hàng triệu bài báo của NYT để huấn luyện các mô hình trí tuệ nhân tạo (AI) tạo sinh như ChatGPT và Bing Chat. Vụ kiện được yêu cầu xét xử trước bồi thẩm đoàn, mang số hiệu 1:23-cv-11195, là một cuộc chiến pháp lý mang tính lịch sử, có khả năng định hình lại tương lai của ngành công nghiệp AI và quyền sở hữu trí tuệ toàn cầu.

”

Quy trình huấn luyện AI - nơi phát sinh cáo buộc vi phạm quyền tác giả

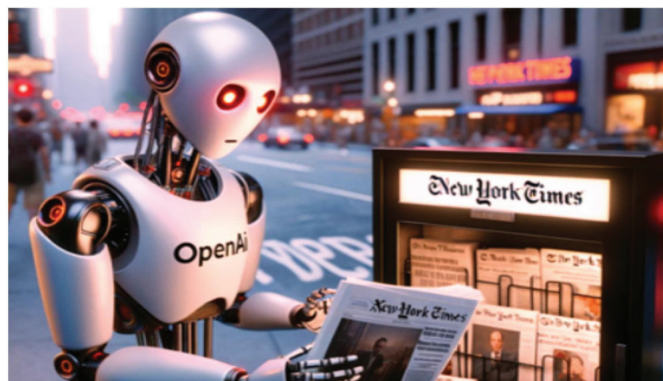
NYT cho rằng, việc bị đơn sử dụng trái phép nội dung của họ đã tạo ra các sản phẩm cạnh tranh trực tiếp, gây suy yếu mô hình kinh doanh của tờ báo. OpenAI và Microsoft phản bác rằng, hành vi của họ được bảo vệ bởi học thuyết “fair use” (sử dụng hợp lý) và có tính “biến đổi” (transformative), mang lại lợi ích chung cho xã hội. Chúng ta sẽ khám phá quy trình kỹ thuật, cáo buộc cụ thể, lập luận tự vệ của bị đơn, các vấn đề pháp lý trung tâm và cập nhật phán quyết của tòa - tất cả đều dựa trên Đơn kiện của NYT [1] và Bản ý kiến của Thẩm phán Sydney H. Stein [2].



Cuộc chiến pháp lý giữa NYT với OpenAI và Microsoft đang được cả thế giới quan tâm.

Để hiểu rõ cáo buộc của NYT, cần nắm được quy trình cơ bản mà các mô hình AI như GPT được huấn luyện. Quy trình này được mô tả chi tiết trong đơn kiện, không chỉ là kỹ thuật mà còn là nền tảng pháp lý cho cáo buộc vi phạm quyền tác giả trực tiếp (quyền sao chép và phân phối). NYT lập luận rằng, từ giai đoạn thu thập dữ liệu đến ghi nhớ, bị đơn đã tạo và sử dụng “bản sao” của tác phẩm NYT mà không xin phép, dẫn đến thiệt hại kinh tế và danh tiếng. Quy trình huấn luyện kết thúc khi mô hình được “đóng gói” (frozen) và không bao gồm giai đoạn tạo trích xuất đầu ra (trích xuất đầu ra) cho người dùng (phần này sẽ được phân tích ở đoạn sau như một biểu hiện của vi phạm). Hãy phân tích ba bước chính của huấn luyện:

Thu thập (Crawl): Các công ty AI sử dụng chương trình tự động (web crawlers) để thu thập hàng tỷ dữ liệu (bài viết, hình ảnh, văn bản...) từ Internet, bao gồm các nguồn như



Trọng tâm của vụ kiện tập trung vào dữ liệu huấn luyện cho mô hình AI là ChatGPT đã tạo và sử dụng “bản sao” các tác phẩm của NYT mà không xin phép.

nytimes.com. Quy trình này tương tự cách Google thu thập dữ liệu cho công cụ tìm kiếm, nhưng NYT cáo buộc OpenAI ưu tiên nội dung chất lượng cao của họ, với tên miền nytimes.com chiếm 1,23% số URL trong tập dữ liệu OpenWebText2 và là nguồn độc quyền lớn thứ ba trong thu thập thường xuyên (Common Crawl). OpenAI thừa nhận “các bộ dữ liệu chất lượng cao hơn sẽ được lấy mẫu thường xuyên hơn” (tài liệu công khai của OpenAI, 2023), củng cố lập luận rằng việc thu thập không phải ngẫu nhiên mà có chủ đích khai thác. Từ góc nhìn pháp lý, giai đoạn này có thể cấu thành vi phạm sao chép ban đầu, vì dữ liệu được lưu trữ tạm thời để xử lý, dù bị đơn cho rằng đây là dữ liệu “công khai”.

