

YÊU CẦU THỂ CHẾ MỚI ĐỐI VỚI QUẢN TRỊ RỦI RO TRÍ TUỆ NHÂN TẠO TRONG CUNG ỨNG DỊCH VỤ HÀNH CHÍNH CÔNG

NGUYỄN QUỲNH TRANG*
NGUYỄN QUỲNH HƯƠNG LY**

Ứng dụng trí tuệ nhân tạo (AI) trong cung ứng dịch vụ hành chính công tạo điều kiện tự động hóa, hỗ trợ ra quyết định và cá nhân hóa dịch vụ, nhưng đồng thời phát sinh rủi ro về sai lệch thuật toán, dữ liệu cá nhân, an ninh mạng và trách nhiệm giải trình. Từ năm 2026, Luật Trí tuệ nhân tạo năm 2025 và các văn bản hướng dẫn đã hình thành khung quản lý AI dựa trên mức độ rủi ro tại Việt Nam. Trên cơ sở phân tích pháp lý, đối chiếu ISO 31000:2018, NIST AI RMF 1.0, Khung năng lực số của UNESCO và tham chiếu trường hợp Hà Nội, bài viết đề xuất sử dụng công cụ phân loại rủi ro quốc gia, chuẩn hóa năng lực theo vị trí việc làm, tăng cường đánh giá độc lập và trách nhiệm giải trình trong toàn bộ vòng đời hệ thống AI.

Từ khóa: Trí tuệ nhân tạo; quản trị rủi ro AI; chính quyền địa phương; dịch vụ hành chính công; năng lực số.

The application of artificial intelligence (AI) in the provision of public administrative services enables automation, supports decision-making, and facilitates service personalization, while also creating risks related to algorithmic bias, personal data, cybersecurity, and accountability. Since 2026, the 2025 Law on Artificial Intelligence and its implementing regulations have established a risk-based AI governance framework in Vietnam. Based on legal analysis, comparison with ISO 31000:2018, the NIST AI RMF 1.0, and UNESCO's Digital Competency Framework, and with reference to the case of Hanoi, this article proposes the use of a national risk classification tool, the standardization of competencies by job position, and the strengthening of independent assessment and accountability throughout the entire AI system lifecycle.

Keywords: Artificial intelligence; AI risk governance; local government; public administrative services; digital competencies.

NGÀY NHẬN: 14/4/2026 NGÀY PHẢN BIỆN, ĐÁNH GIÁ: 25/7/2026 NGÀY DUYỆT ĐĂNG: 17/8/2026

Email: nqtrangnqtrang@gmail.com; lylychan234@gmail.com

DOI: <https://doi.org/10.59394/qlnn.368.2026.1599>

1. Đặt vấn đề

Cuộc cách mạng công nghiệp 4.0 đang làm thay đổi phương thức vận hành của bộ máy hành chính nhà nước, trong đó AI nổi lên như công cụ cải cách hành chính công

có tính đột phá. AI đang chuyển từ công cụ hỗ trợ tác nghiệp sang tham gia trực tiếp

* ThS, Học viện Chính trị quốc gia Hồ Chí Minh

** Công ty CP phát triển công nghệ VSHOME

vào quy trình cung cấp dịch vụ công. Tại TP. Hồ Chí Minh và Hà Nội, chatbot tư vấn thủ tục, sàng lọc giấy tờ và hỗ trợ tra cứu thủ tục đã được thử nghiệm hoặc triển khai trên các kênh số (Báo Pháp luật Việt Nam, 16/6/2025; Việt Nga, 07/5/2026). Các ứng dụng này có thể cải thiện tốc độ xử lý và khả năng tiếp cận, tuy nhiên, hiệu quả quản trị không thể được đánh giá chỉ bằng mức độ tự động hóa.

