

ẢNH HƯỞNG CỦA ÁP LỰC HỌC TẬP ĐẾN HÀNH VI SỬ DỤNG AI VÀ NGUY CƠ VI PHẠM LIÊM CHÍNH HỌC THUẬT: PHÂN TÍCH THỰC NGHIỆM TỪ DỮ LIỆU KHẢO SÁT SINH VIÊN TRƯỜNG ĐẠI HỌC MỞ HÀ NỘI

Nguyễn Anh Hoàn¹
Email: anhhoan5880@hou.edu.vn

Ngày tòa soạn nhận được bài báo: 02/06/2025

Ngày phản biện đánh giá: 15/08/2025

Ngày bài báo được duyệt đăng: 29/08/2025

DOI: 10.59266/houjs.2025.670

Tóm tắt: Trong bối cảnh trí tuệ nhân tạo (AI) ngày càng được tích hợp sâu vào hoạt động dạy - học ở bậc đại học, việc hiểu rõ các yếu tố ảnh hưởng đến hành vi sử dụng AI của sinh viên và nguy cơ vi phạm liêm chính học thuật trở nên cấp thiết. Nghiên cứu này nhằm phân tích mối quan hệ giữa áp lực học tập, ý định sử dụng AI, hành vi sử dụng AI và hành vi vi phạm liêm chính học thuật của sinh viên. Trên cơ sở khảo sát hơn 2.500 sinh viên thuộc nhiều chuyên ngành và năm học khác nhau tại Trường Đại học Mở Hà Nội, kết quả cho thấy áp lực học tập có ảnh hưởng đáng kể đến ý định sử dụng AI; ý định này tác động mạnh mẽ đến hành vi sử dụng thực tế; và cả hai yếu tố đều góp phần làm gia tăng nguy cơ vi phạm chuẩn mực đạo đức học thuật. Nghiên cứu đề xuất các khuyến nghị ở cấp nhà trường, giảng viên và sinh viên nhằm nâng cao nhận thức, kỹ năng và đạo đức sử dụng AI trong học tập, qua đó góp phần xây dựng môi trường học thuật minh bạch, trách nhiệm và thích ứng với chuyển đổi số trong giáo dục đại học.

Từ khóa: áp lực học tập, hành vi sử dụng ai, nhận thức đạo đức học thuật, trí tuệ nhân tạo, vi phạm liêm chính học thuật

I. Đặt vấn đề

Sự phát triển nhanh của trí tuệ nhân tạo (AI), đặc biệt các mô hình ngôn ngữ như ChatGPT, tạo ra cơ hội học tập nhưng đồng thời đặt ra thách thức về đạo đức và liêm chính học thuật trong giáo dục đại học. Việc sinh viên sử dụng AI thiếu tuân thủ quy tắc trích dẫn có thể làm gia tăng

nguy cơ vi phạm, đòi hỏi nghiên cứu sâu về các yếu tố tác động. Bài viết này phân tích ảnh hưởng của áp lực học tập đến hành vi sử dụng AI và nguy cơ vi phạm liêm chính học thuật, đồng thời xem xét vai trò trung gian của ý định sử dụng AI và tác động điều tiết của nhận thức đạo đức học thuật. Khảo sát được thực hiện

¹ Trường Đại học Mở Hà Nội

với 2.538 sinh viên Trường Đại học Mở Hà Nội, nhằm cung cấp bằng chứng thực nghiệm và đề xuất chính sách thúc đẩy việc sử dụng AI có trách nhiệm trong môi trường học thuật số.

II. Cơ sở lý thuyết

Nghiên cứu dựa trên ba nền tảng chính: Thuyết hành vi có kế hoạch - TPB (Ajzen, 1991), Lý thuyết căng thẳng - thích ứng (Lazarus & Folkman, 1984), và Khung liên chính học thuật số (Sagge, 2024; Bennett & Abusalem, 2024).

Theo TPB, ý định hành vi hình thành từ thái độ, chuẩn mực xã hội và nhận thức kiểm soát, giúp lý giải xu hướng sinh viên sử dụng AI nhằm giảm áp lực và nâng cao hiệu quả học tập (Jain et al., 2022; Đặng, 2024). Lý thuyết căng thẳng - thích ứng cho rằng khi áp lực vượt quá nguồn lực cá nhân, AI có thể trở thành cơ chế đối phó. Nhiều nghiên cứu

quốc tế gần đây cũng khẳng định stress học tập gắn liền với xu hướng phụ thuộc AI, trong khi môi trường số hóa hiệu quả và năng lực sử dụng AI giúp giảm tác động tiêu cực (Zhang, 2024; Norabuena-Figueroa et al., 2025; Kim, 2024).

