

PHÁP LUẬT HOA KỲ VỀ TRÁCH NHIỆM DÂN SỰ VÀ BỒI THƯỜNG THIẾT HẠI DO TRÍ TUỆ NHÂN TẠO GÂY RA VÀ GỢI MỞ CHO VIỆT NAM

NGUYỄN TRÍ CƯỜNG *

Bài viết phân tích khung pháp luật của Hoa Kỳ về trách nhiệm dân sự đối với thiệt hại do AI gây ra thông qua các học thuyết trách nhiệm sản phẩm, trách nhiệm nghiêm ngặt và trách nhiệm do sơ suất; đồng thời, làm rõ những xu hướng lập pháp mới. Trên cơ sở đối chiếu với Bộ luật Dân sự năm 2015, Luật Trí tuệ nhân tạo năm 2025 và Luật Công nghiệp công nghệ số năm 2025, bài viết đề xuất một số gợi mở nhằm hoàn thiện cơ chế bồi thường thiệt hại do AI gây ra ở Việt Nam.

Từ khóa: Trí tuệ nhân tạo; trách nhiệm dân sự; bồi thường thiệt hại; trách nhiệm sản phẩm; pháp luật Việt Nam.

This article analyzes the United States legal framework governing civil liability for damages caused by artificial intelligence (AI), focusing on the doctrines of product liability, strict liability, and liability based on negligence, while also examining emerging legislative trends. By comparing these developments with the 2015 Civil Code, the 2025 Law on Artificial Intelligence, and the 2025 Law on the Digital Technology Industry of Vietnam, the article proposes several recommendations for improving the legal framework governing compensation for AI-caused damages in Vietnam.

Keywords: Artificial intelligence; civil liability; compensation for damages; product liability; Vietnamese law.

NGÀY NHẬN: 12/4/2026 NGÀY PHẢN BIỆN, ĐÁNH GIÁ: 25/6/2026 NGÀY DUYỆT ĐĂNG: 08/7/2026

Email: cuongnt@uel.edu.vn

DOI: <https://doi.org/10.59394/qlnn.366.2026.1589>

1. Đặt vấn đề

Sự phát triển nhanh chóng của AI đang đặt ra những thách thức mới đối với pháp luật về trách nhiệm dân sự. Với đặc tính tự học, tính khó giải thích và khả năng tự đưa ra quyết định, AI làm cho việc xác định lỗi, chủ thể chịu trách nhiệm và mối quan hệ nhân quả trở nên phức tạp hơn so với các cơ chế bồi thường thiệt hại truyền thống (Behrendt,

2025). Trong bối cảnh đó, Hoa Kỳ là một trong những quốc gia tiên phong điều chỉnh trách nhiệm pháp lý đối với AI thông qua sự phát triển của án lệ và các sáng kiến lập pháp mới, nổi bật là AI LEAD Act năm 2025 (Durbin, 2025). Trong khi đó, tại Việt Nam, sự ra đời của Luật Trí tuệ nhân tạo năm 2025

* TS, Trường Đại học Kinh tế - Luật, Đại học Quốc gia Thành phố Hồ Chí Minh

trong bối cảnh kỷ nguyên số cho thấy tầm nhìn chiến lược của Nhà nước trong việc kiến tạo hành lang pháp lý cho nền kinh tế số. Tuy nhiên, việc áp dụng các nguyên tắc bồi thường thiệt hại ngoài hợp đồng đối với một công nghệ biến đổi nhanh như AI vẫn là thách thức lớn đối với các cơ quan thực thi pháp luật. Nghiên cứu kinh nghiệm của Hoa Kỳ, từ các học thuyết pháp lý và thực tiễn xét xử đến các mô hình bồi thường để làm rõ các quy định hiện hành của Việt Nam và gợi mở những cơ chế hỗ trợ tài chính cũng như quản trị rủi ro cần thiết để bảo vệ quyền lợi con người mà không làm cản trở đổi mới sáng tạo.