Trong dịch vụ hành chính công, đầu ra AI có thể ảnh hưởng trực tiếp hoặc gián tiếp đến việc xác định tính đầy đủ của hồ sơ, xếp hạng ưu tiên, phát hiện bất thường, xác định đối tượng thụ hưởng hoặc hỗ trợ ban hành quyết định hành chính. Do đó, rủi ro trong quản trị không chỉ là lỗi kỹ thuật mà còn bao gồm thiên kiến, phân biệt đối xử do thuật toán, suy giảm khả năng giải trình, rò rỉ dữ liệu, tấn công mạng và hiện tượng công chức phụ thuộc quá mức vào khuyến nghị của hệ thống. Các nghiên cứu về công cụ chấm điểm tái phạm tại Hoa Kỳ và hệ thống SyRI tại Hà Lan cho thấy, việc sử dụng mô hình thuật toán trong các quyết định tác động đến quyền và lợi ích của cá nhân có thể tạo ra tranh luận nghiêm trọng về tính công bằng, minh bạch và kiểm soát quyền lực (Chouldechova, 2017; Van Bakkum & Zuiderveen Borgesius, 2021).

Khung pháp lý về AI của Việt Nam bao gồm: *Luật Trí tuệ nhân tạo* năm 2025; Nghị định số 142/2026/NĐ-CP của Chính phủ ngày 30/4/2026 cụ thể hóa việc phân loại, rà soát, hồ sơ và đánh giá sự phù hợp của hệ thống AI; Quyết định số 33/2026/QĐ-TTg ngày 30/6/2026 của Thủ tướng Chính phủ ban hành danh mục hệ thống AI có rủi ro cao. Cùng với *Luật Dữ liệu* năm 2024, *Luật Bảo vệ dữ liệu cá nhân* năm 2025, *Luật An ninh mạng* năm 2025 và *Luật Chuyển đổi số* năm 2025. Việc ứng dụng AI trong khu vực công đã được đặt trong một hệ sinh thái

pháp lý rộng hơn về dữ liệu, an toàn, quyền riêng tư và chuyển đổi số.

Vấn đề đặt ra không đơn thuần là chính quyền địa phương có ứng dụng AI hay không mà là chính quyền đó có đủ năng lực để nhận diện, phân loại, kiểm soát và giám sát rủi ro AI trong suốt vòng đời hệ thống hay không. Do các văn bản quy phạm pháp luật mới, thời gian thực thi còn ngắn và chưa có dữ liệu khảo sát đại diện toàn quốc, bài viết không khái quát “thực trạng năng lực công chức” trên phạm vi cả nước. Thay vào đó, bài viết tập trung nhận diện khoảng cách giữa yêu cầu thể chế mới và năng lực căn cơ của chính quyền địa phương, từ đó, đề xuất hàm ý tổ chức thực thi phù hợp với mô hình vận hành chính quyền địa phương hai cấp hiện nay.

2. Cơ sở lý luận và khung phân tích năng lực quản trị rủi ro AI

2.1. Từ “thẩm định rủi ro” đến “quản trị rủi ro AI”

Trong bài viết này, “quản trị rủi ro AI” bao quát hoạt động nhận diện, phân loại, đánh giá, xử lý, giám sát và tái đánh giá rủi ro trong vòng đời hệ thống. “Thẩm định rủi ro” là một bộ phận của quá trình đó, không phải tên của một thủ tục hành chính độc lập. Cách tiếp cận phù hợp với ISO 31000:2018, vốn yêu cầu tích hợp quản lý rủi ro vào quản trị tổ chức (International Organization for Standardization, 2018).

