

الاستخدام الأخلاقي لنماذج الذكاء الاصطناعي التوليدي: دراسة العوامل المؤثرة لدى طلبة الدراسات العليا في قسم علم المعلومات بجامعة الملك سعود

The Ethical Use of Generative Artificial Intelligence
Models: A Study of the Influencing Factors Among
Graduate Students in the Department of Information
Science at King Saud University

د. أحمد عبد الله بن خضير

أستاذ مشارك بجامعة الملك سعود بالرياض

khudair@KSU.EDU.SA

ساره حمد القحطاني

باحثة دكتوراه بجامعة الملك سعود بالرياض

444203667@student.KSU.EDU.SA

مسؤولية المراسلات:

الباحث: ساره حمد القحطاني.

قسم علم المعلومات ، جامعة الملك سعود ، ١١٤٢١ ، الرياض.

البريد الإلكتروني: 444203667@student.KSU.EDU.SA

١ سبتمبر ٢٠٢٥	تاريخ الاستلام
٢٨ سبتمبر ٢٠٢٥	تاريخ القبول
٨ أكتوبر ٢٠٢٥	تاريخ النشر

Abstract:

This study investigates the ethical use of generative artificial intelligence (GenAI) models among doctoral students in the Department of Information Science at King Saud University. Drawing on the extended Unified Theory of Acceptance and Use of Technology (UTAUT2), the research incorporates additional ethical constructs—moral responsibility, perceived privacy, data security, and ethical awareness to provide a comprehensive understanding of behavioral intention in adopting GenAI in higher education. A descriptive methodology was applied, and data were collected through a structured questionnaire administered to the entire cohort of PhD students (N = 15). Validity and reliability of the instrument were confirmed through expert review and Cronbach's alpha testing. Findings revealed that technical factors particularly performance expectancy—positively influenced students' behavioral intention, while social influence, facilitating conditions, and habit showed no significant impact. On the ethical side, moral responsibility contributed to stronger behavioral intention, whereas perceived privacy had limited influence. Notably, data security and ethical awareness emerged as the most critical determinants, underscoring the role of data protection and academic integrity in shaping responsible engagement with GenAI models.

The study highlights the importance of developing institutional policies to regulate the academic use of GenAI and raising awareness of ethical risks and responsibilities among students and researchers. It also calls for investment in secure infrastructures and recommends that future research extend to larger samples across universities and disciplines, thereby enhancing sustainable and ethical adoption of generative AI in academic environments.

Keywords: Generative AI, ethical use, graduate students, Extension of the Unified Theory of Acceptance and Use of Technology (UTAUT2).

مستخلص:

تهدف الدراسة إلى استكشاف العوامل التقنية والأخلاقية المؤثرة على الاستخدام الأخلاقي لنماذج الذكاء الاصطناعي التوليدي من قبل طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود، وذلك بالاعتماد على امتداد النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT2، مع إضافة بعض العوامل الأخلاقية مثل: المسؤولية الأخلاقية، والخصوصية المدركة، وأمن البيانات، والوعي الأخلاقي. واعتمدت الدراسة على المنهج الوصفي. وتمثلت أداة جمع البيانات في الاستبانة. وتكونت عينة الدراسة من طلبة الدراسات العليا (الدكتوراه) في قسم علم المعلومات بجامعة الملك سعود والبالغ عددهم ١٥ طالب وطالبة. وتوصلت الدراسة إلى مجموعة من النتائج أهمها: أن العوامل التقنية وخاصة الأداء المتوقع تؤثر بدرجة إيجابية على النية السلوكية لطلبة الدراسات العليا في استخدامهم الأخلاقي لنماذج الذكاء الاصطناعي التوليدي، في حين لم يظهر للتأثير الاجتماعي والظروف الميسرة والعادة أثر معنوي واضح. أما على المستوى الأخلاقي، فقد تبين أن المسؤولية الأخلاقية تعزز النية السلوكية، إلا أن الخصوصية المدركة لم تحقق تأثيراً يذكر. وفي المقابل، برز كل من أمن البيانات والوعي الأخلاقي كأكثر العوامل تأثيراً على النية السلوكية، وهذا يعكس أهمية قضايا حماية البيانات، والنزاهة الأكاديمية في تشكيل توجهات الطلبة نحو الاستخدام الأخلاقي والمسؤول لهذه النماذج الذكية. وأوصت الدراسة بضرورة تطوير سياسات مؤسسية واضحة لضبط الاستخدام الأكاديمي للذكاء الاصطناعي التوليدي، إلى جانب تعزيز وعي الطلبة والباحثين بالجوانب الأخلاقية والمخاطر المحتملة. كما تدعو إلى الاستثمار في تقنيات حماية البيانات، وأجراء الدراسات المستقبلية لتشمل عينات أوسع من جامعات وتخصصات مختلفة، بما يساهم في بناء بيئة أكاديمية أكثر وعياً واستدامة في التعامل مع هذه النماذج الذكية.