Như vậy, AI vừa là công cụ hỗ trợ vừa là cách ứng phó với áp lực, đồng thời đặt ra yêu cầu tích hợp nguyên tắc đạo đức và liên chính trong môi trường số. Từ cơ sở trên, nghiên cứu đề xuất mô hình quan hệ:

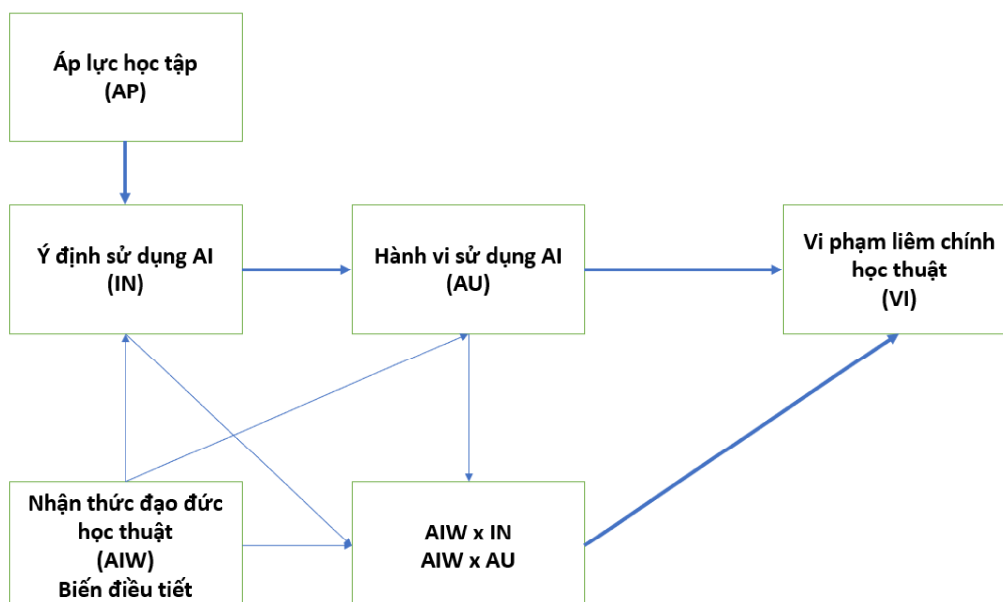
(i) Áp lực học tập (AP) → Ý định sử dụng AI (IN);

(ii) Ý định (IN) → Hành vi sử dụng AI (AU);

(iii) IN & AU → Vi phạm liên chính học thuật (VI);

(iv) Nhận thức đạo đức học thuật (AIW) điều tiết quan hệ AU → VI.

Hình 1. Mô hình nghiên cứu đề xuất



(Nguồn: Tác giả xây dựng)

Trên cơ sở đó, nghiên cứu xây dựng và kiểm định 9 giả thuyết trên mẫu 2.538 sinh viên, bao gồm quan hệ trực tiếp, trung gian và điều tiết (Bảng 1).

Bảng 1. Hệ thống giả thuyết nghiên cứu

Mã	Nội dung	Kiểm định
H1	AP $\rightarrow$ IN (thuận chiều)	Hồi quy đơn
H2	IN $\rightarrow$ AU (thuận chiều)	Hồi quy đơn
H3	IN $\rightarrow$ VI (thuận chiều)	Hồi quy đơn
H4	AU $\rightarrow$ VI (thuận chiều)	Hồi quy đơn
H5	AP $\rightarrow$ VI (trực tiếp)	Hồi quy đa biến
H6	IN trung gian: AP $\rightarrow$ AU	Trung gian (PROCESS 4)
H7	IN & AU trung gian chuỗi: AP $\rightarrow$ VI	Trung gian (PROCESS 6)
H8	AIW điều tiết: IN $\rightarrow$ VI	Điều tiết (PROCESS 1)
H9	AIW điều tiết: AU $\rightarrow$ VI	Điều tiết (PROCESS 1)

(Nguồn: Tác giả tổng hợp từ mô hình nghiên cứu)

III. Phương pháp nghiên cứu

Nghiên cứu mang tính mô tả - giải thích - kiểm định giả thuyết, dựa trên ba nền tảng lý thuyết đã trình bày ở phần trước. Mô hình xem xét các quan hệ trực tiếp, trung gian và điều tiết, phân tích bằng hồi quy hiện đại.

3.1. Thiết kế nghiên cứu

Nghiên cứu mang tính mô tả - giải thích - kiểm định giả thuyết, dựa trên ba nền tảng lý thuyết: Thuyết hành vi có kế hoạch (Ajzen, 1991), Lý thuyết căng thẳng - thích ứng (Lazarus & Folkman, 1984) và Khung liên chính học thuật số (Sagge, 2024; Bennett & Abusalem, 2024). Mô hình xem xét các quan hệ trực tiếp, trung gian và điều tiết, được phân tích bằng kỹ thuật hồi quy hiện đại.