Đối với AI, NIST AI RMF 1.0 tổ chức quản trị rủi ro theo bốn chức năng govern - map - measure - manage, nhấn mạnh: trách nhiệm tổ chức, hiểu bối cảnh sử dụng, đo lường rủi ro và quản lý rủi ro liên tục (Tabassi, 2023). *Đạo luật Trí tuệ Nhân tạo* của Liên minh châu Âu (EU AI Act) cũng sử dụng cách tiếp cận dựa trên rủi ro, trong đó các nghĩa vụ pháp lý tăng lên khi hệ thống có khả năng tác động lớn đến an toàn hoặc

quyền cơ bản (European Parliament & Council of the European Union, 2024). Tại Việt Nam, Nghị định số 142/2026/NĐ-CP xác định ba mức rủi ro: cao, trung bình và thấp; đơn vị triển khai phải phối hợp với nhà cung cấp để rà soát, phân loại lại khi sửa đổi, tích hợp hoặc thay đổi chức năng làm phát sinh rủi ro mới hoặc cao hơn.

2.2. Khung năng lực hai tầng

Năng lực quản trị rủi ro AI không chỉ là kiến thức của từng công chức của chính quyền địa phương. Khung UNESCO nhấn mạnh năng lực con người và thể chế trong chuyển đổi số (UNESCO, 2022). Trong bài viết này, tác giả đề xuất hai tầng năng lực. *Tầng năng lực cá nhân*, gồm: hiểu biết AI và dữ liệu; nhận diện giới hạn, thiên kiến, dấu hiệu bất thường; xác định nghĩa vụ pháp lý; đọc hiểu bằng chứng về độ chính xác và thực hiện giám sát có ý nghĩa. *Tầng năng lực tổ chức*, gồm: phân công trách nhiệm; quy trình, hồ sơ; kiểm thử, giám sát, tái đánh giá; quản lý nhà cung cấp; bảo vệ dữ liệu, an ninh mạng; báo cáo sự cố và đánh giá độc lập.

Để bảo đảm tính kế thừa với cấu trúc ban đầu, hai tầng năng lực được phân tích theo 4 nhóm chức năng: (1) Nhận thức kỹ thuật; (2) Pháp lý và tuân thủ; (3) Tổ chức và quy trình; (4) Giám sát có ý nghĩa của con người và phán quyết độc lập. Bốn nhóm này có quan hệ bổ sung: năng lực cá nhân yếu làm giảm chất lượng kiểm soát, trong khi quy trình tổ chức yếu có thể khiến công chức có năng lực vẫn khó thực hiện trách nhiệm một cách nhất quán.

3. Yêu cầu thể chế mới và những vấn đề đặt ra

(1) Năng lực nhận thức kỹ thuật.

Khung pháp lý mới yêu cầu việc phân loại và quản lý rủi ro dựa vào mục đích sử dụng, mức độ tự động hóa, phạm vi tác

động và khả năng giám sát của con người. Điều này đòi hỏi cán bộ chủ quản và người sử dụng đầu ra AI phải hiểu tối thiểu về dữ liệu đầu vào, độ chính xác, sai số, độ lệch, thay đổi mô hình và sự khác biệt giữa hệ thống hỗ trợ quyết định với hệ thống có khả năng tác động trực tiếp đến quyết định cuối cùng. Tuy nhiên, pháp luật của Việt Nam hiện nay chủ yếu quy định nghĩa vụ của chủ thể mà chưa hình thành một khung năng lực đầu ra thống nhất theo vị trí việc làm đối với đội ngũ công chức địa phương.

Do đó, khoảng cách năng lực cần lưu ý không phải là công chức phải trở thành chuyên gia học máy mà là thiếu “AI literacy” (năng lực nhận biết, sử dụng, đánh giá và làm việc có trách nhiệm với công cụ AI) đủ để đặt câu hỏi đúng, đọc được tài liệu kỹ thuật, nhận biết khi kết quả cần xác minh lại và biết thời điểm phải yêu cầu hỗ trợ chuyên môn. Trường hợp thành phố Hà Nội cho thấy, xu hướng cụ thể hóa trách nhiệm đào tạo về quản trị, đạo đức, pháp lý, an toàn và rủi ro AI, tuy nhiên, đây là quy định của một địa phương cấp tỉnh và chỉ nên được xem như mô hình tham chiếu.

