

HỘI CHỨNG “AI ANXIETY” VÀ HÀNH VI THIẾT KẾ ĐÁNH GIÁ CỦA GIẢNG VIÊN ĐẠI HỌC: VAI TRÒ TRUNG GIAN CỦA CẢM NHẬN ĐE DỌA LIÊM CHÍNH HỌC THUẬT

LÊ VĂN

Khoa Quản trị Kinh doanh, Trường Đại học Văn Lang
Faculty of Business Administration, Van Lang University
Email: van.l@vlu.edu.vn

Ngày nhận bài: 22/6/2026

Ngày sửa bài: 9/7/2026

Ngày duyệt đăng bài: 23/7/2026

Tóm tắt:

Nghiên cứu kiểm định mối quan hệ giữa lo âu về trí tuệ nhân tạo (AI Anxiety) và hành vi thiết kế đánh giá thích ứng của giảng viên đại học, trong đó cảm nhận đe dọa liên chính học thuật đóng vai trò trung gian, còn năng lực số về AI và hỗ trợ thể chế đóng vai trò điều tiết. Mô hình dựa trên Lý thuyết Né tránh mối đe dọa công nghệ được kiểm định bằng PLS-SEM trên mẫu 272 giảng viên. Kết quả cho thấy, AI Anxiety tác động dương đến cảm nhận đe dọa ($\beta = 0,486$) và đến hành vi thiết kế đánh giá ($\beta = 0,305$); cảm nhận đe dọa tác động dương đến hành vi ($\beta = 0,436$) và giữ vai trò trung gian một phần. Có 2 vai trò điều tiết chưa đạt ý nghĩa thống kê. Mô hình giải thích 41,2% phương sai hành vi thiết kế đánh giá.

Từ khóa: AI Anxiety, liên chính học thuật, thiết kế đánh giá, PLS-SEM, giảng viên đại học

AI anxiety and university lecturers' assessment-design behavior: The mediating role of perceived threat to academic integrity

Abstract:

This study examines how AI anxiety influences university lecturers' adaptive assessment-design behavior, with perceived threat to academic integrity as a mediator and AI literacy and institutional support as moderators. Grounded in Technology Threat Avoidance Theory and the Transactional Model of Stress and Coping, the model was tested using PLS-SEM with data from 272 lecturers. AI anxiety positively affects perceived threat ($\beta = 0.486$) and assessment-design behavior ($\beta = 0.305$), while perceived threat positively influences behavior ($\beta = 0.436$) and partially mediates the relationship. Neither moderating effect is statistically significant. The model explains 41.2% of the variance in assessment-design behavior. The findings clarify the anxiety-coping mechanism in AI-mediated higher education and provide implications for institutional policies supporting assessment innovation.

Keywords: AI anxiety, academic integrity, assessment design, PLS-SEM, university lecturers

1. Đặt vấn đề

Sự phổ biến của các công cụ trí tuệ nhân tạo (AI) tạo sinh như ChatGPT, Claude hay Gemini từ cuối năm 2022 đã tạo bước ngoặt cho giáo dục đại học. Khả năng tạo văn bản và giải bài tập cũng như bài luận của các mô hình ngôn ngữ lớn làm lung lay nhiều giả định nền tảng về cách kiểm tra, đánh giá năng lực người học. Với giảng viên, AI vừa là công cụ hỗ trợ, và cũng là nguồn gốc của những lo lắng như: bị thay thế nghề nghiệp, sinh viên gian lận, tụt hậu về công nghệ và mất kiểm soát các ranh giới đạo đức trong lớp học.

Trạng thái này được khái niệm hóa thành hội chứng AI Anxiety (Johnson & Verdicchio, 2017; Wang & Wang, 2022) là một cấu trúc đa chiều phản ánh đồng thời nỗi lo về vị thế nghề nghiệp, năng lực thích ứng và hệ quả đạo đức của công nghệ. Trong môi trường đại học, nơi hoạt động đánh giá gắn chặt với tính liêm chính học thuật thì những lo âu này không dừng ở cảm xúc mà có thể thúc đẩy giảng viên thay đổi cách thiết kế đánh giá người học.

