

XU HƯỚNG CHẤP NHẬN AI TRONG HỌC TẬP VÀ LÀM VIỆC TẠI VIỆT NAM: BÀN LUẬN VỀ CÁC NHÂN TỐ CHI PHỐI CHỦ ĐẠO

● NGUYỄN CÔNG THÀNH

Khoa Kinh doanh

Trường Đại học Quản lý và Công nghệ Thành phố Hồ Chí Minh

Email: thanh.nguyencong@umt.edu.vn

TÓM TẮT:

Bài nghiên cứu phân tích các nhân tố tác động đến Ý định chấp nhận AI (AI Adoption intention - AAI) trên cơ sở mô hình đề xuất bao gồm 5 biến độc lập: (1) Giá trị công năng cảm nhận, (2) mức thuận tiện tương tác, (3) năng lực khai thác AI, (4) độ tin cậy đối với AI, và (5) mối lo ngại rủi ro AI. Kết quả định lượng cho thấy, sau 2 lần hiệu chỉnh thang đo, mô hình kiểm định cuối cùng chỉ còn lại 4 nhân tố độc lập, trong khi đó PIS bị loại khỏi mô hình đánh giá. Cuối cùng, mô hình hồi quy có R^2 hiệu chỉnh = 0,627. PFV, PAT, và PARC có tác động dương và ý nghĩa thống kê đến AAI, còn AUC không có ý nghĩa thống kê. Trong đó, PAT là yếu tố ảnh hưởng mạnh nhất. Kết quả này cho thấy nếu muốn người dùng tại Việt Nam sẵn sàng chấp nhận AI nhiều hơn cần làm nổi bật lợi ích thực tế của việc sử dụng AI, đồng thời tăng độ tin cậy và đưa ra cách quản lý rủi ro phù hợp.

Từ khóa: trí tuệ nhân tạo, ý định chấp nhận AI, niềm tin vào AI, rủi ro cảm nhận, học tập và làm việc.

1. Đặt vấn đề

AI đang làm thay đổi cách con người học tập và làm việc khi có thể hỗ trợ viết, tóm tắt, phân tích và gợi ý các ý tưởng. Trong học tập, AI giúp cá nhân hóa việc học và phản hồi nhanh hơn; còn trong công việc, AI giúp giảm bớt các công việc lặp đi lặp lại và giúp nâng cao hiệu quả làm việc. Tuy nhiên, lợi ích đó đi kèm những lo ngại về độ chính xác, thiên lệch nội dung, quyền riêng tư, lệ thuộc vào công nghệ và trách nhiệm sử dụng. Vì vậy, hành vi chấp nhận AI khó được giải thích đầy đủ nếu chỉ dựa trên các mô hình chấp nhận công nghệ truyền thống vốn nhấn mạnh tính hữu ích và tính dễ sử dụng của công nghệ (Davis, 1989; Venkatesh et al., 2003).

Các công trình nghiên cứu gần đây cho thấy, việc chấp nhận AI còn phụ thuộc vào niềm tin, năng lực sử dụng, và nhận thức rủi ro (Dwivedi et al., 2023; Kasneci et al., 2023; Sallam, 2023; Tlili et al., 2023). Tuy nhiên, tại Việt Nam, các bằng chứng thực nghiệm kết hợp đồng thời các nhân tố này trong cả 2 bối cảnh học tập và làm việc vẫn còn tương đối ít. Xuất phát từ khoảng trống nghiên cứu này, bài viết đặt ra 2 câu hỏi: (1) điều gì ảnh hưởng đến ý định chấp nhận AI trong học tập và làm việc tại Việt Nam; (2) liệu mô hình 5 yếu tố tác động có phù hợp với bối cảnh Việt Nam hay không. Mục tiêu của nghiên cứu là kiểm định các nhân tố tác động đến AAI và so sánh mức độ ảnh hưởng của

từng nhân tố đối với các mẫu khảo sát được tiến hành tại TP. Hồ Chí Minh.