(2) Năng lực pháp lý và tuân thủ.

Nghị định số 142/2026/NĐ-CP đã tạo cơ chế hỗ trợ rõ hơn so với giai đoạn chưa có luật chuyên ngành. Điều 6 quy định phân loại hệ thống AI theo ba mức rủi ro; Điều 10 quy định Bộ Khoa học và Công nghệ thiết lập, duy trì Công cụ hỗ trợ phân loại rủi ro tự động trên Cổng thông tin điện tử một cửa về AI nhằm hỗ trợ nhà cung cấp, bên triển khai tự đánh giá, phân loại hệ thống AI; Điều 11 quy định rà soát, cập nhật mức độ rủi ro; Điều 12 yêu cầu lập hồ sơ phân loại đối với hệ thống rủi ro cao và trung bình. Trường hợp chưa xác định được mức độ rủi ro sau khi sử dụng công cụ, nhà cung cấp có quyền đề nghị Bộ Khoa học và Công nghệ hướng dẫn.

Như vậy, thách thức trong thực thi là khả năng mô tả đúng chức năng hệ thống, chuẩn hóa thông tin đầu vào, nhận diện trường hợp ranh giới, phối hợp với nhà cung cấp và lưu vết căn cứ phân loại. Đối với dịch vụ hành chính công, yêu cầu tuân thủ còn phải được đối chiếu với pháp luật về dữ liệu, bảo vệ dữ liệu cá nhân, an ninh mạng và quy định chuyên ngành. Quyết định số 33/2026/QĐ-TTg tiếp tục làm tăng yêu cầu nhận diện rủi ro khi liệt kê các hệ thống AI rủi ro cao trong nhiều lĩnh vực, trong đó có những hệ thống tự động chấm điểm, phân loại hoặc ra quyết định ảnh hưởng đến hồ sơ, quyền lợi và nguồn lực công.

(3) Năng lực tổ chức và quy trình.

Quản trị rủi ro AI đòi hỏi chuỗi trách nhiệm từ lựa chọn giải pháp, ký hợp đồng, triển khai, vận hành đến ngừng sử dụng. Nghị định số 142/2026/NĐ-CP phân định nghĩa vụ giữa nhà cung cấp và bên triển khai; đồng thời, yêu cầu phân loại lại khi thay đổi làm phát sinh rủi ro mới. Ở cấp tỉnh, khoảng cách năng lực tổ chức dễ xuất hiện là thiếu danh mục hệ thống AI; trách nhiệm giữa nghiệp vụ, công nghệ và nhà cung cấp chưa rõ; tiêu chí kiểm thử, lưu nhật ký và báo cáo sự cố chưa thống nhất.

Trường hợp thành phố Hà Nội cung cấp một minh họa cụ thể về nỗ lực cụ thể hóa yêu cầu quản trị rủi ro AI ở cấp tỉnh. Ngày 21/7/2026, UBND thành phố ban hành Quyết định số 102/2026/QĐ-UBND quy định về bảo đảm an toàn, quản lý rủi ro và biện pháp ứng dụng AI trong quản lý nhà nước trên địa bàn. Quyết định yêu cầu các cơ quan, đơn vị xây dựng quy trình nội bộ đối với AI tạo sinh và tác nhân AI, xác định rõ dữ liệu được phép sử dụng, giới hạn thực thi, trách nhiệm kiểm chứng đầu ra, lưu vết quá trình xử lý và cơ chế chuyển sang cán bộ chuyên trách khi hệ thống tương tác với người dân nhưng không giải quyết đúng yêu cầu. Cách

tiếp cận này thể hiện nỗ lực gắn trách nhiệm tổ chức với giám sát có ý nghĩa của con người; đồng thời, tạo hành lang cho việc phân định vai trò giữa đơn vị nghiệp vụ, đơn vị công nghệ và nhà cung cấp. Tuy nhiên, việc thực thi vẫn phụ thuộc vào năng lực xây dựng và vận hành quy trình nội bộ của từng cơ quan; nếu thiếu danh mục hệ thống AI thống nhất, thiếu tiêu chí kiểm thử và thiếu cơ chế đánh giá độc lập, quy định chi tiết vẫn có nguy cơ dừng lại ở mức hình thức (Ủy ban nhân dân thành phố Hà Nội, 2026).