3.2. Đối tượng và mẫu nghiên cứu

Mẫu gồm 2.538 sinh viên chính quy Trường Đại học Mở Hà Nội thuộc nhiều ngành và khóa học, chọn theo phương pháp thuận tiện có kiểm soát. Dữ liệu được thu thập bằng bảng hỏi trực tuyến trong tháng 5-6/2025.

3.3. Thang đo

Các biến quan sát được đo bằng thang Likert 5 mức, kế thừa từ các nghiên

cứu trước và được điều chỉnh cho phù hợp với bối cảnh Việt Nam. Hệ thống thang đo gồm năm nhóm chính:

Áp lực học tập (AP): 4 biến quan sát phản ánh tình trạng căng thẳng do hạn nộp, khối lượng môn học lớn, áp lực điểm số và xu hướng dùng AI thay cho tự học (Nguyen et al., 2024; Fatima et al., 2021; Zhytnyk & Gryniuk, 2024; Steare et al., 2023).

Ý định sử dụng AI (IN): 4 biến quan sát thể hiện dự định tăng cường sử dụng AI, niềm tin vào hiệu quả, sẵn sàng giới thiệu cho bạn học và coi AI là công cụ tất yếu (Jain et al., 2022; tác giả tự xây dựng).

Hành vi sử dụng AI (AU): 3 biến quan sát phản ánh mức độ thường xuyên dùng AI, sử dụng khi gấp hoặc nhằm nâng cao hiệu suất học tập (Pitts et al., 2025).

Vi phạm liên chính học thuật (VI): 5 biến quan sát liên quan đến hành vi nộp bài do AI viết, sao chép không trích dẫn, xem nhẹ tính vi phạm, bỏ qua gian lận của bạn học và sẵn sàng sử dụng AI nếu không bị phát hiện (Ngo et al., 2024; Chan et al., 2025; NMU, n.d.).

Nhận thức đạo đức học thuật (AIW): 5 biến quan sát phản ánh việc được học về đạo đức học thuật, hiểu biết về đạo văn,

quan điểm về công khai sử dụng AI, nhận thức về công cụ phát hiện và sự ủng hộ quy định rõ ràng (University of Oxford, 2025; Overono, 2025; Weber-Wulff et al., 2023; Chan, 2023).

3.4. Phân tích dữ liệu

Dữ liệu được xử lý bằng SPSS 27.0 và macro PROCESS 4.2 (Hayes, 2018). Các bước gồm: thống kê mô tả; kiểm định độ tin cậy và giá trị thang đo (Cronbach's Alpha, EFA); hồi quy tuyến tính kiểm định H1-H5; phân tích trung gian (PROCESS Model 4 và 6, H6-H7); và phân tích điều tiết (PROCESS Model 1, H8-H9).

IV. Kết quả nghiên cứu và thảo luận

4.1. Đặc điểm mẫu và hành vi sử dụng AI của sinh viên

Khảo sát được tiến hành với 2.538 sinh viên Trường Đại học Mở Hà Nội. Về giới tính, mẫu gồm 1.735 nam (68,4%), 800 nữ (31,5%) và 3 thuộc nhóm khác (0,1%). Xét theo năm học, sinh viên năm nhất chiếm tỷ lệ cao nhất với 918 người (36,2%), tiếp đến là năm hai 698 người (27,5%), năm ba 548 người (21,6%), năm tư 370 người (14,6%) và năm năm 4 người (0,2%). Về phân bố ngành học, tỷ lệ lớn nhất thuộc Khoa Kinh tế (16,2%), Khoa Tiếng Trung Quốc (14,0%), Khoa Tài chính - Ngân hàng (12,6%) và Khoa Du lịch (12,6%). Các đơn vị khác chiếm tỷ lệ thấp hơn, gồm Khoa Công nghệ thông tin (9,8%), Khoa Tạo dáng công nghiệp (9,2%), Khoa Luật (8,0%), Khoa Điện - Điện tử (7,7%), Khoa Tiếng Anh (7,2%) và Viện CNSH & CNTP (2,7%).

Hành vi sử dụng AI: Kết quả cho thấy hơn 90% sinh viên đã sử dụng AI trong học tập, chủ yếu cho các hoạt động

như tra cứu tài liệu, viết báo cáo, dịch thuật và tóm tắt nội dung. Phần lớn xem AI là công cụ hỗ trợ nâng cao hiệu quả, đặc biệt trong tình huống khẩn cấp hoặc khi chịu áp lực thời gian. Tuy nhiên, cũng có một bộ phận sinh viên dùng AI thay cho tự học hoặc để hoàn thành toàn bộ bài tập, phản ánh nguy cơ phụ thuộc công nghệ và tiềm ẩn rủi ro vi phạm liêm chính học thuật.