2. Trách nhiệm dân sự và bồi thường thiệt hại do AI gây ra theo pháp luật Hoa Kỳ

Hệ thống pháp luật Hoa Kỳ giải quyết các thiệt hại do AI gây ra thông qua một mạng lưới phức tạp các học thuyết của Tort Law (pháp luật về hành vi xâm phạm quyền lợi ngoài hợp đồng), pháp luật bảo vệ người tiêu dùng và các đạo luật chuyên biệt cấp bang và liên bang mới được ban hành. Một trong những căn cứ phổ biến nhất để yêu cầu bồi thường là học thuyết trách nhiệm sản phẩm. Theo học thuyết này, bất kỳ chủ thể nào trong chuỗi sản xuất, bao gồm: nhà sản xuất, nhà lắp ráp hoặc nhà phân phối đều có thể phải chịu trách nhiệm đối với thiệt hại do sản phẩm có khiếm khuyết gây ra (ICLG, 2025). Tuy nhiên, một rào cản pháp lý truyền thống tại Hoa Kỳ cho rằng, phần mềm chỉ là dịch vụ hoặc tài sản phi vật thể nên không thuộc đối tượng điều chỉnh của trách nhiệm sản phẩm nghiêm ngặt (Morrison Foerster, 2025). Thực tiễn xét xử giai đoạn 2024 - 2025 cho thấy, quan điểm này đã thay đổi đáng kể, thậm chí mang tính lịch sử.

Trong phán quyết tiêu biểu của vụ *Garcia v. Character Technologies, Inc.* (2025), Tòa án quận Middle District của Florida đã bác bỏ lập luận cho rằng, chatbot AI chỉ là biểu đạt ngôn ngữ được bảo vệ bởi Tu chính án thứ nhất (Frank Young, 2025). Tòa án khẳng định, ứng dụng chatbot có thể được coi là một “sản phẩm” cho mục đích bồi thường thiệt hại nếu các khiếm khuyết trong thiết kế

của nó, chẳng hạn thiếu cơ chế ngăn chặn sự ám ảnh độc hại đối với trẻ vị thành niên, dẫn đến thiệt hại thực tế như hành vi tự sát của người sử dụng (Kimberly Chew & cộng sự, 2025). Phán quyết này mở đường cho nhiều vụ kiện đáng chú ý khác nhằm vào các nền tảng AI như OpenAI và Google khi các hệ thống AI gây ra tổn thương về tâm lý hoặc thể chất (Chat GPT Is Eating the World, 2025).

Để xác lập trách nhiệm sản phẩm, nguyên đơn phải chứng minh sự tồn tại của một trong 3 dạng khiếm khuyết cơ bản: (1) Khiếm khuyết thiết kế (Design Defect) phát sinh khi thuật toán được xây dựng theo cách tạo ra rủi ro không hợp lý cho người sử dụng mặc dù vận hành đúng mục đích; (2) Khiếm khuyết sản xuất (Manufacturing Defect) xảy ra khi một phiên bản AI cụ thể sai lệch so với thiết kế ban đầu hoặc phát sinh lỗi trong quá trình triển khai, chẳng hạn thiết bị y tế tích hợp AI bị lỗi phần sụn khiến phân phối quá liều thuốc (Elizabeth & cộng sự, 2024); (3) Khiếm khuyết cảnh báo xuất hiện khi nhà phát triển không cung cấp đầy đủ hướng dẫn hoặc cảnh báo về giới hạn và nguy cơ của hệ thống AI, ví dụ không thông báo rằng, thuật toán chẩn đoán hình ảnh chỉ đạt độ chính xác cao khi sử dụng hình ảnh có độ phân giải phù hợp (Villasenor, 2019).

Sự phức tạp của AI nằm ở chỗ, các rủi ro thường phát sinh từ quá trình “tự học” của máy. Phân tích pháp lý từ Viện Brookings cho thấy, các công ty AI thường đổ lỗi cho dữ liệu đầu vào hoặc cách sử dụng của người dùng (Villasenor, 2019). Tuy nhiên, các Tòa án Hoa Kỳ đang dần áp dụng tiêu chuẩn rằng, nếu một hệ thống AI được thiết kế để thích nghi, nhà sản xuất phải chịu trách nhiệm về các “kỹ năng” hoặc “hành vi” không mong muốn xuất hiện trong quá trình đó (U.S. Congress, 2025).