Nhiều công bố gần đây ghi nhận làn sóng điều chỉnh việc đánh giá để ứng phó với AI như chuyển đổi hình thức đánh giá sang vấn đáp, thi tại lớp, thiết kế bối cảnh thực hay ban hành chính sách khai báo AI khi làm bài nhưng phần lớn mang tính mô tả hoặc khuyến nghị thực hành (Cotton và cộng sự, 2024). Cơ chế tâm lý kết nối trạng thái lo âu với hành vi tái thiết kế đánh giá vẫn ít được kiểm định trong các nghiên cứu. Cụ thể, chưa rõ AI Anxiety tác động trực tiếp đến hành vi hay chủ yếu qua một đánh giá nhận thức trung gian là sự cảm nhận rằng AI đe dọa tính liêm chính của phương pháp đánh giá hiện hành, đồng thời vai trò của năng lực số và hỗ trợ thể chế cũng chưa được làm rõ.

Trong bối cảnh giáo dục đại học tại Việt Nam, sinh viên tiếp cận các công cụ AI tạo sinh miễn phí một cách rộng rãi, trong khi phần lớn cơ sở đào tạo mới bắt đầu xây dựng chính sách và năng lực đánh giá trong môi trường có AI. Giảng viên vì thế thường phải tự xoay xở điều chỉnh cách kiểm tra dựa trên cảm nhận cá nhân về rủi ro, khiến việc hiểu rõ động lực tâm lý đằng sau hành vi này có giá trị trực tiếp cho quản lý đào tạo.

Nghiên cứu này xây dựng và kiểm định một mô hình tích hợp nhằm để (i) đánh giá tác động trực tiếp của AI Anxiety đến hành vi thiết kế đánh giá, (ii) kiểm định vai trò trung gian của cảm nhận đe dọa liêm chính học thuật, và (iii) xem xét vai trò điều tiết của năng lực số về AI và hỗ trợ thể chế. Nghiên cứu kết nối lý thuyết Né tránh mối đe dọa công nghệ với Mô hình giao dịch về căng thẳng và ứng phó để xây dựng mô hình chuỗi lo âu → đánh giá đe dọa → ứng phó trong một bối cảnh nghề nghiệp cụ thể là giảng viên đại học. Kết quả cung cấp cơ sở cho các trường đại học trong việc hỗ trợ giảng viên đổi mới cách thức đánh giá một cách chủ động.

2. Cơ sở lý thuyết và giả thuyết

AI Anxiety là trạng thái lo âu và căng thẳng của cá nhân trước sự phát triển và ứng dụng của AI (Johnson & Verdicchio, 2017). Wang và Wang (2022) thao tác hóa khái niệm này thành cấu trúc đa chiều qua thang đo AI Anxiety. Nghiên cứu khái niệm hóa AI Anxiety là cấu trúc bậc 2 gồm 4 chiều, đó là lo bị thay thế nghề nghiệp, về liêm chính học thuật, thiếu năng lực công nghệ và mất kiểm soát về đạo đức. Lý thuyết Né tránh mối đe dọa công nghệ (Liang & Xue, 2009) lập luận khi nhận thức công nghệ là mối đe dọa thì cá nhân hình thành động cơ né tránh và thực hiện hành vi thích ứng thông qua hai quá trình đánh giá mối đe dọa và đánh giá khả năng ứng phó. Bổ trợ cho khung lý thuyết trên thì Lazarus và Folkman (1984) cho rằng lo âu là kết quả của đánh giá nhận thức về mối đe dọa và từ đó cá nhân triển khai chiến lược ứng phó, việc tái thiết kế đánh giá được xem là ứng phó tập trung vào vấn đề trước áp lực do AI tạo ra.

Khi giảng viên lo âu về AI thì họ có xu hướng diễn giải sự hiện diện của công cụ AI trong tay sinh viên như mối đe dọa đối với tính công bằng và tin cậy của việc đánh giá. Theo logic đánh giá mối đe dọa thì mức lo âu càng cao thì cảm nhận đe dọa liên chính càng lớn. Vì vậy, lo âu cũng có thể trực tiếp thúc đẩy hành vi thích ứng như một phản ứng phòng vệ chủ động. Khi cảm nhận các hình thức đánh giá hiện tại dễ bị AI vô hiệu hóa, giảng viên có động cơ để tái thiết kế đánh giá nhằm khôi phục giá trị đo lường cho bài kiểm tra. Từ đó, cảm nhận đe dọa được kỳ vọng là cầu nối tâm lý chuyển hóa lo âu thành hành động. Ngoài ra, theo quá trình đánh giá khả năng ứng phó, nguồn lực cá nhân (năng lực số về AI) và tổ chức (hỗ trợ thể chế) được kỳ vọng điều chỉnh cách lo âu và cảm nhận đe dọa chuyển thành hành vi. 6 giả thuyết được đề xuất như sau:

H1: AI Anxiety có tác động cùng chiều đến cảm nhận đe dọa liên chính học thuật.