4.2. Đánh giá độ tin cậy và giá trị thang đo

Kết quả kiểm định cho thấy tất cả thang đo đều đạt độ tin cậy rất cao (Cronbach's Alpha > 0,98; hệ số tương quan hiệu chỉnh > 0,9), đủ điều kiện cho các phân tích tiếp theo. Phân tích EFA cũng khẳng định tính hội tụ và phân biệt tốt (KMO = 0,915; Bartlett's Test $p < 0,001$), cấu trúc đơn chiều giải thích hơn 83% phương sai, phản ánh tính ổn định của thang đo.

4.3. Kết quả hồi quy và kiểm định mô hình

Quan hệ trực tiếp

Phân tích hồi quy tuyến tính cho thấy áp lực học tập (AP) có ảnh hưởng mạnh đến ý định sử dụng AI (IN) ($\beta \approx 0,89$; $R^2 = 0,842$), đồng thời IN dự báo đáng kể hành vi sử dụng AI (AU) ($\beta \approx 0,91$; $R^2 = 0,862$). Mô hình hồi quy đa biến khẳng định cả AP, IN và AU đều tác động thuận chiều đến vi phạm liêm chính học thuật (VI), trong đó IN là yếu tố chi phối mạnh nhất ($\beta = 0,47$), tiếp đến là AP ($\beta = 0,34$) và AU ($\beta = 0,17$). Kết quả này phù hợp với Thuyết hành vi có kế hoạch (TPB), nhấn mạnh vai trò trung gian của ý định trong hình thành hành vi.

Hiệu ứng trung gian

Kết quả phân tích PROCESS và Bootstrap (5.000 mẫu) chỉ ra rằng tổng ảnh hưởng gián tiếp là đáng kể (0,5755; $p < 0,001$). Đường $AP \rightarrow IN \rightarrow VI$ có tác động mạnh nhất (0,4171), trong khi $AP \rightarrow AU \rightarrow VI$ có ảnh hưởng trung bình (0,1106) và chuỗi $AP \rightarrow IN \rightarrow AU \rightarrow VI$ nhỏ hơn (0,0478) nhưng vẫn có ý nghĩa thống kê. Điều này khẳng định ý định sử dụng AI là biến trung gian then chốt, truyền tải tác động từ áp lực học tập đến hành vi vi phạm.

Hiệu ứng điều tiết

Nhận thức đạo đức học thuật (AIW) điều tiết đáng kể các mối quan hệ $IN \rightarrow VI$ và $AU \rightarrow VI$, song theo hướng khuếch đại thay vì làm suy giảm. Nghĩa là, ngay cả khi sinh viên có hiểu biết về chuẩn mực

đạo đức, việc tin tưởng vào AI và áp lực học tập vẫn có thể khiến hành vi vi phạm gia tăng.

Mô hình nghiên cứu giải thích 89,6% phương sai của biến VI, cho thấy độ phù hợp rất cao. Kết quả khẳng định rằng áp lực học tập là nguồn gốc thúc đẩy ý định sử dụng AI, trong khi ý định lại đóng vai trò quyết định, trực tiếp dẫn đến hành vi vi phạm liên chính học thuật. AU có tác động bổ sung nhưng ở mức thấp hơn. Những phát hiện này không chỉ củng cố giá trị ứng dụng của các nền tảng lý thuyết (TPB, Stress-Coping, Integrity Framework), mà còn cung cấp cơ sở thực nghiệm quan trọng cho việc xây dựng chính sách quản lý áp lực học tập, định hướng hành vi và giáo dục đạo đức học thuật trong kỷ nguyên số.

Bảng 2: Bảng tổng hợp kết quả nghiên cứu

Nội dung kiểm định	Chỉ số/Kết quả	Diễn giải
Độ tin cậy thang đo (EFA)	KMO = 0,915; Bartlett's Test $p < 0,001$; communalities 0,741-0,877; tổng phương sai giải thích = 83,3%	Thang đo đạt độ tin cậy và tính hội tụ cao, cấu trúc đơn nhân tố, phản ánh "thái độ và hành vi tích hợp AI trong học tập đại học".
$AP \rightarrow IN$	$\beta = 0,8879$; $R^2 = 0,842$	Áp lực học tập tác động rất mạnh đến ý định sử dụng AI, giải thích 84,2% phương sai của IN.
$IN \rightarrow AU$	$\beta = 0,909$ -0,937; $R^2 = 0,862$	Ý định là tiền đề trực tiếp và mạnh mẽ của hành vi sử dụng AI.
$AP \rightarrow VI$	$\beta = 0,3370$	AP có tác động trực tiếp và có ý nghĩa thống kê đến vi phạm liên chính học thuật.
$IN \rightarrow VI$	$\beta = 0,4697$; $R^2 = 0,896$	IN là yếu tố chi phối nhất, ảnh hưởng mạnh đến vi phạm.
$AU \rightarrow VI$	$\beta = 0,1720$	AU có tác động độc lập, trung bình nhưng vẫn có ý nghĩa.
Mô hình hồi quy đa biến	$R^2 = 0,896$; $F = 7.241,817$; $p < 0,001$	AP, IN, AU đồng thời giải thích 89,6% phương sai VI $\rightarrow$ mô hình có độ phù hợp rất cao.
Hiệu ứng trung gian (Bootstrap 5.000 mẫu)	Tổng ảnh hưởng gián tiếp = 0,5755 ($p < 0,001$); $AP \rightarrow IN \rightarrow VI = 0,4171$ (mạnh); $AP \rightarrow AU \rightarrow VI = 0,1106$ (trung bình); $AP \rightarrow IN \rightarrow AU \rightarrow VI = 0,0478$ (có ý nghĩa)	IN là biến trung gian then chốt, vừa kết nối AP với AU, vừa dẫn tới VI.