Bên cạnh học thuyết trách nhiệm sản phẩm, pháp luật Hoa Kỳ còn phân biệt giữa trách nhiệm nghiêm ngặt và trách nhiệm do sơ suất. Đối với trách nhiệm nghiêm ngặt, bị đơn phải bồi thường nếu sản phẩm tồn tại khiếm khuyết nguy hiểm một cách bất hợp

lý, bất kể họ đã áp dụng mức độ cẩn trọng nào trong quá trình phát triển (ICLG, 2025). Đối với AI, điều này giúp loại bỏ gánh nặng cho nạn nhân trong việc phải giải mã các dòng mã nguồn phức tạp để tìm ra sai sót của lập trình viên (Villasenor, 2019). Theo định hướng này, Dự thảo *Đạo luật AI LEAD Act* năm 2025 đề xuất thiết lập cơ chế trách nhiệm nghiêm ngặt ở cấp liên bang, theo đó, các nhà phát triển AI sẽ phải chịu trách nhiệm đối với các thiệt hại phát sinh từ khiếm khuyết của hệ thống, ngay cả khi họ chứng minh đã thực hiện đầy đủ các biện pháp bảo đảm an toàn (U.S. Senate Judiciary Committee, 2025).

Trong những trường hợp không thể áp dụng trách nhiệm sản phẩm, chẳng hạn AI chỉ đóng vai trò là công cụ hỗ trợ chuyên môn, các tranh chấp thường được giải quyết theo học thuyết trách nhiệm do sơ suất (Smith & cộng sự, 2024). Trọng tâm khi đó là xác định liệu nhà phát triển AI hoặc người sử dụng AI có thực hiện đầy đủ nghĩa vụ cẩn trọng hay không. Các Tòa án Hoa Kỳ thường tham chiếu các tiêu chuẩn ngành, như “Khung Quản trị rủi ro AI của NIST” để đánh giá việc thử nghiệm, giám sát và kiểm định hệ thống AI (Smith & cộng sự, 2024). Chẳng hạn, trong lĩnh vực y tế, bác sĩ có thể bị coi là sơ suất nếu hoàn toàn phụ thuộc vào kết quả chẩn đoán của AI mà bỏ qua các dấu hiệu lâm sàng rõ ràng hoặc ngược lại, bác bỏ một khuyến nghị chính xác của AI mà một bác sĩ có năng lực thông thường phải lưu tâm (Sommers Schwartz, 2025).

Ngoài các thiệt hại về tính mạng, sức khỏe và tài sản, pháp luật Hoa Kỳ còn mở rộng phạm vi trách nhiệm dân sự đối với các hành vi phân biệt đối xử bằng thuật toán. Các doanh nghiệp sử dụng AI để đưa ra những quyết định mang tính hệ quả, như tuyển dụng, cho vay hoặc bảo hiểm có thể phải đối mặt với các vụ kiện tập thể nếu thuật toán tạo ra tác động bất lợi đối với các nhóm được pháp luật bảo vệ, chẳng hạn dựa trên chủng tộc, giới tính hoặc độ tuổi (Anjali

Das, 2026). Vụ kiện *Mobley v. Workday, Inc.* (2025) là một ví dụ tiêu biểu khi các nguyên đơn trên 40 tuổi cáo buộc hệ thống AI của Workday tự động loại bỏ họ khỏi hàng trăm cơ hội việc làm. Tòa án quận Bắc California đã phán quyết rằng, Workday có thể bị kiện với tư cách là “đại lý” của các doanh nghiệp sử dụng lao động. Phán quyết này xác lập nguyên tắc quan trọng rằng, việc doanh nghiệp giao quyền ra quyết định cho thuật toán AI không làm mất đi nghĩa vụ tuân thủ các đạo luật về quyền dân sự và chống phân biệt đối xử (Quinn Emanuel, 2026).

3. Pháp luật Việt Nam về trách nhiệm dân sự và bồi thường thiệt hại do AI gây ra

Tại Việt Nam, khung pháp lý điều chỉnh trách nhiệm dân sự đối với thiệt hại do AI gây ra đang từng bước được hình thành trên cơ sở kết hợp giữa các nguyên tắc chung của *Bộ luật Dân sự* năm 2015 với các quy định chuyên ngành trong *Luật Trí tuệ nhân tạo* năm 2025 và *Luật Công nghiệp công nghệ số* năm 2025. Mặc dù chưa xây dựng một chế định trách nhiệm dân sự riêng dành cho AI, hệ thống pháp luật hiện hành đã bước đầu thiết lập cơ sở pháp lý để xác định chủ thể chịu trách nhiệm, phạm vi bồi thường và cơ chế phân bổ rủi ro đối với các hệ thống AI.