- **Trích xuất/Lọc (Extract):** Dữ liệu thô được xử lý, lọc bỏ các phần không cần thiết (như quảng cáo, liên kết lỗi) và chuẩn bị cho huấn luyện thông qua tokenization (phân tách thành đơn vị ngôn ngữ) và embedding (chuyển đổi thành vector toán học). NYT lập luận rằng, ngay cả sau lọc, dữ liệu có quyền tác giả của họ vẫn được giữ lại dưới dạng mã hóa, tạo “bản sao trung gian” (intermediate copies) vi phạm quyền tác giả. Điều này giống như vụ MAI Systems Corp kiện Peak Computer, Inc. (1991), nơi tòa án công nhận bản sao tạm thời trong RAM cũng là vi phạm. Bị đơn phản bác rằng, lọc là để “làm sạch” dữ liệu, không lưu trữ lâu dài, nhưng tòa án đã giữ nguyên cáo buộc này để xét xử.

- **Huấn luyện (Training):** Đây là giai đoạn cốt lõi, nơi mô hình AI được “dạy” bằng cách xử lý lượng lớn dữ liệu qua học sâu (deep learning), sử dụng thuật toán như transformer để điều chỉnh tham số (parameters). Trong quá trình này, mô hình “ghi nhớ” và mã hóa thông tin từ dữ liệu huấn luyện vào cấu trúc nội bộ. NYT cáo buộc việc tạo và lưu trữ hàng triệu bản sao tác phẩm của họ trong bộ nhớ máy tính để huấn luyện AI đã cấu thành hành vi vi phạm quyền tác giả trực tiếp. Ví dụ, nội dung NYT chiếm tỷ lệ đáng kể trong Common Crawl (một tập dữ liệu khổng lồ với hàng tỷ trang web), và ghi nhớ cho phép AI tái tạo nội dung gốc. Bị đơn lập luận đây là “biến đổi” (transformative), chỉ trích xuất mẫu ngôn ngữ chứ không sao chép để tái bản, nhưng NYT trích dẫn nghiên cứu (như từ Carlini et al., 2023) chứng minh các mô hình ngôn ngữ lớn (Large language models - LLMs) có thể “nhớ” và tái tạo 95% nội dung huấn luyện nếu prompt đúng.

Sau huấn luyện, mô hình sẵn sàng cho giai đoạn sử dụng, nơi trích xuất đầu ra được tạo ra - một biểu hiện của ghi nhớ nhưng không thuộc quy trình huấn luyện. Quy trình huấn luyện này không chỉ là kỹ thuật mà còn có tính pháp lý: NYT lập luận nó tạo “bản sao mã hóa” trong tham số AI, tương đương sao chép toàn bộ tác phẩm. Bị đơn cho rằng đây là “nghiên cứu” và đề nghị không xem xét, nhưng bị bác bỏ với lý do cần xét xử đầy đủ để đánh giá tính thương mại.

NYT cáo buộc cụ thể gì và bằng chứng nào?