الكلمات المفتاحية: الذكاء الاصطناعي التوليدي، الاستخدام الأخلاقي، طلبة الدراسات العليا، امتداد النظرية الموحدة لقبول واستخدام التكنولوجيا (UTAUT2).

المقدمة

شهد العالم في السنوات الأخيرة تطوراً متسارعاً في مجال التقنيات الحديثة والابتكار، وبرزت التقنيات الناشئة Emerging Technologies وهي التقنيات المتقدمة التي تمر في مرحلة التطوير حالياً أو سيتم تطويرها خلال السنوات القادمة، ولها تأثير ملحوظ في العديد من المجالات (Ganesamoorthy & Selvakamal, 2024). (Ganesamoorthy & Selvakamal, 2024). ويمثل الذكاء الاصطناعي (Artificial Intelligence (AI أحد أبرز هذه التقنيات والذي يملك القدرات التحليلية الهائلة، واتخاذ القرارات بشكل مستقل يحاكي الذكاء البشري (The European Commission's, 2018) وتوجهت الجامعات إلى استثمار إمكانياته في المجالات التعليمية والإدارية والبحثية، فمثلاً أنشأت جامعة الملك سعود مركزاً للدراسات المتقدمة في الذكاء الاصطناعي (ذكاء) لدعم البحث والتطوير وابتكار الحلول للذكاء الاصطناعي تحقيقاً لمستهدفات رؤية المملكة ٢٠٣٠ (جامعة الملك سعود، ٢٠٢٣) واستخدمت جامعة ستانفورد محرك البحث الدلالي Yewno في مكتباتها لتحسين عملية البحث عن المصطلحات واستخراج العلاقات بين المفاهيم، وعرضها بطرق تفاعلية وبصرية (Schreur, 2019).

ومع تطور هذه التقنيات، برزت نماذج الذكاء الاصطناعي التوليدي Generative AI (Gen AI) (GAI) كأحد نماذج التعلم الآلي Machine Learning (ML) القادرة على إنتاج محتوى جديد وأصيل استناداً إلى البيانات الضخمة (Ahmed et al., 2024)، واستثمرت العديد من الجامعات إمكانيات هذه النماذج الذكية في عمليات البحث، والتعليم، وصياغة الأفكار، والتلخيص، والترجمة وتصميم العروض (Ahmed et al., 2024; Chan & Hu, 2023; Yusuf et al., 2024). وبالرغم من التوسع في استخدام هذه النماذج عالمياً وتشجيع بعض الجامعات لاستخدامها، إلا أن بعض الجامعات حظرت هذا الاستخدام، ووضع البعض منها سياسات لتنظيم هذا الاستخدام (Mcdonald et al., 2025).

ويواجه المجتمع الأكاديمي عند تبني نماذج الذكاء الاصطناعي التوليدي في التعليم العالي تحديات أخلاقية تتمثل في ضرورة الاستخدام الآمن والمسؤول بما يضمن النزاهة الأكاديمية، وتجنب التحيزات الكامنة في بيانات التدريب التي قد تعيد إنتاج التحيزات بين أفراد