Về nền tảng pháp lý chung, *Bộ luật Dân sự* năm 2015 vẫn là đạo luật trung tâm điều chỉnh trách nhiệm bồi thường thiệt hại ngoài hợp đồng. Điều 584 xác lập nguyên tắc mọi chủ thể phải bồi thường khi có hành vi xâm phạm tính mạng, sức khỏe, danh dự, nhân phẩm, uy tín, tài sản hoặc quyền, lợi ích hợp pháp của chủ thể khác. Trong tiến trình hoàn thiện pháp luật năm 2025, một số quy định mới đã tạo tiền đề quan trọng để thích ứng với sự phát triển của AI. Trước hết, *Luật Công nghiệp công nghệ số* năm 2025 lần đầu tiên thừa nhận tài sản số là một loại tài sản theo nghĩa của *Bộ luật Dân sự*, qua đó, mở rộng phạm vi thiệt hại có thể được bồi thường khi AI xâm phạm dữ liệu cá nhân, tài sản số, tiền mã hóa hoặc các thực thể số có giá trị kinh tế. Đồng thời, *Bộ luật Dân sự* năm

2015 vẫn duy trì nguyên tắc chỉ cá nhân và pháp nhân mới có địa vị pháp lý để gánh chịu trách nhiệm dân sự. Theo đó, AI không được coi là chủ thể của quan hệ pháp luật mà chỉ là công cụ hoặc tài sản; mọi trách nhiệm phát sinh từ hoạt động của AI đều được truy nguyên về cá nhân hoặc tổ chức sở hữu, phát triển, quản lý hoặc vận hành hệ thống AI (Hương Giang, 2025).

Trên cơ sở đó, *Luật Trí tuệ nhân tạo* năm 2025 đã thiết lập khung quản trị AI dựa trên mức độ rủi ro, với cách tiếp cận tương đồng với EU AI Act nhưng được điều chỉnh phù hợp với điều kiện Việt Nam. Điều 29 *Luật Trí tuệ nhân tạo* năm 2025 thiết lập một quy trình phân định trách nhiệm rõ ràng trong chuỗi giá trị AI. Theo đó: (1) Trách nhiệm của bên triển khai (đây là điểm tương đồng lớn nhất với xu hướng Hoa Kỳ) được quy định tại khoản 2 Điều 29: “Trường hợp hệ thống trí tuệ nhân tạo có rủi ro cao được quản lý, vận hành và sử dụng đúng quy định nhưng vẫn phát sinh thiệt hại thì bên triển khai phải chịu trách nhiệm bồi thường cho người bị thiệt hại. Sau khi bồi thường, bên triển khai yêu cầu nhà cung cấp, nhà phát triển hoặc các bên liên quan hoàn trả khoản tiền bồi thường nếu có thỏa thuận giữa các bên”; (2) Quyền truy đòi: sau khi bồi thường, bên triển khai có quyền yêu cầu nhà cung cấp hoặc nhà phát triển hoàn trả số tiền đó nếu chứng minh được thiệt hại bắt nguồn từ khiếm khuyết thiết kế hoặc lỗi sản xuất của bên đó; (3) Sự can thiệp của bên thứ ba: nếu hệ thống bị tấn công mạng hoặc bị chiếm quyền điều khiển trái pháp luật, bên thứ ba gây ra hành vi đó phải bồi thường. Tuy nhiên, nếu bên triển khai hoặc nhà cung cấp có lỗi trong việc bảo đảm an toàn, an ninh hệ thống dẫn đến việc AI bị xâm nhập, các chủ thể này vẫn phải chịu trách nhiệm liên đới.

Một trong những vấn đề gây tranh cãi nhất là việc liệu có nên xếp AI vào nhóm “nguồn nguy hiểm cao độ” hay không (Thu Giang, 2025). Theo quy định hiện hành, nguồn nguy hiểm cao độ truyền thống bao