2. Giả thuyết và phương pháp nghiên cứu

Dựa trên TAM (mô hình chấp nhận công nghệ) và UTAUT (lý thuyết hợp nhất về chấp nhận và sử dụng công nghệ), nghiên cứu cho rằng ý định sử dụng công nghệ chịu ảnh hưởng bởi cách người dùng đánh giá lợi ích và sự dễ dàng khi sử dụng (Davis, 1989; Venkatesh et al., 2003). Trong bối cảnh với sự hiện diện của AI, mô hình được mở rộng thêm các nhân tố, như: năng lực khai thác AI (AI utilization competency - AUC), độ tin cậy đối với AI (perceived AI trust - PAT), và mối lo ngại rủi ro AI (perceived AI risk concern - PARC) để phản ánh đặc thù của AI. Mô hình lý thuyết của nghiên cứu đề xuất 5 giả thuyết:

- H1: Giá trị công năng cảm nhận (perceived functional value - PFV) tác động cùng chiều đến Ý định chấp nhận AI (AI Adoption intention - AAI).
 - H2: mức thuận tiện tương tác (perceived interaction simplicity - PIS) tác động cùng chiều đến AAI.
 - H3: AUC tác động cùng chiều đến AAI.
 - H4: PAT tác động cùng chiều đến AAI.
 - H5: PARC tác động ngược chiều đến AAI.
- Nghiên cứu sử dụng thiết kế định lượng theo

hướng diễn dịch. Đối tượng khảo sát là học sinh - sinh viên và nhân viên văn phòng tại TP. Hồ Chí Minh, được tiếp cận bằng phương pháp chọn mẫu thuận tiện. Dữ liệu được thu thập bằng bảng hỏi và chủ yếu thực hiện khảo sát trực tuyến. Các biến được đo bằng thang đo Likert 5 mức. Quy trình phân tích bao gồm (1) thống kê mô tả, (2) kiểm định độ tin cậy bằng Cronbach's Alpha, (3) phân tích nhân tố (EFA), và (4) hồi quy tuyến tính đa biến. Mẫu phân tích cuối cùng gồm 92 quan sát hợp lệ. Ban đầu mô hình gồm 5 biến độc lập; tuy nhiên, sau quá trình tinh lọc thang đo, mô hình cuối cùng chỉ còn lại 4 biến phù hợp nhất.

3. Kết quả nghiên cứu và thảo luận

Phân tích EFA cho các biến độc lập được thực hiện nhiều vòng. Sau 2 lần hiệu chỉnh thang đo, mô hình không giữ nguyên 5 nhân tố như đề xuất ban đầu. Cụ thể, PIS không đạt yêu cầu phân biệt và bị loại khỏi mô hình kiểm định cuối cùng. Bảng 1 trình bày ma trận nhân tố xoay hiệu chỉnh cuối cùng.

Kết quả cho thấy, dữ liệu phù hợp để thực hiện EFA với KMO = 0,826 và Bartlett có Sig. = 0,000. 4 nhân tố cuối cùng được trích ra với tổng phương sai

Bảng 1. Ma trận nhân tố xoay hiệu chỉnh cuối cùng

Biến quan sát	PAT	PARC	PFV	AUC
Tôi cảm thấy yên tâm khi sử dụng AI trong học tập/công việc	0,875			
Tôi tin rằng AI có thể hỗ trợ tôi thực hiện nhiệm vụ tốt hơn	0,845			
Tôi tin rằng AI thường đưa ra câu trả lời hữu ích	0,763			
Tôi tin rằng AI có thể hỗ trợ tôi đưa ra quyết định phù hợp	0,634			
Tôi lo dữ liệu của tôi có thể bị lộ/lạm dụng		0,819		
Tôi lo AI có thể đưa thông tin sai		0,813		
Tôi lo về quyền riêng tư khi dùng AI		0,796		
Tôi lo dùng AI quá nhiều làm giảm kỹ năng của tôi		0,753		
AI giúp tôi hoàn thành nhiệm vụ học tập/công việc nhanh hơn			0,875	
AI giúp tôi nâng cao chất lượng kết quả học tập/công việc			0,830	
AI mang lại lợi ích thiết thực trong học tập và công việc của tôi			0,782	
Tôi có thể đánh giá kết quả AI đưa ra là đúng hay sai				0,899
Tôi có thể kiểm soát việc sử dụng AI để tránh lạm dụng hoặc phụ thuộc				0,764
Tôi biết cách đặt câu hỏi/yêu cầu để AI trả lời đúng trọng tâm				0,696
Tôi biết cách sử dụng AI để hỗ trợ học tập/công việc hiệu quả				0,556