NYT cáo buộc OpenAI và Microsoft đã xâm phạm quyền tác giả hàng loạt bằng cách sử dụng trái phép hàng triệu bài báo có bản quyền của NYT để huấn luyện mô hình GPT mà



Vấn đề về bản quyền dữ liệu sử dụng để huấn luyện cho các mô hình AI đã bàn đến từ lâu nhưng chưa có phán quyết nào được đưa ra.

không xin phép hay trả phí. Các cáo buộc cốt lõi bao gồm vi phạm trực tiếp, trách nhiệm liên đới của Microsoft, và vi phạm Đạo luật Bản quyền thiên niên kỷ kỹ thuật số (Digital Millennium Copyright Act - DMCA) về xóa thông tin quản lý quyền tác giả. NYT yêu cầu bồi thường thiệt hại thực tế, lợi nhuận của bị đơn và lệnh cấm sử dụng dữ liệu NYT trong tương lai. Các cáo buộc cụ thể:

Sao chép và ghi nhớ: Bị đơn đã thu thập và lưu trữ “hàng triệu bài báo có bản quyền của tờ New YorkTimes”, ưu tiên nội dung NYT vì “chất lượng cao”. NYT chứng minh qua phân tích thu thập thường xuyên: nytimes.com chiếm tỷ lệ lớn và LLMs như GPT-4 “ghi nhớ” đoạn văn bản nguyên văn. Nghiên cứu độc lập (từ Đại học Stanford, 2023) cho thấy, LLMs có tỷ lệ ghi nhớ lên đến 10% dữ liệu huấn luyện, dẫn đến trích xuất đầu ra vi phạm.

- **Trích xuất đầu ra - biểu hiện của vi phạm và các ví dụ thực tế:** Mặc dù không thuộc quy trình huấn luyện, trích xuất đầu ra là giai đoạn suy luận sau huấn luyện, nơi mô hình sử dụng tham số đã mã hóa để tạo phản hồi từ prompt (yêu cầu do người dùng nhập vào). Đây là “biểu hiện” của ghi nhớ, cấu thành vi phạm quyền phân phối nếu tái tạo nội dung gốc. NYT liệt kê hàng trăm ví dụ có thể tái tạo (reproducible), như ChatGPT tái tạo bài đánh giá nhà hàng của Pete Wells về Guy Fieri khi nhắc lại từng câu (sao chép câu châm biếm “một mớ hỗn độn của phô mai tan chảy và thịt xông khói giả tạo”), hoặc các đoạn từ loạt bài điều tra đoạt giải Pulitzer về ngành taxi New York (tái tạo “các quan chức thành phố đã miễn trừ nó khỏi các quy định”). Một ví dụ khác: Loạt bài về chuỗi cung ứng iPhone (2012), AI sao chép chi tiết “Apple sử dụng 700.000 người ở nước ngoài” và “8.000 công nhân Trung Quốc huy động 96 giờ để sản xuất 10.000 iPhone”. Những ví dụ này không chỉ sao chép nội dung mà còn bắt chước phong cách biểu đạt sáng tạo (expressive style) của NYT, vi phạm yếu tố thứ 2 của “sử dụng hợp lý” là “tính chất tác phẩm sáng tạo”. Trích xuất đầu ra giống 95-100% bài báo gốc nhưng không trích nguồn hoặc liên kết, thay thế trực tiếp lưu lượng truy cập NYT. Đến 08/09/2025, trong quá trình trao đổi chứng cứ hay còn gọi là khám phá chứng cứ (discovery) theo Điều 26(b)(1)

Luật thủ tục tố tụng liên bang, toà buộc bị đơn bảo tồn bản ghi nhật ký hệ thống (logs), cụ thể là tòa án ra lệnh bảo tồn 120 triệu đoạn hội thoại (tháng 7/2025) để chứng minh vi phạm có hệ thống.

- *Gây nhầm lẫn và suy giảm giá trị*: AI tạo ra các “ảo giác” (hallucinations) - thông tin sai lệch nhưng được gán sai cho NYT, gây nhầm lẫn cho người đọc và làm xói mòn uy tín của tờ báo. Ví dụ: AI bịa “danh sách 15 thực phẩm tốt cho tim từ NYT” bao gồm rượu vang đỏ (không tồn tại), hoặc trích dẫn sai “con gái Steve Forbes về Amazon”. Những “ảo giác” này bị gán cho NYT, dẫn đến mất lòng tin độc giả – thiệt hại danh tiếng ước tính hàng triệu đô la.