gồm các thực thể vật lý có khả năng gây thiệt hại lớn mà con người không thể kiểm soát tuyệt đối (xe cơ giới, hệ thống tải điện, chất nổ...). Vì vậy, xét về bản chất, phần mềm AI không phải là nguồn nguy hiểm theo cách hiểu truyền thống. Tuy nhiên, đối với các hệ thống AI được tích hợp vào các thiết bị vật lý, như: xe tự hành, robot phẫu thuật hoặc dây chuyền sản xuất tự động, mức độ nguy hiểm có thể tương đương với các nguồn nguy hiểm cao độ. Do đó, *Luật Trí tuệ nhân tạo* năm 2025 và các báo cáo lập pháp liên quan đã tiến tới sự thống nhất: thay vì coi mọi loại AI là nguồn nguy hiểm cao độ, Nhà nước chỉ áp dụng cơ chế bồi thường tương tự như nguồn nguy hiểm cao độ (trách nhiệm nghiêm ngặt) đối với các hệ thống AI rủi ro cao. Điều này bảo đảm sự công bằng cho nạn nhân khi không phải chứng minh lỗi; đồng thời, phù hợp với logic của *Bộ Luật Dân sự* năm 2015.

Hiện nay, Việt Nam chưa ghi nhận các phán quyết landmark về lỗi thuật toán như Hoa Kỳ, nhưng các vụ tranh chấp về lỗi phần mềm ngân hàng dẫn đến mất tiền trong tài khoản đã bắt đầu xuất hiện. Rào cản lớn nhất đối với nạn nhân tại Việt Nam là việc chứng minh mối quan hệ nhân quả giữa quyết định của thuật toán và thiệt hại thực tế, trong khi việc giải thích cơ chế vận hành của AI thường đòi hỏi chuyên gia giám định công nghệ. Ngoài ra, do chưa có án lệ hoặc hướng dẫn xét xử chuyên biệt, các Tòa án vẫn chủ yếu tìm kiếm hành vi vi phạm cụ thể của cá nhân hoặc pháp nhân thay vì đánh giá các khiếm khuyết nội tại của hệ thống AI như một căn cứ độc lập để xác lập trách nhiệm dân sự. Những hạn chế này cho thấy, mặc dù pháp luật Việt Nam đã bước đầu tiếp cận mô hình quản trị AI dựa trên rủi ro, nhưng vẫn cần tiếp tục hoàn thiện cơ chế chứng minh, phân bổ trách nhiệm và hướng dẫn áp dụng pháp luật nhằm bảo đảm khả năng bồi thường hiệu quả đối với các thiệt hại do AI gây ra.

4. Một số kinh nghiệm cho Việt Nam

Thứ nhất, cần cụ thể hóa tiêu chuẩn “khiếm khuyết” và “nghĩa vụ cẩn trọng”. Hoa

Kỳ đang dịch chuyển mạnh mẽ sang việc coi AI là một sản phẩm để áp dụng trách nhiệm nghiêm ngặt. Việt Nam đã thiết lập nguyên tắc này tại Điều 29 *Luật Trí tuệ nhân tạo* năm 2025 nhưng cần các hướng dẫn chi tiết về nội hàm “khiếm khuyết AI”. Do đó, Việt Nam nên ban hành Nghị định hướng dẫn giải thích rõ các loại khiếm khuyết AI, đặc biệt là khiếm khuyết cảnh báo. Ví dụ, một nhà cung cấp AI y tế phải bị coi là có khiếm khuyết nếu họ không công bố các giới hạn về dữ liệu huấn luyện (như dữ liệu thiếu tính đại diện cho một nhóm dân cư nhất định), dẫn đến chẩn đoán sai. Đồng thời, xây dựng tiêu chuẩn “Người vận hành hợp lý” cho các lĩnh vực chuyên môn. Việt Nam có thể tham khảo cách Hoa Kỳ xử lý trách nhiệm y khoa: bác sĩ không bị coi là có lỗi nếu họ tuân thủ đúng quy trình giám sát AI và thực hiện kiểm chứng con người.

Thứ hai, cần xây dựng cơ chế bảo hiểm trách nhiệm AI bắt buộc. Tại Hoa Kỳ, thị trường bảo hiểm đang tiên phong trong việc quản trị rủi ro AI thông qua các sản phẩm “AI Security Riders” và điều chỉnh phí bảo hiểm dựa trên mức độ an toàn của thuật toán. Do đó, đối với các hệ thống AI rủi ro cao, như: ngân hàng, y tế, hạ tầng, Việt Nam nên quy định bảo hiểm trách nhiệm dân sự bắt buộc cho bên triển khai. Điều này giúp san sẻ gánh nặng tài chính khi xảy ra thiệt hại hàng loạt; đồng thời, tạo ra một lớp màng lọc rủi ro: các công ty AI có hệ thống kém an toàn sẽ phải chịu phí bảo hiểm rất cao hoặc không thể mua bảo hiểm, từ đó bị loại bỏ khỏi thị trường.