H2: AI Anxiety có tác động cùng chiều đến hành vi thiết kế đánh giá thích ứng.

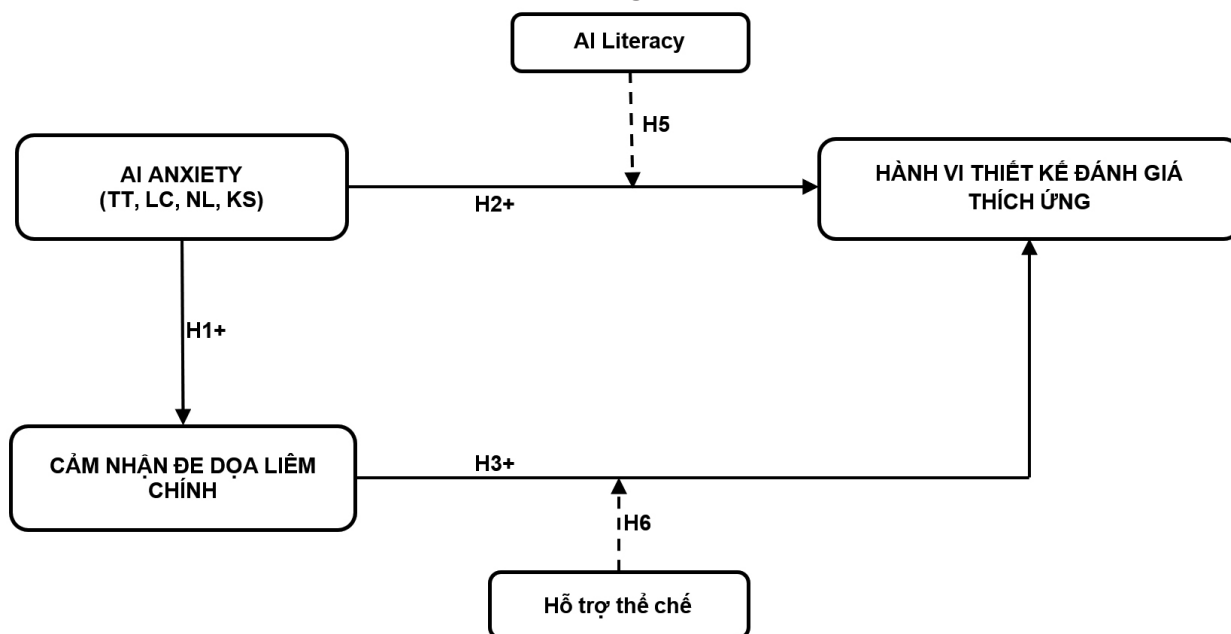
H3: Cảm nhận đe dọa liên chính có tác động cùng chiều đến hành vi thiết kế đánh giá thích ứng.

H4: Cảm nhận đe dọa liên chính đóng vai trò trung gian giữa AI Anxiety và hành vi thiết kế đánh giá.

H5: Năng lực số về AI làm suy giảm tác động của AI Anxiety đến hành vi thiết kế đánh giá.

H6: Hỗ trợ thể chế điều tiết mối quan hệ giữa cảm nhận đe dọa và hành vi thiết kế đánh giá.

Hình 1. Mô hình nghiên cứu đề xuất



3. Phương pháp nghiên cứu

Nghiên cứu sử dụng thiết kế định lượng cắt ngang với bảng hỏi tự điền. Sau khi làm sạch dữ liệu cỡ mẫu phân tích là 272 giảng viên, phù hợp với khuyến nghị cho PLS-SEM có cấu trúc bậc 2 và biến điều tiết. Mẫu đa dạng về giới tính (nam 44,1%; nữ 53,3%), nhóm tuổi (30-39: 37,5%; 40-49: 36,4%), thâm niên (11-20 năm: 33,8%; 5-10 năm: 29,4%), học vị (thạc sĩ 53,3%; tiến sĩ 36,8%) và nhóm ngành. Tất cả khái niệm được đo bằng thang Likert 5 mức. AI Anxiety là cấu trúc bậc 2 gồm 4 chiều với 16 biến quan sát, kế thừa Thang đo AI Anxiety (Wang & Wang, 2022) và bổ sung chiều liên chính dựa trên nghiên cứu của Cotton và cộng sự (2024). Cảm nhận đe dọa liên chính (4 biến) vận dụng thang đo của Perceived