Nguồn: Trích xuất từ kết quả phân tích dữ liệu bằng SPSS 28.

giải thích đạt 72,781%. Tất cả hệ số tải nhân tố đều lớn hơn 0,5, cho thấy mức hội tụ chấp nhận được. Như vậy, mô hình lý thuyết ban đầu gồm 5 biến độc lập nhưng mô hình kiểm định cuối cùng chỉ còn 4 nhân tố là PFV, AUC, PAT và PARC. Đối với biến phụ thuộc AAI, EFA cho KMO = 0,812, Bartlett có Sig. = 0,000 và tổng phương sai giải thích đạt 79,339%, xác nhận thang đo đơn hướng và đủ điều kiện cho hồi quy. (Bảng 2)

Mô hình hồi quy có mức phù hợp khá tốt với $R^2 = 0,644$ và R^2 hiệu chỉnh = 0,627; kiểm định F = 39,321 với Sig. = 0,000 cho thấy, mô hình có ý nghĩa thống kê tổng thể. Khi xét từng biến, PFV có tác động dương và có ý nghĩa thống kê đến AAI (Beta = 0,411; Beta chuẩn hóa = 0,395; Sig. = 0,000), cho thấy giá trị sử dụng cảm nhận vẫn là nền tảng của chấp nhận AI. AUC không có ý nghĩa thống kê (Beta = -0,096; Sig. = 0,275), hàm ý kỹ năng sử dụng AI chưa đủ để tạo khác biệt về ý định chấp nhận khi đã kiểm soát các yếu tố khác.

PAT có tác động dương và là nhân tố mạnh nhất trong mô hình (Beta = 0,466; Beta chuẩn hóa = 0,439; Sig. = 0,000). Kết quả này cho thấy, người dùng sẵn sàng chấp nhận AI hơn khi họ cảm thấy yên tâm và tin vào khả năng hỗ trợ nhiệm vụ của công nghệ. Trong khi đó, PARC cũng có tác động dương và có ý nghĩa thống kê (Beta = 0,285; Beta chuẩn hóa = 0,284; Sig. = 0,000), trái với kỳ vọng lý thuyết ban đầu. Điều này cho thấy, người hiểu rõ rủi ro chưa hẳn sẽ ngại dùng AI; ngược lại, họ có thể vẫn chấp

nhận vì đánh giá cao lợi ích và độ tin cậy của công nghệ. Tuy vậy, kết quả này cần được diễn giải thận trọng do nghiên cứu dùng mẫu thuận tiện và quy mô mẫu còn hạn chế.

4. Kết luận

Nghiên cứu cho thấy, mô hình đề xuất ban đầu gồm 5 biến độc lập không được dữ liệu thực nghiệm ủng hộ đầy đủ. Sau 2 lần hiệu chỉnh EFA, PIS bị loại và mô hình cuối cùng còn 4 nhân tố độc lập. Trong mô hình này, PFV, PAT và PARC có tác động dương, có ý nghĩa thống kê đến AAI, còn AUC không có ý nghĩa. PAT là yếu tố có ảnh hưởng mạnh nhất, nhấn mạnh vai trò trung tâm của niềm tin trong chấp nhận AI.

Kết quả cho thấy, các tổ chức cần làm rõ lợi ích của AI, đồng thời xây dựng cách kiểm chứng, hướng dẫn sử dụng rõ ràng và quản lý dữ liệu minh bạch để tăng độ tin cậy. Việc đào tạo AI cũng nên hướng đến cách sử dụng có trách nhiệm như biết đặt câu hỏi đúng, đánh giá kết quả và nhận diện rủi ro, thay vì chỉ dừng lại ở việc học cách dùng công cụ.