Thứ ba, nghiên cứu xây dựng Quỹ bồi thường thiệt hại do AI gây ra. Thực tiễn Hoa Kỳ cho thấy, các mô hình quỹ bồi thường theo cơ chế “không lỗi”, tiêu biểu là Quỹ bồi thường thiệt hại do vắc xin đã góp phần bảo đảm quyền lợi của nạn nhân thông qua việc chi trả bồi thường nhanh chóng mà không buộc họ phải chứng minh lỗi của nhà sản xuất hoặc nhà cung cấp. Đổi lại, người yêu cầu bồi thường bị giới hạn quyền khởi kiện đối với một số loại và mức thiệt hại nhất định. Kinh nghiệm này có thể được nghiên cứu vận dụng phù hợp với điều kiện Việt Nam.

Trong bối cảnh *Luật Trí tuệ nhân tạo* năm 2025 đã quy định thành lập Quỹ Phát triển AI quốc gia, có thể xem xét bổ sung chức năng hỗ trợ hoặc bồi thường thiệt hại cho quỹ này, hoặc thành lập một quỹ chuyên biệt để xử lý các trường hợp thiệt hại khó xác định chủ thể chịu trách nhiệm, thiệt hại phát sinh từ các hệ thống AI xuyên biên giới hoặc các trường hợp nạn nhân không có khả năng theo đuổi các vụ kiện quốc tế. Nguồn tài chính của quỹ có thể được hình thành từ một tỷ lệ nhất định trên doanh thu của các doanh nghiệp cung cấp hệ thống AI rủi ro cao, kết hợp với nguồn thu từ xử phạt vi phạm hành chính trong lĩnh vực AI và các nguồn tài chính hợp pháp khác. Cơ chế này không chỉ tăng cường khả năng bồi thường kịp thời cho người bị thiệt hại mà còn góp phần phân bổ rủi ro hợp lý giữa Nhà nước, doanh nghiệp và xã hội trong quá trình phát triển và ứng dụng AI.

Thứ tư, cần giảm bớt khó khăn trong việc chứng minh thông qua cơ chế “Đảo ngược nghĩa vụ”. Một trong những đóng góp lớn nhất của các đạo luật AI mới tại Hoa Kỳ (như Colorado AI Act) và EU AI Act là việc thiết lập các giả định có thể bị bác bỏ về lỗi hoặc khiếm khuyết. Do đó, đối với Việt Nam, trong các vụ kiện liên quan đến AI rủi ro cao, pháp luật nên quy định đảo ngược nghĩa vụ chứng minh. Nghĩa là, một khi nạn nhân chứng minh được có thiệt hại và hệ thống AI có tham gia vào quá trình gây ra thiệt hại đó thì nhà cung cấp/bên triển khai phải chứng minh, hệ thống của họ không có khiếm khuyết hoặc thiệt hại do lỗi cố ý của nạn nhân. Điều này là cần thiết để đối phó với sự bất đối xứng về thông tin giữa người tiêu dùng và các tập đoàn công nghệ.

5. Kết luận

Trách nhiệm dân sự đối với thiệt hại do trí tuệ nhân tạo gây ra đang trở thành một trong những vấn đề trung tâm của pháp luật trong kỷ nguyên số. Nghiên cứu cho thấy, pháp luật Hoa Kỳ đang dịch chuyển từ việc vận dụng các học thuyết truyền thống sang xây dựng cơ chế trách nhiệm phù hợp với đặc tính của AI; đồng thời, tăng cường các công cụ phân bổ và quản

trị rủi ro. Đối với Việt Nam, việc ban hành *Luật Trí tuệ nhân tạo* năm 2025 đã tạo tiền đề quan trọng cho việc hình thành khung pháp lý về trách nhiệm dân sự đối với AI. Tuy nhiên, để bảo đảm tính khả thi trong thực thi, cần tiếp tục hoàn thiện cơ chế pháp lý trên cơ sở kế thừa có chọn lọc kinh nghiệm quốc tế và phù hợp với điều kiện phát triển của Việt Nam. Đây cũng là tiền đề quan trọng để xây dựng môi trường pháp lý an toàn, minh bạch và thúc đẩy phát triển AI có trách nhiệm □

Tài liệu tham khảo:

- Anjali Das (14/01/2026). *Artificial Intelligence Legislative Update*. <https://www.wilsonelser.com/publications/artificial-intelligence-legislative-update>
- Chat GPT Is Eating the World (02/11/2025). *Tracker of Tort Lawsuits v. AI companies*. <https://chatgptiseatingtheworld.com/2025/11/02/tracker-of-tort-lawsuits-v-ai-companies/>
- Durbin (29/9/2025). *Durbin, Hawley Introduce Bill Allowing Victims To Sue AI Companies*. <https://www.durbin.senate.gov/newsroom/press-releases/durbin-hawley-introduce-bill-allowing-victims-to-sue-ai-companies>
- Elizabeth Chiarello, Laura Craig, and Anna Boardman, Sidley Austin LLP (10/01/2024). *Product Liability Considerations For AI-Enabled Medtech*. <https://www.meddeviceonline.com/doc/product-liability-considerations-for-ai-enabled-medtech-0001>
- Frank Young (13/10/2025). *AI as a Product: The Next Frontier in Product Liability Law*. <https://library.law.uic.edu/news-stories/ai-as-a-product-the-next-frontier-in-product-liability-law/>
- Hương Giang (07/5/2025). *AI “gây họa”, ai chịu trách nhiệm? Giải mã pháp lý về thiệt hại do trí tuệ nhân tạo*. <https://xn--luts-tib7365b.net/ai-gay-hoa-ai-chiu-trach-nhiem-giai-ma-phap-ly-ve-thiet-hai-do-tri-tue-nhan-cao.html>
- Thu Giang (10/12/2025). *Việt Nam chính thức có Luật Trí tuệ nhân tạo (AI)*. *Báo Chính phủ*. <https://baochinhphu.vn/viet-nam-chinh-thuc-co-luat-tri-tue-nhan-cao-ai-102251210164948585.htm>
- Gregory Smith, Karlyn D. Stanley, Krystyna Marcinek, Paul Cormarie, Salil Gunashekar (20/11/2024). *Liability for Harms from AI Systems*. https://www.rand.org/pubs/research_reports/RRA3243-4.html
- ICLG (15/06/2026). *USA - Product Liability Laws and Regulations 2026*. <https://iclg.com/practice-areas/product-liability-laws-and-regulations/usa>
- John Villasenor. (31/10/2019). *Products liability law as a way to address AI harms*. <https://www.brookings.edu/articles/products-liability-law-as-a-way-to-address-ai-harms/>
- Kimberly Chew, Hilda Akopyan & Nick Morgan (20/10/2025). *Emerging Legal Challenges: Artificial Intelligence and Product Liability*. <https://www.productlawperspective.com/2025/10/emerging-legal-challenges-artificial-intelligence-and-product-liability/>
- Morrison Foerster (18/6/2025). *Software Gains New Status as a Product Under Strict Liability Law / Morrison Foerster*. <https://www.mofo.com/resources/insights/250618-software-gains-new-status-as-a-product-under-strict-liability-law>
- Philipp Behrendt (07/1/2025). *AI liability - who is accountable when artificial intelligence malfunctions?* <https://www.taylorwessing.com/en/insights-and-events/insights/2025/01/ai-liability-who-is-accountable-when-artificial-intelligence-malfunctions>
- Quinn Emanuel (14/01/2026). *Lead Article: When Machines Discriminate: The Rise of AI Bias Lawsuits - Quinn Emanuel*. <https://www.quinnemanuel.com/the-firm/publications/when-machines-discriminate-the-rise-of-ai-bias-lawsuits/>
- Sommers Schwartz (10/01/2026). *AI Misdiagnosed My Illness. Who is Liable? - Sommers Schwartz*. <https://www.sommerspc.com/blog/2025/03/ai-misdiagnosed-my-illness-who-is-liable/>
- United States Congress (2025). *S. 2937: Aligning Incentives for Leadership, Excellence, and Advancement in Development Act (AI LEAD Act)*. [Congress.gov](https://www.congress.gov)
- U.S. Senate Judiciary Committee (2025). *Aligning Incentives for Leadership, Excellence, and Advancement in Development (AI LEAD) Act - Section - by - Section Summary Senator Dick Durbin (D-IL), Senator Josh Hawley (R-MO)*. <https://www.judiciary.senate.gov/imo/media/doc/Section-by-Section%20-%20AI%20LEAD%20Act.pdf>