المجتمع (OECD, 2023)، بالإضافة إلى أن الإفراط في الاعتماد على هذه النماذج يضعف القدرات البحثية والمهارات التحليلية لدى الأفراد (Ahmed et al., 2024; Bouteraa et al., 2024; Chan & Hu, 2023a; Mironova et al., 2024; Zlateva et al., 2024). كما أن نماذج الذكاء الاصطناعي التوليدي تميل إلى نسخ أجزاء من المقالات من مصادر الإنترنت دون إشارة إلى العمل الأصلي أو تقديم الاستشهادات (Michel-Villarreal et al., 2023). وتمتلك هذه النماذج الذكية القدرة على إنشاء محتوى يمكن أن يعتبر نسخة من أعمال محمية بحقوق الملكية الفكرية دون الإشارة إلى الأعمال الأصلية، مما يثير العديد من التساؤلات حول أصالة الأعمال الأكاديمية وحقوق المؤلفين (Helberger & Diakopoulos, 2023).

وأظهرت مراجعة الدراسات السابقة عن غياب الدراسات العربية التي تناولت الاستخدام الأخلاقي لنماذج الذكاء الاصطناعي التوليدي GAI في التعليم العالي، وركزت معظم الدراسات الأجنبية على استخدام عوامل النظرية الموحدة لقبول واستخدام التكنولوجيا بإصدارها UTAUT1، UTAUT2 لقياس قبول واستخدام Gen AI في التعليم العالي دون استخدام UTAUT كمقياس للاستخدام الأخلاقي. مما يبرز الحاجة لدراسة شاملة تجمع ما بين الجانبين التقني والأخلاقي.

مشكلة الدراسة

أحدثت نماذج الذكاء الاصطناعي التوليدي Gen AI تحولاً في التعليم العالي من خلال الفرص التي قدمتها مثل تخصيص التعلم، ودعم البحث العلمي، وتحسين جودة المخرجات الأكاديمية (Ahmed et al., 2024; Chan & Hu, 2023). إلا أن هذا التحول أثار العديد من التحديات الأخلاقية، والتي تتمثل في انتهاك الخصوصية، وتهديد أمن البيانات، ومخالفة النزاهة الأكاديمية، والتحيزات الخوارزمية (Huallpa et al., 2023; Saxena et al., 2023)، مما دفع بعض الجامعات إلى تقييد أو حظر استخدام هذه النماذج الذكية (Herman, 2023; Yeralan & Lee, 2023).

وبالرغم من توظيف العديد من الدراسات للنظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT1، UTAUT2 في دراسة قبول واستخدام هذه النماذج الذكية، إلا أن معظم الدراسات ركزت على العوامل التقنية دون قياس للاستخدام الأخلاقي. أما الدراسات التي تناولت القضايا

الأخلاقية فلم تعتمد على إطار نظري متكامل (Gallent-Torres et al., 2023). وبناءً على ذلك، تتمثل مشكلة هذه الدراسة في سد هذه الفجوة من خلال تطبيق UTAUT2 مع دمج عوامل أخلاقية إضافية مثل المسؤولية الأخلاقية، والخصوصية المدركة، وأمن البيانات، والوعي الأخلاقي كمقياس للاستخدام الأخلاقي لنماذج الذكاء الاصطناعي التوليدي Gen AI في التعليم العالي.

أهمية الدراسة

تستمد الدراسة أهميتها من أهمية التحولات التي أحدثتها نماذج الذكاء الاصطناعي التوليدي GAI في التعليم العالي، حيث أصبحت أداة فعالة لدعم التعلم المخصص، وتحسين العمليات الأكاديمية، وتعزيز الابتكار في البحث والتعليم (Ahmed et al., 2024; Aljuaid, 2024; Chan & Hu, 2023; Gallent-Torres et al., 2023). ضرورة توظيف هذه النماذج بما يتماشى مع القيم والمعايير الأخلاقية مثل الشفافية، والعدالة، وأمن البيانات، والمحافظة على الملكية الفكرية، حيث برزت العديد من التحديات الأخلاقية مثل الاعتماد المفرط على هذه النماذج الذكية، ومخاطر الخصوصية، والنزاهة الأكاديمية (Hualpa et al., 2023; Liu et al., 2023; Luk et al., n.d.; UNESCO, 2023). كما تكتسب الدراسة أهميتها أيضاً من توظيفها لامتداد النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT2 مع دمج عوامل أخلاقية مثل المسؤولية الأخلاقية، والخصوصية المدركة، وأمن البيانات، والوعي الأخلاقي كمقياس للاستخدام الأخلاقي والمسؤول لنماذج الذكاء الاصطناعي التوليدي في البيئة الأكاديمية العربية. وتسعى الدراسة إلى تقديم توصيات عملية يمكن أن تسهم في صياغة سياسات تعليمية لدعم الابتكار، مع المحافظة على القيم الأخلاقية (Gallent-Torres et al., 2023; UNESCO, 2023).