Threat của TTAT (Liang & Xue, 2009). Hành vi thiết kế đánh giá gồm 12 biến trên 3 khía cạnh (chuyển đổi hình thức, tái thiết kế đề, chính sách sử dụng AI), phát triển dựa trên khung đánh giá xác thực của Wiggins (1998). Năng lực số về AI (5 biến) và hỗ trợ thể chế (4 biến) tham chiếu nhân tố Điều kiện thuận lợi trong UTAUT (Venkatesh và cộng sự, 2003). Giới tính, tuổi, thâm niên, ngành và tần suất sử dụng AI là biến kiểm soát.

Mô hình được kiểm định bằng PLS-SEM. Cấu trúc bậc 2 AI Anxiety được xử lý theo cách tiếp cận 2 giai đoạn. Mô hình đo lường được đánh giá qua độ tin cậy (Cronbach's α và $CR \geq 0,70$), giá trị hội tụ ($AVE \geq 0,50$) và giá trị phân biệt ($HTMT < 0,85$). Mô hình cấu trúc được đánh giá qua hệ số đường dẫn chuẩn hóa, R^2 , f^2 và Q^2 . Ý nghĩa thống kê được ước lượng bằng bootstrapping 5.000 mẫu, kiểm định trung gian dựa trên tác động gián tiếp bootstrap (Zhao và cộng sự, 2010) và kiểm định điều tiết sử dụng biến tương tác.

4. Kết quả nghiên cứu

Mô hình đo lường (Bảng 1) đạt yêu cầu: Cronbach's α từ 0,77 đến 0,86, CR từ 0,86 đến 0,90 (đều vượt 0,70), AVE từ 0,60 đến 0,71 (vượt 0,50), hệ số tải ngoài đều lớn hơn 0,75. 2 cấu trúc bậc 2 của AI Anxiety ($\alpha = 0,77$; $CR = 0,86$; $AVE = 0,60$) và hành vi thiết kế đánh giá ($\alpha = 0,80$; $CR = 0,88$; $AVE = 0,71$) đều đạt chuẩn. Giá trị phân biệt được khẳng định với các chỉ số HTMT giữa cảm nhận đe dọa, năng lực số và hỗ trợ thể chế đều dưới 0,12, thấp hơn nhiều so với ngưỡng 0,85.

Bảng 1. Độ tin cậy và giá trị hội tụ

Khái niệm	Số biến	α	CR	AVE
AI Anxiety (bậc hai)	4 chiều	0,774	0,855	0,596
Cảm nhận đe dọa liên chính (DD)	4	0,844	0,896	0,682
Hành vi thiết kế đánh giá (bậc hai)	3 chiều	0,797	0,881	0,711
Năng lực số về AI (AL)	5	0,855	0,897	0,634
Hỗ trợ thể chế (IS)	4	0,849	0,898	0,688

Kết quả mô hình cấu trúc với bootstrapping 5.000 mẫu (Bảng 2) cho thấy, 3 tác động chính đều được ủng hộ ở mức ý nghĩa cao. AI Anxiety tác động dương đến cảm nhận đe dọa ($\beta = 0,486$; $t = 9,67$; $p < 0,001$, ủng hộ H1) và đến hành vi thiết kế đánh giá ($\beta = 0,305$; $t = 5,57$; $p < 0,001$, ủng hộ H2); cảm nhận đe dọa tác động dương đến hành vi ($\beta = 0,436$; $t = 7,74$; $p < 0,001$, ủng hộ H3). Tác động gián tiếp của AI Anxiety đến hành vi qua cảm nhận đe dọa có ý nghĩa (hiệu ứng gián tiếp = 0,212; $p < 0,001$; khoảng tin cậy 95%). Vì cả tác động trực tiếp lẫn gián tiếp đều có ý nghĩa cho nên cảm nhận đe dọa đóng vai trò trung gian một phần, với tỷ lệ phương sai giải thích qua trung gian (VAF) khoảng 41% trên tổng tác động 0,517 (H4 được ủng hộ).