- *Thiệt hại tài chính*: NYT yêu cầu bồi thường hàng tỷ USD vì cho rằng, việc sử dụng trái phép này đã giúp Microsoft tăng vốn hóa thị trường thêm 1.000 tỷ USD và OpenAI đạt mức định giá 90 tỷ USD (nay lên 157 tỷ USD theo cập nhật 2025). Ngược lại, Wirecutter (phụ san của NYT) mất lượng truy cập do AI sao chép khuyến nghị sản phẩm; tổng thiệt hại ước tính 100 triệu USD/năm từ giảm số lượng đăng ký. NYT trích dẫn dữ liệu nội bộ: Lượng truy cập từ truy vấn AI tăng nhưng không chuyển đổi thành doanh thu, vì trích xuất đầu ra thay thế nội dung gốc.

Bằng chứng của NYT thể hiện mạnh mẽ với hơn 100 ví dụ “có thể tái tạo” qua prompt, phân tích dữ liệu huấn luyện (từ thu thập thường xuyên) và chuyên gia chứng minh ghi nhớ dẫn đến trích xuất đầu ra vi phạm. Đến 08/09/2025, thủ tục trao đổi chứng cứ được toà xác nhận qua nhật ký hệ thống trích xuất đầu ra (bị đơn phải bảo tồn 120 triệu đoạn hội thoại theo lệnh tháng 7/2025), củng cố cáo buộc vi phạm hệ thống.

OpenAI & Microsoft tự vệ bằng cách nào?

Phía bị đơn không phủ nhận việc sử dụng dữ liệu từ Internet để huấn luyện mô hình, nhưng phản bác mạnh mẽ với các lập luận chính, chủ yếu dựa trên “sử dụng hợp lý” như một “lá chắn”. Họ lập luận hành vi là “biến đổi” và không gây thiệt hại, tương tự vụ một nhóm tác giả kiện Google (2015) nơi quét sách là “sử dụng hợp lý”.

- *Dữ liệu “công khai”*: Dữ liệu đầu vào đều là có sẵn công khai, tức là không bị chặn, không cần đăng nhập hay trả phí, tương tự như cách Google vẫn thu thập dữ liệu. OpenAI trích dẫn robots.txt (file hướng dẫn thu thập) của NYT không cấm thu thập và dữ liệu từ thu thập thường xuyên là “dữ liệu mở” (open data). Họ lập luận thu thập là “tạm thời” (transient), không vi phạm và tỷ lệ NYT chỉ 1,23% không đáng kể so với hàng tỷ nguồn khác.

- *Mục đích “biến đổi” (transformative)*: Việc sử dụng dữ liệu nhằm mục đích phân tích và trích rút mẫu ngôn ngữ, không phải tái bản hay sao chép nội dung. Trích xuất đầu ra của AI là kết quả sáng tạo mới, không phải copy nguyên văn, và trong nhiều trường hợp có yếu tố “chuyển hóa” - ví dụ,



NYT căn cứ nhật ký hệ thống trích xuất đầu ra với nội dung các đoạn hội thoại để chứng minh OpenAI đã vi phạm bản quyền.
Nguồn: Internet.

tóm tắt hoặc tổng hợp kiến thức. Microsoft nhấn mạnh Azure chỉ là “công cụ trung lập”, có sử dụng không vi phạm đáng kể như huấn luyện mô hình y tế hoặc giáo dục, theo án lệ vụ án giữa Sony và Universal (1984). Họ trích dẫn nghiên cứu nội bộ: Sau khi tinh chỉnh, tỷ lệ ghi nhớ giảm 90% và trích xuất đầu ra thường thêm giá trị mới (như phân tích cá nhân hóa).

- *Khái niệm “sử dụng hợp lý”*: Quá trình này được bảo vệ bởi khái niệm “sử dụng hợp lý” trong luật quyền tác giả Hoa Kỳ, cho phép sử dụng tác phẩm có quyền tác giả với những giới hạn nhất định (như phê bình, học thuật, nghiên cứu) không phải xin phép và cũng không phải trả tiền. Bị đơn phân tích 4 yếu tố:

+ *Mục đích*: Biến đổi, không thương mại thuần túy (dù ChatGPT có phí cho bản cao cấp, nhưng huấn luyện là nghiên cứu).