Nội dung kiểm định	Chỉ số/Kết quả	Diễn giải
Hiệu ứng điều tiết (PROCESS Model 1)	AIW điều tiết quan hệ $IN \rightarrow VI$ ($\beta = 0,12$) và $AU \rightarrow VI$ ($\beta = 0,09$), theo hướng khuếch đại	Nhận thức đạo đức không hạn chế mà làm gia tăng tác động của ý định và hành vi lên vi phạm.

4.4. Thảo luận kết quả

Kết quả cho thấy, ngay cả khi nhận thức được chuẩn mực liên chính học thuật, sinh viên vẫn có thể vi phạm khi chịu áp lực học tập cao hoặc tin tưởng mạnh vào hiệu quả của AI. Điều này khẳng định giá trị của Thuyết hành vi có kế hoạch - TPB (Ajzen, 1991), đồng thời mở rộng Lý thuyết căng thẳng - thích ứng (Lazarus & Folkman, 1984) sang bối cảnh học thuật số.

Đáng chú ý, nhận thức đạo đức học thuật (AIW) không hạn chế mà lại khuếch đại tác động của ý định và hành vi sử dụng AI đến vi phạm. Hiện tượng này có thể giải thích bởi cơ chế *moral self-licensing*, khi cá nhân tin rằng “chứng chỉ đạo đức” cho phép họ linh hoạt hơn trong hành vi lệch chuẩn (Merritt et al., 2010). Kết quả này cho thấy TPB cần được mở rộng bằng cách tích hợp các yếu tố tâm lý đạo đức để giải thích hành vi trong kỷ nguyên AI.

Phát hiện cũng phù hợp với nghiên cứu trước: Cotton, Cotton và Shipway (2023) cho rằng sinh viên vẫn gian lận bằng AI dù ý thức rủi ro; Borge (2024) chứng minh áp lực học tập làm suy yếu vai trò điều chỉnh của nhận thức đạo đức; còn Nguyen và Goto (2024) cho thấy tỷ lệ sinh viên thừa nhận gian lận bằng AI tăng gần gấp ba lần khi khảo sát bằng *list experiment* so với hỏi trực tiếp. Điều này phản ánh tính phức tạp của hành vi học

(Nguồn: Tác giả tổng hợp từ kết quả nghiên cứu)

thuật trong môi trường số, chịu tác động đồng thời từ nhận thức, áp lực và cơ chế giám sát.

Hàm ý lý thuyết

Nghiên cứu củng cố tính dự báo của TPB (Ajzen, 1991) nhưng cũng chỉ ra giới hạn khi đối diện với hành vi trong bối cảnh AI. Sự hiện diện của cơ chế *moral self-licensing* cho thấy nhận thức đạo đức có thể mang tính hai mặt: vừa định hướng hành vi, vừa tiềm ẩn nguy cơ khuếch đại vi phạm (Merritt et al., 2010). Do đó, TPB cần được mở rộng bằng cách tích hợp thêm các yếu tố tâm lý đạo đức và bối cảnh xã hội - văn hóa, đồng thời kết nối với lý thuyết căng thẳng - thích ứng để giải thích toàn diện hơn hành vi học thuật trong môi trường số.

Hàm ý thực tiễn

Kết quả mang lại cơ sở thực nghiệm cho quản lý giáo dục và giảng dạy đại học. Thứ nhất, nâng cao nhận thức đạo đức cần đi kèm cơ chế giám sát, chế tài và đánh giá rõ ràng. Thứ hai, các trường đại học cần ban hành hướng dẫn minh bạch về sử dụng AI, yêu cầu sinh viên khai báo trung thực. Thứ ba, giảng viên nên đổi mới phương pháp dạy và đánh giá, khuyến khích sản phẩm cá nhân hóa, hạn chế phụ thuộc vào AI. Cuối cùng, ở cấp quốc gia, cần phát triển khung quản lý AI trong giáo dục để vừa tận dụng lợi ích công nghệ, vừa duy trì chuẩn mực liên chính học thuật.