Ngược lại, giả thuyết điều tiết không được ủng hộ. Tác động điều tiết của năng lực số về AI tác động của AI Anxiety $\rightarrow$ hành vi không có ý nghĩa ($\beta = -0,042$; $t = 0,83$; $p = 0,408$, bác bỏ H5). Tác động điều tiết của hỗ trợ thể chế lên tác động của cảm nhận đe dọa $\rightarrow$ hành vi có dấu dương nhưng chỉ tiệm cận ngưỡng ý nghĩa 5% ($\beta = 0,088$; $t = 1,81$; $p = 0,070$), do đó H6 chưa được ủng hộ.

Về năng lực giải thích và dự báo của mô hình giải thích 23,6% phương sai của cảm nhận đe dọa ($R^2 = 0,236$) và 41,2% phương sai của hành vi thiết kế đánh giá ($R^2 = 0,412$). Các hệ số tác động cho thấy AI Anxiety ảnh hưởng mạnh đến cảm nhận đe dọa ($f^2 = 0,309$) và ảnh hưởng vừa đến hành vi ($f^2 = 0,121$), trong khi cảm nhận đe dọa ảnh hưởng vừa đến mạnh lên hành vi ($f^2 = 0,246$). Các giá trị Q^2 dương (0,225 và 0,398) khẳng định độ liên quan dự báo tốt. Điểm trung bình các câu trúc trên thang 1-5 đều xoay quanh mức 3,0. Điều đó cho thấy, giảng viên trong mẫu ở trạng thái lo âu và mức độ điều chỉnh đánh giá ở mức độ trung bình.

Bảng 2. Kết quả kiểm định giả thuyết (bootstrapping 5.000 mẫu)

Giả thuyết	Quan hệ	β	t	p	Kết luận
H1	AI Anxiety $\rightarrow$ Cảm nhận đe dọa	0,486	9,67	< 0,001	Ủng hộ
H2	AI Anxiety $\rightarrow$ Hành vi	0,305	5,57	< 0,001	Ủng hộ
H3	Cảm nhận đe dọa $\rightarrow$ Hành vi	0,436	7,74	< 0,001	Ủng hộ
H4	AI Anxiety $\rightarrow$ Đe dọa $\rightarrow$ Hành vi (gián tiếp)	0,212	5,75	< 0,001	Ủng hộ (một phần)
H5	AI Literacy $\times$ (AI Anxiety $\rightarrow$ Hành vi)	-0,042	0,83	0,408	Không ủng hộ
H6	Hỗ trợ thể chế $\times$ (Đe dọa $\rightarrow$ Hành vi)	0,088	1,81	0,070	Không ủng hộ

5. Thảo luận và hàm ý

Kết quả xác nhận một cơ chế tâm lý là sự lo âu về AI không chỉ là trạng thái cảm xúc thụ động mà là động lực thúc đẩy hành vi nghề nghiệp. Việc AI Anxiety làm tăng cảm nhận đe dọa (H1) và cảm nhận đe dọa thúc đẩy hành vi (H3) phù hợp với logic đánh giá mối đe dọa của Liang và Xue (2009) và chuỗi đánh giá nhận thức ứng phó của Lazarus và Folkman (1984). Điều này có nghĩa là giảng viên lo âu diễn giải sự hiện diện của AI như mối đe dọa đối với sự tin cậy của việc đánh giá người học và phản ứng bằng cách tập trung vào vấn đề là tái thiết kế cách kiểm tra đánh giá.

Đóng góp lý thuyết nổi bật là làm rõ vai trò trung gian một phần của cảm nhận đe dọa (H4). Việc tác động trực tiếp vẫn có ý nghĩa bên cạnh đường gián tiếp cho thấy lo âu vận hành theo 2 cơ chế song song (i) qua đánh giá nhận thức có chủ đích về đe dọa liên chính và (ii) qua phản ứng phòng vệ tương đối tự động. Điều này mở rộng AI Anxiety từ một khái niệm mô tả sang một biến giải thích, đồng thời định vị liên chính học thuật như nút thắt nhận thức trung tâm liên kết lo âu với hành động.