Kinh nghiệm của Hà Nội cho thấy, chính quyền địa phương có thể chủ động cụ thể hóa yêu cầu pháp lý nhưng cần kết hợp với hỗ trợ kỹ thuật từ Sở Khoa học và Công nghệ cấp tỉnh, chuẩn hóa hồ sơ tuân thủ và tăng cường đào tạo theo vị trí việc làm thì mới bảo đảm tính bền vững của cơ chế giám sát rủi ro.

(4) Giám sát có ý nghĩa của con người và phán quyết độc lập.

Nghị định số 142/2026/NĐ-CP coi khả năng xem xét độc lập, can thiệp, từ chối hoặc thay đổi quyết định của hệ thống là yếu tố quan trọng khi đánh giá rủi ro; trong một số trường hợp chuyển tiếp, bên triển khai phải duy trì cơ chế giám sát thực chất của con người và lưu nhật ký can thiệp. Quyết định số 33/2026/QĐ-TTg cũng xếp nhiều hệ thống tự động ra quyết định hoặc thực thi mà thiếu bước xem xét độc lập của con người vào nhóm rủi ro cao.

Khoảng cách cốt lõi trong trường hợp này là xu hướng thiên lệch do tự động hóa: là xu hướng tâm lý khiến con người tin tưởng và quá phụ thuộc vào các kết quả, đề xuất hoặc quyết định từ hệ thống máy tính, phần mềm hay AI. Về hình thức, công chức vẫn chịu trách nhiệm nhưng trên thực tế gần như mặc nhiên chấp nhận khuyến nghị của AI. Chỉ số tỷ lệ cán bộ đưa ra quyết định

khác AI có thể là một tín hiệu, nhưng không đủ để kết luận mức độ phụ thuộc. Giám sát nên dựa trên tập hợp chỉ số, gồm: tỷ lệ ghi đề và lý do ghi đề; tỷ lệ sai dương, sai âm; sai lệch giữa các nhóm đối tượng; kết quả kiểm tra ngẫu nhiên; tỷ lệ quyết định bị sửa sau phản ánh, khiếu nại; thời gian thực tế dành cho việc xem xét; số sự cố hoặc cảnh báo được phát hiện. Cách đo đa chỉ số phù hợp hơn với yêu cầu map - measure - manage của NIST AI RMF.

4. Một số đề xuất

Thứ nhất, xây dựng khung năng lực quản trị rủi ro AI theo vị trí việc làm. Bộ Khoa học và Công nghệ cần phối hợp với cơ quan quản lý công chức ban hành khung năng lực tối thiểu theo vị trí việc làm, tránh tình trạng mỗi tỉnh xây dựng chuẩn riêng. Ở cấp địa phương, UBND tổ chức đánh giá nhu cầu và phân tầng đào tạo: lãnh đạo, người quyết định đầu tư/chủ quản hệ thống cần có năng lực đọc báo cáo rủi ro và ra quyết định tạm dừng; cán bộ kỹ thuật cần năng lực kiểm thử, giám sát, quản lý dữ liệu và nhà cung cấp; cán bộ trực tiếp xử lý hồ sơ cần năng lực kiểm chứng đầu ra, nhận diện sai lệch và thực hiện giám sát có ý nghĩa.