Hàm ý chính sách - quản lý giáo dục

Từ những phân tích trên, nghiên cứu đề xuất một số hàm ý quản lý:

1. Không chỉ tuyên truyền: Cần bổ sung cơ chế giám sát hành vi sử dụng AI thay vì chỉ dựa vào ý thức cá nhân.

2. Chính sách kiểm soát phân tầng: Xây dựng biện pháp quản lý dựa trên động cơ và tần suất sử dụng, tránh cấm đoán tuyệt đối gây phản tác dụng.

3. Minh bạch sử dụng AI: Yêu cầu sinh viên khai báo rõ ràng việc sử dụng AI trong quá trình học tập và nghiên cứu.

4. Ứng dụng công cụ phát hiện AI: Tích hợp công cụ dò tìm AI trong đánh giá học thuật, hỗ trợ giảng viên đảm bảo tính công bằng và minh bạch.

Như vậy, nghiên cứu không chỉ củng cố cơ sở lý thuyết hiện có mà còn đóng góp mở rộng khung phân tích, khi chỉ ra rằng nhận thức đạo đức có thể mang tính hai mặt: vừa là nền tảng định hướng, vừa có khả năng tạo ra hiệu ứng “giấy phép” cho hành vi lệch chuẩn.

V. Kết luận và kiến nghị

5.1. Kết luận

Nghiên cứu khảo sát hơn 2.500 sinh viên Trường Đại học Mở Hà Nội đã phân tích mối quan hệ giữa áp lực học tập, ý định sử dụng AI, hành vi sử dụng AI và vi phạm liên chính học thuật. Kết quả cho thấy:

- Áp lực học tập có tác động thuận chiều đáng kể đến ý định sử dụng AI, và ý định này tiếp tục dự báo mạnh mẽ hành vi sử dụng thực tế.

- Cả áp lực học tập, ý định và hành vi sử dụng AI đều góp phần gia tăng nguy

cơ vi phạm liên chính học thuật, trong đó ý định giữ vai trò chi phối.

- Nhận thức đạo đức học thuật (AIW) không làm giảm mà khuếch đại mối quan hệ giữa ý định, hành vi sử dụng AI và vi phạm, phản ánh cơ chế *moral self-licensing* (Merritt et al., 2010).

Về mặt lý thuyết, nghiên cứu bổ sung bằng chứng cho TPB (Ajzen, 1991), mở rộng khung căng thẳng - thích ứng (Lazarus & Folkman, 1984), đồng thời chỉ ra tác động hai mặt của nhận thức đạo đức trong hành vi học thuật. Về thực tiễn, kết quả nhấn mạnh sự cần thiết của chính sách quản lý và giáo dục liên chính học thuật phù hợp với bối cảnh AI.

5.2. Kiến nghị

Từ kết quả nghiên cứu, một số khuyến nghị được đề xuất:

1. Ở cấp cơ sở giáo dục đại học:

- Ban hành quy định rõ ràng và yêu cầu sinh viên khai báo minh bạch việc sử dụng AI.

- Tích hợp giáo dục đạo đức số và liên chính học thuật vào chương trình đào tạo chính khóa, kết hợp tập huấn và trải nghiệm.

- Áp dụng công cụ phát hiện AI trong đánh giá để hỗ trợ giảng viên bảo đảm tính công bằng.

2. Ở cấp giảng viên:

- Đổi mới phương pháp dạy học và đánh giá, khuyến khích tư duy phản biện và sản phẩm cá nhân hóa.

- Phát huy vai trò cố vấn học tập, hỗ trợ sinh viên cân bằng áp lực và định hướng sử dụng AI đúng mục đích.

3. Ở cấp sinh viên:

- Rèn luyện kỹ năng số và năng lực đạo đức trong sử dụng AI, nhận thức rõ rủi ro vi phạm và hệ quả lâu dài.

- Thực hành trung thực học thuật, coi khai báo minh bạch AI là chuẩn mực bắt buộc.

4. Ở cấp chính sách quốc gia:

- Bộ GD&ĐT xây dựng khung chính sách AI trong giáo dục đại học, quy định minh bạch học thuật, chia sẻ dữ liệu vi phạm và khuyến khích áp dụng công nghệ quản lý.

- Thúc đẩy nghiên cứu liên ngành về AI - giáo dục - đạo đức học thuật làm cơ sở cho chính sách dài hạn.

Tổng thể, nghiên cứu không chỉ cung cấp bằng chứng thực nghiệm cho Việt Nam mà còn đóng góp góc nhìn mới vào diễn đàn quốc tế, khẳng định rằng nhận thức đạo đức cao chưa đủ để ngăn ngừa vi phạm nếu thiếu giám sát, minh bạch và cơ chế kiểm soát hiệu quả.