Việc 2 giả thuyết điều tiết không được ủng hộ cũng có ý nghĩa. Kỳ vọng năng lực số làm dịu tác động của lo âu không được xác nhận, gợi ý rằng ngay cả giảng viên thành thạo AI vẫn điều chỉnh đánh giá vì sự lo âu có thể là hiểu quá sâu về AI của giảng viên thì càng lo âu mức độ đề thi truyền thống dễ bị vô hiệu hóa. Tương tự, hỗ trợ thể chế chỉ điều tiết tiệm cận ý nghĩa, cho thấy các nguồn lực tổ chức hiện có thể chưa đủ mạnh hoặc chưa đúng hướng, nhất quán với quan sát rằng nhiều trường vẫn đang ở giai đoạn đầu xây dựng chính sách về AI (Cotton và cộng sự, 2024).

Vấn đề liên chính học thuật là điểm then chốt cần được can thiệp và nhà trường nên hỗ trợ giảng viên chuyển hóa cảm nhận đe dọa thành đổi mới phong cách đánh giá có chất lượng, như cung cấp nguồn lực

thiết kế đánh giá xác thực, ban hành hướng dẫn rõ ràng về khai báo và sử dụng AI, và xây dựng cộng đồng giảng viên thực hành AI. Vì hỗ trợ thể chế chưa phát huy vai trò điều tiết rõ rệt nên các trường phải rà soát để hỗ trợ thực chất hơn và gắn với nhu cầu đánh giá cụ thể của từng ngành học.

Nghiên cứu này có một số điểm hạn chế cần lưu ý (i) phân tích chỉ dựa trên dữ liệu từ thiết kế cắt ngang nên bị hạn chế suy luận nhân quả nên một nghiên cứu theo chiều dọc sẽ củng cố cơ chế trung gian trên. (ii) Mô hình có thể mở rộng bằng cách tách tác động của 4 chiều AI Anxiety và phân tích đa nhóm theo ngành, thâm niên hoặc tần suất sử dụng AI.

6. Kết luận

Nghiên cứu xây dựng và kiểm định một mô hình tích hợp giải thích cơ chế qua đó AI Anxiety chuyển hóa thành hành vi thiết kế đánh giá thích ứng của giảng viên đại học. Lo âu về AI vừa tác động trực tiếp đến hành vi, vừa tác động gián tiếp thông qua cảm nhận đe dọa liên chính học thuật giữ vai trò trung gian một phần, trong khi vai trò điều tiết của năng lực số và hỗ trợ thể chế chưa được xác nhận. Bằng việc kết nối Lý thuyết Né tránh mối đe dọa công nghệ với Mô hình giao dịch về căng thẳng và ứng phó, nghiên cứu định vị liên chính học thuật như nút thắt nhận thức trung tâm liên kết lo âu với hành động, đồng thời gợi mở rằng hỗ trợ thể chế cần đi vào thực chất để phát huy vai trò trong bối cảnh giáo dục đại học ở kỷ nguyên AI ■

TÀI LIỆU THAM KHẢO:

- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228-239. <https://doi.org/10.1080/14703297.2023.2190148>
- Johnson, D. G., & Verdicchio, M. (2017). AI anxiety. *Journal of the Association for Information Science and Technology*, 68(9), 2267-2270. <https://doi.org/10.1002/asi.23867>
- Lazarus, R. S., & Folkman, S. (1984). *Stress, appraisal, and coping*. Springer.
- Liang, H., & Xue, Y. (2009). Avoidance of information technology threats: A theoretical perspective. *MIS Quarterly*, 33(1), 71-90. <https://doi.org/10.2307/20650279>
- 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>
- Wang, Y.-Y., & Wang, Y.-S. (2022). Development and validation of an artificial intelligence anxiety scale: An initial application in predicting motivated learning behavior. *Interactive Learning Environments*, 30(4), 619-634. <https://doi.org/10.1080/10494820.2019.1674887>
- Wiggins, G. (1998). *Educative Assessment. Designing Assessments To Inform and Improve Student Performance*. Jossey-Bass Publishers, 350 Sansome Street, San Francisco, CA 94104.
- Zhao, X., Lynch, J. G., & Chen, Q. (2010). Reconsidering Baron and Kenny: Myths and truths about mediation analysis. *Journal of Consumer Research*, 37(2), 197-206. <https://doi.org/10.1086/651257>