Thứ hai, tổ chức tuân thủ dựa trên công cụ quốc gia thay vì phát triển công cụ phân loại riêng ở từng địa phương. Sở Khoa học và Công nghệ nên đóng vai trò đầu mối hỗ trợ sử dụng Công cụ phân loại rủi ro của Bộ Khoa học và Công nghệ, xây dựng bộ tình huống mẫu, danh mục câu hỏi kiểm tra và cơ chế chuyển tiếp trường hợp chưa rõ. Mỗi cơ quan cần duy trì sổ đăng ký hệ thống AI, ghi nhận mục đích sử dụng, cơ quan chủ quản, nhà cung cấp, dữ liệu chính, mức rủi ro, thời điểm rà soát và tình trạng tuân thủ. Hồ sơ tuân thủ phải liên kết với yêu cầu về dữ liệu cá nhân, an ninh mạng, bảo mật và pháp luật chuyên ngành.

Thứ ba, chuẩn hóa quản trị liên ngành và đánh giá độc lập theo mức độ rủi ro. Với hệ thống tác động lớn đến quyền, lợi ích hoặc nguồn lực công, cần có sự tham gia của đơn vị nghiệp vụ, công nghệ, pháp chế/an toàn thông tin và khi cần phải có chuyên gia độc lập tham gia đánh giá. Hợp đồng phải quy định quyền tiếp cận tài liệu kỹ thuật, nghĩa vụ thông báo thay đổi mô hình, tiêu chí chất lượng, lưu nhật ký, xử lý sự cố và trách nhiệm khắc phục. Đánh giá định kỳ hệ thống rủi ro cao cần có mức độ độc lập phù hợp, không chỉ do đơn vị vận hành tự thực hiện.

Thứ tư, củng cố cơ chế giám sát con người, giải trình và phản hồi của người dân. Quy trình nghiệp vụ phải xác định rõ điểm nào AI chỉ hỗ trợ, điểm nào bắt buộc con người xem xét và ai chịu trách nhiệm quyết định cuối cùng. Cơ quan chủ quản nên thực hiện kiểm tra mẫu định kỳ, theo dõi bộ chỉ số chất lượng và công bằng, lưu lý do khi cán bộ ghi đề hoặc chấp nhận khuyến nghị trong các trường hợp nhạy cảm. Người dân cần được thông tin phù hợp khi tương tác với AI, được tiếp cận kênh phản ánh, kiến nghị, khiếu nại và yêu cầu xem xét lại theo pháp luật hiện hành; không nên tự thiết lập một cơ chế “miễn trừ trách nhiệm” hoặc thẩm quyền khiếu nại mới nếu chưa có căn cứ pháp lý.

Thứ năm, hình thành cơ chế học hỏi và minh bạch có kiểm soát. UBND có thể công bố định kỳ nhóm ứng dụng AI, mục đích sử dụng, mức rủi ro và thông tin tổng hợp về sự cố trong phạm vi pháp luật cho phép. Sở Khoa học và Công nghệ cần chia sẻ cảnh báo và bài học giữa các cơ quan nhà nước với nhau. Cơ chế báo cáo nội bộ nên bảo vệ cán bộ thực hiện đúng quy trình khi chủ động cảnh báo hoặc kiến nghị tạm dừng hệ thống, hướng tới văn hóa phòng ngừa rủi ro.

5. Kết luận

Khung pháp lý AI của Việt Nam đã chuyển trọng tâm từ khuyến khích ứng dụng sang quản trị theo mức độ rủi ro, kéo theo yêu cầu năng lực mới đối với chính quyền địa phương cấp tỉnh. Năng lực này không thể thu hẹp vào kiến thức kỹ thuật của cán bộ, công chức mà phải bao gồm cả pháp lý, quy trình, dữ liệu, giám sát con người, trách nhiệm giải trình và năng lực tổ chức. Trọng tâm chính sách trong giai đoạn đầu nên là chuẩn hóa năng lực theo vị trí việc làm, tận dụng công cụ phân loại rủi ro cấp quốc gia, thiết lập danh mục và hồ sơ hệ thống AI, tăng cường đánh giá độc lập và bảo đảm giám sát có ý nghĩa của con người. Do nghiên cứu chưa sử dụng khảo sát đại diện toàn quốc, các khoảng cách được nêu là kết quả phân tích thể chế chứ chưa phải kết luận định lượng về năng lực công chức; đây cũng là hướng cần tiếp tục kiểm chứng trong các nghiên cứu thực nghiệm tại địa phương trong thời gian tới.