أهداف الدراسة

تهدف هذه الدراسة إلى التعرف على العوامل المؤثرة في الاستخدام الأخلاقي لنماذج الذكاء الاصطناعي التوليدي لدى طلبة الدراسات العليا (مرحلة الدكتوراه) في قسم علم المعلومات بجامعة الملك سعود، وذلك من خلال:

١. فحص تأثير العوامل التقنية (الأداء المتوقع، التأثير الاجتماعي، الظروف الميسرة، العادة).

٢. فحص تأثير العوامل الأخلاقية (المسؤولية الأخلاقية، الخصوصية المدركة، أمن البيانات، الوعي الأخلاقي).

مصطلحات الدراسة

Generative AI Models

نماذج الذكاء الاصطناعي التوليدي

تُعرف نماذج الذكاء الاصطناعي التوليدي بأنها أنظمة الحوار بين الإنسان والحاسوب عبر الإنترنت باللغة الطبيعية. وهي أنظمة محادثة ذكية يمكنها التواصل مع البشر باستخدام لغتهم الطبيعية، وفي الوقت الفعلي حيث تقوم بمعالجة مدخلات المستخدم، ومن ثم تقديم المخرجات ذات الصلة بالمدخلات في شكل نص باللغة الطبيعية (Hasal et al., 2021). وتُعرف الدراسة نماذج الذكاء الاصطناعي التوليدي إجرائياً بأنها تطبيقات ذكية تولد إجابات فورية للأسئلة المدخلة من قبل طلبة الدكتوراة في جامعة الملك سعود، ويمكن استخدامها في المهام الأكاديمية والبحثية مثل ChatGPT, Gemini, Claude, LLaMA.

Theory Acceptance and Use of

النظرية الموحدة لقبول واستخدام التكنولوجيا of Unified Technology (UTAUT)

نظرية تهدف إلى تفسير سلوك استخدام الأفراد للتكنولوجيا الجديدة وتبنيم لها. تم تطويرها عام ٢٠٠٣ م من قبل (Venkatesh et al., 2003). وتجمع هذه النظرية عناصر من نظريات قبول التكنولوجيا، وتحتوي على أربع محددات رئيسية للنية السلوكية، والاستخدام الفعلي للتكنولوجيا مثل: الأداء المتوقع، الجهد المتوقع، التأثير الاجتماعي، والظروف الميسرة. وتم توسيع هذه النظرية في عام ٢٠١٢ م من قبل (Venkatesh et al., 2012) بإضافة عناصر أخرى، وهي: دافع المتعة، والعادة، والظروف الميسرة، وقيمة السعر، لفهم استخدام المستهلك Consumer Use للتكنولوجيا الجديدة (Venkatesh et al., 2012).

حدود الدراسة

طبقت الدراسة على طلبة الدراسات العليا (مرحلة الدكتوراة) في قسم علم المعلومات بجامعة الملك سعود، في شهر أغسطس من العام ٢٠٢٥ م.

الدراسات السابقة-

تم تصنيف الدراسات السابقة المرتبطة بموضوع الدراسة على قسمين كالتالي:

أولاً- الدراسات المتعلقة باستخدام وقبول نماذج الذكاء الاصطناعي التوليدي Generative AI في التعليم العالي بالاعتماد على النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT شهدت السنوات الأخيرة تزايداً ملحوظاً في الدراسات التي تناولت قبول واستخدام نماذج الذكاء الاصطناعي التوليدي في التعليم العالي، وذلك بالاعتماد على النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT1 وامتدادها UTAUT2. وقد ركزت هذه الدراسات بصورة أساسية على العوامل التقنية التقليدية مثل الأداء المتوقع، والجهد المتوقع، والتأثير الاجتماعي، والظروف الميسرة، وهناك محاولات محدودة لبعض الدراسات لإضافة عوامل جديدة ضمن إطار UTAUT؛ وفيما يلي أبرز ما توصلت إليه الدراسات في هذا الموضوع:

أجمعت العديد من الدراسات السابقة على أن عامل الأداء المتوقع Performance Expectancy (PE) يمثل العامل الأكثر مركزية في قبول واستخدام نماذج الذكاء الاصطناعي التوليدي في التعليم العالي. فإدراك الطلبة والأساتذة للفوائد التعليمية التي توفرها هذه النماذج من تحسين الكفاءة إلى رفع جودة المخرجات الأكاديمية يرتبط بشكل مباشر بنيتهم السلوكية لتبنيها واعتمادها. وقد دعمت هذا الاتجاه دراسات أجريت في العديد من الدول مثل المملكة العربية السعودية (Alzahrani, 2025; Elshaer et al., 2024; Sobaih et al., 2024)، والصين (Tian et al., 2024)، والشرق الأوسط (Khlaif et al., 2024)، وتركيا (Yilmaz et al., 2023)، وباكستان (Parveen et al., 2024)، وكرواتيا (Biloš & Budimir, 2024)، وكندا (Budhathoki et al., 2024) إضافةً إلى دراسة مقارنة في نيبال والمملكة المتحدة (Bazelaïs et al., 2024).

ومع هذا الإجماع الواسع، إلا أن بعض الدراسات أظهرت استثناءات مثيرة للاهتمام. ففي دراسة (Bahadur et al., 2024) في نيبال لم تجد أثراً يُذكر للأداء المتوقع على النية السلوكية، بل ركزت على دور العادة والتأثير الاجتماعي. وأظهرت دراسة (Bazelaïs et al., 2024) في كندا أن الأداء المتوقع كان العامل المؤثر الوحيد، في حين لم يثبت للجهد المتوقع أو التأثير الاجتماعي أو الظروف الميسرة أي دور. هذا التباين يعكس أن مركزية الأداء المتوقع قد

تتأثر باختلاف السياقات التعليمية والثقافية، إضافةً إلى تباين مستوى المهارات الرقمية والدعم المؤسسي.

وبناءً على ذلك، يمكن القول إن الأداء المتوقع يمثل "المحرك الرئيسي" لقبول واستخدام نماذج الذكاء الاصطناعي التوليدي في معظم البيئات الأكاديمية، إلا أن قوة تأثيره ليست مطلقة، بل مرتبطة بعوامل أخرى مثل البيئة المؤسسية، ومستوى الجاهزية الرقمية، والعادات التعليمية السائدة.

ويمثل عامل الجهد المتوقع (Effort Expectancy (EE أحد العوامل الجوهرية في نظرية UTAUT، حيث يفترض هذا العامل أن سهولة استخدام التقنية تزيد من تبنيها وبالتالي تؤثر على النية السلوكية لاعتمادها. غير أن الدراسات المتعلقة بالـ GAI كشفت عن بعض التباين. حيث برز الجهد المتوقع كعامل مؤثر إيجابياً على النية السلوكية، إذ أظهرت دراسات في المملكة العربية السعودية (Sobaih et al., 2024)، والصين (Tian et al., 2024)، والشرق الأوسط (Khlaif et al., 2024)، وتركيا (Yilmaz et al., 2023)، ونيبال/المملكة المتحدة (Budhathoki et al., 2024) أن سهولة التعامل مع ChatGPT تزيد من احتمالية اعتماده في العمليات الأكاديمية، سواءً في مجال البحث أو التعلم أو إعداد المهام. وتؤكد هذه النتائج أن جانب "سهولة الاستخدام" ما يزال قائم كحافز أساسي خاصةً في البيئات التي يفتقر فيها الطلبة للخبرة التقنية الكافية.

في المقابل، أظهرت دراسات أخرى أن أثر الجهد المتوقع ضعيف أو حتى غائب في بعض السياقات. فدراسة (Bouteraa et al., 2024) في دول جنوب شرق آسيا لم تجد تأثيراً واضحاً للجهد المتوقع على الاستخدام الفعلي، كما لم تظهر دراسة (Elshaer et al., 2024) في الجامعات السعودية، ودراسة (Biloš & Budimir, 2024) في كرواتيا أي دور يُذكر لهذا العامل. ويعكس هذا التباين اختلاف طبيعة العينات والبيئات، حيث قد تصبح سهولة الاستخدام أمراً بديهياً خصوصاً في البيئات ذات البنية الرقمية المتقدمة، أو بين المستخدمين ذوي الكفاءة التقنية العالية.