Lời cảm ơn

Tác giả cảm ơn Trường Đại học Mở Hà Nội đã hỗ trợ điều kiện khảo sát và cung cấp dữ liệu; cảm ơn giảng viên, sinh viên đã tham gia và đóng góp phản hồi. Nghiên cứu cũng ghi nhận sự hỗ trợ kỹ thuật của ChatGPT (OpenAI, GPT-4.0) trong hệ thống hóa tài liệu, gợi ý mô hình, soạn thảo và định dạng văn bản.

Ghi chú sử dụng ChatGPT

Trong nghiên cứu này, ChatGPT được sử dụng như công cụ hỗ trợ học thuật nhằm tóm tắt tài liệu, gợi ý phân tích và chỉnh sửa ngôn ngữ. Tuy nhiên, toàn bộ nội dung, phân tích và kết luận cuối cùng đều

do tác giả trực tiếp kiểm tra, hiệu chỉnh và chịu trách nhiệm. Việc sử dụng ChatGPT không thay thế tư duy khoa học hay trách nhiệm học thuật của tác giả; công cụ này chỉ đóng vai trò hỗ trợ về ngôn ngữ và tổ chức nội dung, không tạo dữ liệu, đưa ra kết luận hay đảm bảo độ chính xác tuyệt đối của kết quả nghiên cứu.

Tài liệu tham khảo:

- [1]. Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179-211. [https://doi.org/10.1016/0749-5978\(91\)90020-T](https://doi.org/10.1016/0749-5978(91)90020-T)
- [2]. Lazarus, R., & Folkman, S. (1984). *Stress, Appraisal, and Coping*. New York: Springer.
- [3]. <https://doi.org/10.54855/ijte.24414>
- [4]. Merritt, A. C., Effron, D. A., & Monin, B. (2010). *Moral selflicensing: When being good frees us to be bad*. *Social and Personality Psychology Compass*, 4(5), 344-357. <https://doi.org/10.1111/j.1751-9004.2010.00263.X>
- [5]. Jain, K., & Raghuram, R. (2024). Gen-AI integration in higher education: Predicting intentions using SEM-ANN approach. *International Journal of Information and Communication Technology Education*, 20(3), 1-18. <https://doi.org/10.1007/s10639-024-12506-4>
- [6]. UNESCO. (2021). *Recommendation on the ethics of artificial intelligence*. Paris: United Nations Educational, Scientific and Cultural Organization. Retrieved from <https://unesdoc.unesco.org/ark:/48223/pf0000379920>
- [7]. Đặng, T. T. (2024). *Khởi nghiệp trong đào tạo thời công nghệ trí tuệ nhân tạo (AI): Khó khăn và giải pháp*. Trong Kỷ yếu Hội thảo quốc gia: Nâng cao hiệu quả hoạt động khởi nghiệp của sinh viên (tr. 611-629).
- [8]. Sagge, R. G. (2024). Exploring Students' Perceptions of Academic Integrity in the Digital Classroom