+ *Tính chất tác phẩm*: Nội dung NYT là “thông tin thực tế”, dễ “sử dụng hợp lý” hơn sáng tạo hư cấu.

+ *Lượng sử dụng*: Cần toàn bộ dữ liệu để huấn luyện hiệu quả, không phải “toàn bộ” một tác phẩm cụ thể.

+ *Tác động thị trường*: Không thay thế NYT; doanh thu NYT tăng 10% năm 2024 nhờ tin tức về AI (AI buzz) và AI cũng khuyến khích người dùng kiểm tra nguồn gốc.

OpenAI còn lập luận NYT “đạo đức giả” (hypocritical), vì NYT cũng dùng AI nội bộ cho chỉnh sửa. Đến 08/09/2025, bị đơn đề xuất cung cấp 20 triệu đoạn hội thoại cho giai đoạn trao đổi chứng cứ (tháng 8/2025).

Các vấn đề pháp lý trung tâm

Tòa án sẽ phải xem xét nhiều khía cạnh phức tạp, nhưng ba điểm then chốt nhất sẽ định đoạt số phận của vụ kiện:

- *Trích xuất đầu ra của AI có thực sự là “biến đổi” không?* Nếu AI tạo ra nội dung có ý nghĩa, mục đích mới và không gây tổn hại thị trường gốc, nó có thể được coi là “sử dụng hợp lý”. Ngược lại, nếu chỉ “nhại lại” hoặc thay thế trực tiếp sản phẩm gốc, nó có thể vẫn là vi phạm.

- Việc huấn luyện AI có phải là hành vi “sử dụng hợp lý”? NYT lập luận rằng, việc huấn luyện AI không chỉ trích một phần nhỏ để bình luận hay nghiên cứu, mà dùng toàn bộ tập dữ liệu để tái cấu trúc hệ thống ngôn ngữ, đây là sử dụng mang tính khai thác thương mại, không phải sử dụng hợp lý. Họ so sánh với vụ kiện liên quan tới Google Books (“sử dụng hợp lý” chỉ tạo các chỉ mục), nhưng huấn luyện GPT là “sao chép toàn diện”. Bị đơn phản bác đây là “nghiên cứu mang tính biến đổi”, tương tự huấn luyện mô hình y tế bằng dữ liệu bệnh án (không vi phạm đạo luật HIPAA nếu ẩn danh).

- Microsoft có phải chịu trách nhiệm liên đới? Vì Microsoft là nhà đầu tư lớn và đối tác chiến lược, cung cấp cơ sở hạ tầng siêu máy tính Azure để huấn luyện GPT nên NYT cho rằng, Microsoft đã “tiếp tay và hưởng lợi” từ hành vi vi phạm và có thể bị xem là góp phần vi phạm quyền tác giả. Yếu tố: Microsoft biết qua đầu tư và hỗ trợ vật chất (Azure độc quyền). Bị đơn lập luận Microsoft chỉ là “nhà đầu tư thụ động”, không kiểm soát trích xuất đầu ra và Azure có sử dụng không vi phạm (như cloud cho doanh nghiệp). Ý kiến Thẩm phán giữ lại cáo buộc này, yêu cầu trao đổi chứng cứ về hợp đồng giữa OpenAI và Microsoft.

Các vấn đề này chạm đến tiền lệ lớn: Liệu huấn luyện AI có phải “nghiên cứu công bằng” hay “ăn cắp dữ liệu”?

Tòa đã nói gì?

Vào ngày 04/04/2025, bản Ý kiến của Thẩm phán Sydney H. Stein từ chối kiến nghị của bị đơn bác bỏ vụ kiện đối với hầu hết cáo buộc, đánh dấu chiến thắng lớn cho NYT. Tòa chấp nhận lập luận từ NYT có đủ cơ sở để xét xử, đặc biệt vi phạm quyền tác giả trực tiếp (huấn luyện bằng dữ liệu NYT) và trách nhiệm liên đới của Microsoft (“tiếp tay qua Azure”). Tuy nhiên, tòa bác cáo buộc cạnh tranh không lành mạnh và một phần DMCA (thiếu chi tiết về xóa CMI), vì trích xuất đầu ra AI không phải “bản sao toàn bộ”.