Tài liệu tham khảo:

- Báo Pháp luật Việt Nam (16/6/2025). *Tro-ly GenAI: “Cánh tay số” giúp thực hiện thủ tục hành chính thuận tiện hơn sau sáp nhập*. <https://baophapluat.vn/tro-ly-genai-canhtay-so-giup-thuc-hien-thu-tuc-hanh-chinh-thuan-tien-hon-sau-sap-nhap-post551927.html>
- Chính phủ (2026). *Nghị định số 142/2026/NĐ-CP ngày 30/4/2026 quy định chi tiết một số điều và biện pháp thi hành Luật Trí tuệ nhân tạo*.
- Chouldechova, A. (2017). *Fair prediction with disparate impact: A study of bias in recidivism prediction instruments*, *Big Data*, 5(2), pp.153-163. <https://doi.org/10.1089/big.2016.0047>
- Estonia Government (2026). *Prime Minister Michal: Estonia to become first country to create digital identities for AI agents*. <https://www.valitsus.ee/en/news/prime-minister-michal-estonia-become-first-count-ry-create-digital-identities-ai-agents>
- European Parliament & Council of the European Union (2024). *Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. Official Journal of the European Union.
- International Organization for Standardization (2018). *ISO 31000:2018 Risk management - Guidelines*. ISO.
- Quốc hội (2024). *Luật Dữ liệu năm 2024*.
- Quốc hội (2025). *Luật Bảo vệ dữ liệu cá nhân năm 2025*.
- Quốc hội (2025). *Luật An ninh mạng năm 2025*.
- Quốc hội (2025). *Luật Trí tuệ nhân tạo năm 2025*.
- Quốc hội (2025). *Luật Chuyển đổi số năm 2025*.
- Tabassi, E (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0) (NIST AI 100-1)*, National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.100-1>
- Thủ tướng Chính phủ (2026). *Quyết định số 33/2026/QĐ-TTg ngày 30/6/2026 ban hành Danh mục hệ thống trí tuệ nhân tạo có rủi ro cao*.
- UNESCO (2022). *Artificial intelligence and digital transformation: Competencies for civil servants*. UNESCO.
- Ủy ban nhân dân thành phố Hà Nội (2026). *Quyết định số 102/2026/QĐ-UBND ngày 21/7/2026 ban hành Quy định về bảo đảm an toàn, quản lý rủi ro và biện pháp ứng dụng trí tuệ nhân tạo trong quản lý nhà nước trên địa bàn thành phố Hà Nội*.
- Van Bekkum, M., & Zuiderveen Borgesius, F. (2021). *Digital welfare fraud detection and the Dutch SyRI judgment*, *Computer Law & Security Review*, 41, 105572. <https://doi.org/10.1016/j.clsr.2021.105572>
- Việt Nga (07/5/2026). *Hà Nội xác lập mô hình quản trị thông minh dựa trên dữ liệu và theo thời gian thực*. <https://hanoimoi.vn/ha-noi-xac-lap-mo-hinh-quan-tri-thong-minh-dua-tren-du-lieu-va-theo-thoi-gian-thuc-748675.html>
- Vũ Thị Hòa, Nguyễn Vũ Thành (2025). *Ứng dụng trí tuệ nhân tạo trong quản lý hành chính công*. *Tạp chí Quản lý nhà nước*. <https://www.quanlynhanuoc.vn/2025/08/19/ung-dung-tri-tue-nhan-tao-trong-quan-ly-hanh-chinh-cong/>