وبالتالي، يمكن القول إن أثر الجهد المتوقع نسبي: فهو مهم في البيئات التي في طور التبنّي الأولي لنماذج الذكاء الاصطناعي التوليدي، أو بين الطلبة ذوي المهارات التقنية

المحدودة، بينما يتراجع أثره تدريجياً في البيئات الرقمية المتطورة. وهذا يفتح نقاشاً مهماً حول ثبات عوامل القبول التقنية وتغيرها مع الزمن، ومع تطور الممارسات التعليمية. ويُعد التأثير الاجتماعي (SI) أحد العوامل التي تقيس مدى استجابة الأفراد لتشجيع أو ضغط محيطهم الأكاديمي، سواء من الأقران أو الزملاء أو القيادات. وقد أظهرت الدراسات نتائج متباينة بشأن دوره.

فمن جانب، أوضحت عدد من الدراسات أن التأثير الاجتماعي يعزز من النية السلوكية لاستخدام ChatGPT خاصةً بين الطلبة الذين يتأثرون بتجارب زملائهم، أو بحوافز جماعية غير مباشرة (Bahadur et al., 2024; Bouteraa et al., 2024; Budhathoki et al., 2024; Elshaer et al., 2024; Khlaif et al., 2024; Parveen et al., 2024; Sobaih et al., 2024) وفي هذه الحالات، يبدو أن الممارسات الأكاديمية السائدة، والبيئة الاجتماعية المحيطة تشكل عاملاً محفزاً لتبني هذه النماذج الذكية.

في المقابل، هذا العامل قد يكون محدود الأثر في بعض السياقات. فقد أظهرت دراسات في المملكة العربية السعودية (Alzahrani, 2025)، وتركيا (Yilmaz et al., 2023)، وكندا (Bazelais et al., 2024) أن التأثير الاجتماعي لم يكن له دورٌ يُذكر. ويُفسر ذلك بأن الاعتماد على النماذج الذكية قد يرتبط أكثر بالتجربة الفردية، والمهارة الشخصية بعيداً عن ضغط الأقران أو التوجهات المؤسسية. وعليه، يمكن النظر إلى التأثير الاجتماعي كعامل انتقالي؛ فدوره يبدو واضحاً في المراحل الأولى للتبني، لكنه يفقد أهميته بشكل تدريجي مع الممارسة المستمرة، والاعتماد على استخدام النماذج الذكية.

وتشير الظروف الميسرة (Facilitating Conditions (FC) إلى مدى توفر البنية التحتية التقنية والدعم المؤسسي، وقد ارتبطت في دراسات سابقة بتبني الأنظمة الرقمية. غير أن غالبية الدراسات حول الذكاء الاصطناعي التوليدي لم تجد له أثراً يُذكر على النية السلوكية أو الاستخدام الفعلي، كما هو الحال في دراسات أجريت في المملكة العربية السعودية (Elshaer et al., 2024; Sobaih et al., 2024)، وتركيا (Yilmaz et al., 2023)، وكندا (Bazelais et al., 2024) وكرواتيا (Biloš & Budimir, 2024) ويمكن تفسير ذلك بأن الوصول إلى النماذج

الذكية أصبح متاحاً عبر الإنترنت، مما قلل من أهمية وجود دعم تقني مباشر أو بنية تحتية خاصة.

لكن بعض الدراسات أظهرت نتائج مغايرة؛ إذ بينت دراسة (Alzahrani, 2025) في المملكة السعودية أن الظروف الميسرة كان لها تأثير إيجابي على النية السلوكية بشكل مباشر، مما يشير إلى أن الدعم المؤسسي والتنظيمي قد يعزز نية التبني في بعض السياقات. بينما في دراسات الشرق الأوسط (Khlaif et al., 2024)، وباكستان (Parveen et al., 2024) لم يظهر تأثير للظروف الميسرة على النية السلوكية، لكنه ارتبط بشكل مباشر بالاستخدام الفعلي، أي أن توفر السياسات الواضحة، أو الدعم التقني الملموس في الجامعات يساعد في تحويل النية إلى ممارسة عملية.