Cập nhật đến 08/09/2025, vụ kiện tiến triển ở giai đoạn trao đổi chứng cứ, với lệnh bảo tồn nhật ký hệ thống ChatGPT yêu cầu OpenAI giữ và phân loại trích xuất đầu ra dữ liệu để chứng minh sao chép. OpenAI kháng cáo tháng 6/2025 (vi phạm quyền riêng tư), nhưng bị bác bỏ tháng 7/2025, buộc cung cấp dữ liệu (NYT đòi 120 triệu đoạn hội thoại, bị đơn đề xuất 20 triệu tháng 8/2025). Đơn kiện từ chối xuất tài liệu nội bộ NYT (như sử dụng AI của họ), cũng cố “sử dụng hợp lý” là nghĩa vụ của bị đơn. Hiện tại vẫn chưa có lịch xét xử tiếp theo. Dự đoán, các bên có thể dàn xếp trước năm 2026, với áp lực từ các vụ tương tự (như

vụ Getty kiện Stability AI). Phán quyết vụ án trên tạo tiền lệ: Huấn luyện AI bằng dữ liệu có quyền tác giả không tự động được coi là “sử dụng hợp lý” mà cần chứng minh tại tòa.

Tác động rộng lớn của vụ kiện



Cả OpenAI và Microsoft đều cho rằng việc huấn luyện ChatGPT tuân thủ nguyên tắc “sử dụng hợp lý” đối với dữ liệu. Nguồn: Internet.

Vụ kiện sẽ tạo tiền lệ pháp lý quan trọng, định hình lại cách các công ty AI hoạt động, cách các nhà sáng tạo nội dung được bảo vệ và đền bù, cụ thể:

- **Ngành xuất bản và báo chí:** Nếu NYT thắng, các công ty AI có thể bị buộc phải xin phép hoặc trả phí quyền tác giả cho việc sử dụng nội dung để huấn luyện mô hình, giúp bảo vệ nguồn thu và mô hình kinh doanh của các tổ chức báo chí. Điều này có thể khuyến khích các mô hình cấp phép mới (ví dụ: cấp phép dữ liệu huấn luyện) và đảm bảo sự đền bù công bằng cho người sáng tạo.

- **Ngành đổi mới sáng tạo và công nghệ AI:** Kết quả sẽ định hình lại cách AI được phát triển, huấn luyện và thương mại hóa. Nó có thể làm tăng chi phí phát triển AI nhưng cũng tạo ra một môi trường pháp lý rõ ràng hơn, thúc đẩy sự minh bạch và trách nhiệm giải trình trong việc sử dụng dữ liệu có quyền tác giả.

- **Bối cảnh Việt Nam:** Đối với Việt Nam, một quốc gia đang thúc đẩy đổi mới AI, việc theo dõi sát sao vụ kiện này là cực kỳ quan trọng. Phán quyết sẽ cung cấp những bài học quý giá để Việt Nam xây dựng khung pháp lý về quyền tác giả AI phù hợp, cân bằng giữa việc bảo vệ quyền tác giả và khuyến khích sự phát triển của công nghệ AI trong nước. Nó cũng sẽ là cơ sở để các nhà sáng tạo nội dung Việt Nam bảo vệ tác phẩm của mình trước sự phát triển của AI, đặc biệt trong bối cảnh các công ty AI toàn cầu hoạt động xuyên biên giới.

Tài liệu tham khảo

[1] E. Roth (2023), “The New York Times is suing OpenAI and Microsoft for copyright infringement”, <https://www.theverge.com/2023/12/27/24016212/new-york-times-openai-microsoft-lawsuit-copyright-infringement>, truy cập ngày 05/08/2025.

[2] Bản ý kiến của Thẩm phán Sydney H. Stein, <https://ffingfx.thomsonreuters.com/gfx/legaldocs/znpnjnkgqpl/NYT%20OPENAI%20COPYRIGHT%20LAWSUIT%20mtdruling.pdf>, truy cập ngày 05/08/2025.