Nghiên cứu còn giới hạn ở việc chọn mẫu thuận tiện, địa bàn khảo sát giới hạn tại TP. Hồ Chí Minh, và 92 quan sát hợp lệ; thiết kế cắt ngang cũng chưa phản ánh được sự thay đổi ý định theo thời gian. Trong các nghiên cứu trong tương lai thì các tác giả khác có thể xem xét mở rộng quy mô mẫu, so sánh giữa các nhóm người dùng khác nhau, và phân tích kỹ hơn tác động ngược chiều của PARC bằng những phương pháp chuyên sâu hơn■

Bảng 2. Kết quả mô hình hồi quy các nhân tố tác động đến ý định chấp nhận AI

Nhân tố	Hệ số hồi quy (beta)	Hệ số hồi quy chuẩn hóa (beta chuẩn hóa)	t	Sig.
Hằng số	-0,236		0,606	0,546
Giá trị công năng cảm nhận (PFV)	0,411	0,395	4,814	0,000
Năng lực khai thác AI (AUC)	-0,096	-0,082	-1,098	0,275
Độ tin cậy đối với AI (PAT)	0,466	0,439	5,193	0,000
Mối lo ngại rủi ro AI (PARC)	0,285	0,284	4,220	0,000
Kiểm định	Giá trị			
R	0,802			
R^2	0,644			
R^2 hiệu chỉnh	0,627			
Sai số chuẩn của ước lượng	0,46631			
Thống kê F	39,321			
Sig. của mô hình	0,000			

Nguồn: Trích xuất từ kết quả phân tích dữ liệu bằng SPSS 28.

TÀI LIỆU THAM KHẢO:

- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319-340. <https://doi.org/10.2307/249008>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., et al. (2023). So what if ChatGPT wrote it? Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2010). *Multivariate data analysis* (7th ed.). Pearson Education.
- Kasneci, E., Sessler, K., Küchemann, S., et al. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Nunnally, J. C., & Bernstein, I. H. (1994). *Psychometric theory* (3rd ed.). McGraw-Hill.
- Sallam, M. (2023). ChatGPT utility in healthcare education, research, and practice: Systematic review on the promising perspectives and valid concerns. *Healthcare*, 11(6), 887. <https://doi.org/10.3390/healthcare11060887>
- Tlili, A., Shehata, B., Adarkwah, M. A., Bozkurt, A., Hickey, D. T., Huang, R., & Agyemang, B. (2023). What if the devil is my guardian angel: ChatGPT as a case study of using chatbots in education. *Smart Learning Environments*, 10, 15. <https://doi.org/10.1186/s40561-023-00237-x>
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425-478. <https://doi.org/10.2307/30036540>

Ngày nhận bài: 9/02/2026

Ngày phản biện đánh giá và sửa chữa: 24/02/2026

Ngày chấp nhận đăng bài: 11/3/2026

TRENDS IN AI ADOPTION IN LEARNING AND WORKING IN VIETNAM: A DISCUSSION OF KEY DETERMINANTS

● NGUYEN CONG THANH

School of Business,
University of Management and Technology Ho Chi Minh City
Ho Chi Minh City, Vietnam

ABSTRACT:

This study examines the factors influencing AI Adoption Intention (AAI) based on a proposed model consisting of five independent variables: (1) Perceived Functional Value, (2) Perceived Interaction Convenience, (3) AI Utilization Capability, (4) Trust in AI, and (5) Perceived AI Risk Concerns. The quantitative results indicate that, after two rounds of scale refinement, the final validated model retains only four independent factors, with PIS being excluded from the evaluation model. The final regression model shows an adjusted R^2 of 0.627. PFV, PAT, and PARC have positive and statistically significant effects on AAI, while AUC is not statistically significant. Among these, PAT is the most influential factor. These findings suggest that, to enhance AI adoption among users in Vietnam, it is essential to emphasize the practical benefits of AI usage, strengthen trust in AI, and establish appropriate risk management strategies.

Keywords: artificial intelligence, AI adoption intention, trust in AI, perceived risk, learning and working.