- through Exploratory Factor Analysis (EFA). *Indian Journal of Science and Technology*, 17(46), 4907-4920. <https://doi.org/10.17485/ijst/v17i46.2384>
- [9]. Bennett, L., & Abusalem, A. (2024). Building Academic Integrity and Capacity in Digital Assessment in Higher Education. *Athens Journal of Education*. <https://doi.org/10.30958/aje.11-1-5>
- [10]. Kim, B. J. (2024). The mental health implications of artificial intelligence: Job stress, burnout, and self-efficacy in AI learning. *Humanities and Social Sciences Communications*, X, Article
- [11]. NorabuenaFigueroa, R. P., et al. (2025). Digital teaching practices and student academic stress in the era of digitalization in higher education. *Applied Sciences*, 15(3), 1487. <https://doi.org/10.3390/app15031487>
- [12]. Zhang, S. (2024). The roles of academic self-efficacy, academic stress, performance expectations, and AI dependency. *Educational Technology Research and Development*, 72, Article
- [13]. Merritt, A. C., Effron, D. A., & Monin, B. (2010). Moral Self-Licensing: When Being Good Frees Us to Be Bad. *Social and Personality Psychology Compass*, 4(5), 344-357. <https://doi.org/10.1111/J.1751-9004.2010.00263.X>
- [14]. Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2023). 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>
- [15]. Borge, Samuel, "Academic Cheating and Stressors at the University Level" (2024). *Honors Theses*. 143. <https://digitalcommons.assumption.edu/honorstheses/143>
- [16]. Nguyen, H.M., Goto, D. Unmasking academic cheating behavior in the artificial intelligence era: Evidence from Vietnamese undergraduates. *Educ Inf Technol* 29, 15999-16025 (2024). <https://doi.org/10.1007/s10639-024-12495-4>
- [17]. Nguyen, V. P., Nguyen, Q. T., Pham, V. V., Vo, T. H., Linna, T., Le, K. T., & Trần, T. T. H. (2024). Factors associated with depression, anxiety, and stress among medical college students. *Tạp Chí Nghiên Cứu Y Học*, 184(11E15), 96-105. <https://doi.org/10.52852/tencyh.v184i11e15.2704>
- [18]. Fatima, F., Naz, S., Shabnam, N., Ali, S., & Fatima, S. (2021). *Investigation of Depression and Anxiety in Students due to Burdenized Curriculum and Shortage of Time at University level*. 2(1), 21-35. <https://doi.org/10.52053/JPAP.V2I1.34>
- [19]. Zhytnyk, V., & Gryniuk, O. (2024). Artificial intelligence and education. *Constructive Geography and Rational Use of Natural Resources*, 5 (1), 74-82. <https://doi.org/10.17721/2786-4561.2024.5.1.-9/12>
- [20]. Steare, T., Gutiérrez Muñoz, C., Sullivan, A., & Lewis, G. (2023). The association between academic pressure and adolescent mental health problems: A systematic review. *Journal of Affective Disorders*, 339, 302-317. <https://doi.org/10.1016/j.jad.2023.07.028>
- [21]. Nguyen, N.A. (2025). *Bộ công cụ khảo sát về hành vi sử dụng AI trong học tập và liên chính học thuật* [Tài liệu chưa công bố]. Trường Đại học Mở Hà Nội.
- [22]. Pitts, M., Marcus, K., & Motamedi, V. (2025). *AI chatbots and the future of higher education: Student perspectives on benefits and risks*. arXiv. <https://arxiv.org/abs/2505.02198>
- [23]. Ngo, C.-L., Tran, T. N., & Nguyen, T. T. (2024). Academic integrity in the age of generative AI: Perceptions and responses of Vietnamese EFL teachers. *Teaching English with Technology*, 24(1), 1-15. <https://doi.org/10.56297/FSYB3031/MXNB7567>
- [24]. Northern Michigan University. (n.d.). *Academic dishonesty using generative AI*. Center for Teaching and Learning.

- Retrieved August 28, 2025, from <https://nmu.edu/ctl/academic-dishonesty-using-generative-ai>
- [25]. Chan, C. K. Y. (2025). Students' perceptions of 'AI-giarism': Investigating changes in understandings of academic misconduct. *Education and Information Technologies*, 30(4), 8087-8108. <https://doi.org/10.1007/s10639-024-13151-7>
- [26]. University of Oxford. (2025). *Plagiarism*. Academic guidance for students. <https://www.ox.ac.uk/students/academic/guidance/skills/plagiarism>
- [27]. Overono, O. (2025). Integrating AI disclosure statements into psychology education: Ethical communication and teaching practices. *Teaching of Psychology*. Advance online publication. <https://doi.org/10.1177/00986283241275664>
- [28]. Weber-Wulff, D., Albrecht, S., Gipp, B., & Meuschke, N. (2023). Testing of detection tools for AI-generated text. *International Journal for Educational Integrity*, 19(9), 1-25. <https://doi.org/10.1007/s40979-023-00146-z>
- [29]. Chan, C. K. Y. (2023). An ecological framework for generative AI in education: Pedagogical, institutional, and governance perspectives. *International Journal of Educational Technology in Higher Education*, 20(1), 52. <https://doi.org/10.1186/s41239-023-00408-3>

THE IMPACT OF ACADEMIC PRESSURE ON AI USE BEHAVIOR AND THE RISK OF ACADEMIC INTEGRITY VIOLATIONS: AN EMPIRICAL ANALYSIS FROM A STUDENT SURVEY AT HANOI OPEN UNIVERSITY

Nguyen Anh Hoan²

Abstract: *In the context of artificial intelligence (AI) being increasingly integrated into teaching and learning in higher education, understanding the factors influencing students' AI usage behavior and the associated risk of academic integrity violations has become essential. This study analyzes the relationships among academic pressure, intention to use AI, actual AI usage behavior, and academic integrity violations among university students. Based on a survey of more than 2,500 students from various majors and academic years at Hanoi Open University, the findings indicate that academic pressure significantly influences the intention to use AI; this intention strongly predicts actual usage behavior; and both factors contribute to an increased risk of violating academic integrity standards. The study proposes recommendations at the institutional, faculty, and student levels to enhance awareness, skills, and ethical practices in AI usage for learning, thereby fostering a transparent, responsible academic environment that adapts effectively to digital transformation in higher education.*

Keywords: *academic pressure, ai use behavior, artificial intelligence, ethical awareness, academic integrity violation*

² Hanoi Open University