وبناءً على ذلك، يبدو أن دور الظروف الميسرة لم يعد محصوراً بالبنية التقنية وحدها، بل أصبح أكثر ارتباطاً بالسياسات المؤسسية، والإرشادات التنظيمية التي تضبط استخدام الذكاء الاصطناعي التوليدي في الجامعات، سواء عبر تعزيز النية أو تمكين الاستخدام الفعلي.

وعليه، فإن ما تكشفه الدراسات حول عامل الظروف الميسرة يوضح أنه عامل متحول، ففي بعض السياقات يظل غائب الأثر بفعل انتشار الوصول الرقمي، بينما في سياقات أخرى يصبح أداة تمكين مؤسسي، تضمن الانتقال من مجرد النية إلى الاستخدام الفعلي. وهذا يضع الجامعات أمام تحدٍ أساسي هل تكتفي بتوفير البنية التقنية؟ أم تسعى لصياغة سياسات وإرشادات تجعل الظروف الميسرة جزءاً من استراتيجية تبني الذكاء الاصطناعي التوليدي؟

وإلى جانب العوامل التقنية التقليدية، ظهرت ثلاث إضافات لافتة في بعض الدراسات. فقد تناولت دراسة (Budhathoki et al., 2024) عامل القلق Anxiety، باعتباره شعور نفسي قد يعيق الطلبة من التفاعل مع النماذج الذكية لتجنب الفشل أو سوء الاستخدام، مما يضعف من نيتهم السلوكية للتبني. في المقابل، أبرزت دراسة (Bahadur et al., 2024) قيمة التعلم Learning Value كعامل يعكس إدراك الطلبة للفائدة التعليمية المباشرة، مثل تحسين الفهم وتيسير إنجاز المهام، وهو ما يشكل دافع قوي لاستخدام النماذج الذكية. أما دراسة (Alzahrani, 2025) فقدمت إضافة مختلفة حيث أدخلت مخاطر الخصوصية Privacy Risk

كعامل أخلاقي، مشيرةً إلى أن القلق من تسريب البيانات أو انتهاكها قد يكون عامل حاسم في قرار المجتمع الأكاديمي بتبني هذه النماذج أو رفضها. ورغم أن هذه الإضافات الثلاث كشفت عن أبعاد نفسية وتربوية وأخلاقية مهمة، فإنها ظلت متفرقة وغير مدمجة في إطار شامل، مما يعزز الحاجة إلى نماذج أكثر تكاملاً قادرة على تفسير سلوك التبني في التعليم العالي.

يتضح من الدراسات السابقة أن العوامل التقنية في النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT ليست ثابتة، بل تتأثر بالسياق والبيئة المؤسسية. فالأداء المتوقع يظل المحرك الرئيس للتبني، بينما يتراجع أثر الجهد المتوقع مع ارتفاع الكفاءة الرقمية، ويُعد التأثير الاجتماعي عاملاً مرحلياً أكثر منه دائماً، أما الظروف الميسرة فتفقد قوتها في البيئات الرقمية المفتوحة لكنها تستعيد أهميتها حين تقترن بسياسات ودعم مؤسسي. هذا التباين يكشف أن اعتماد النماذج الذكية لا يمكن تفسيره بالعوامل التقنية وحدها، بل يتطلب دمج الأبعاد الأخلاقية والتنظيمية لفهم الظاهرة بصورة أشمل.

ثانياً- الدراسات المتعلقة بالقضايا الأخلاقية والتحديات لاستخدام نماذج الذكاء الاصطناعي التوليدي Generative AI في سياق التعليم العالي

تشير مراجعة الدراسات السابقة في موضوع القضايا الأخلاقية المرتبطة باستخدام نماذج الذكاء الاصطناعي التوليدي في التعليم العالي إلى تباين واضح في بؤرة الاهتمام؛ فبعضها ركّز على إبراز الإمكانيات التعليمية للنماذج الذكية، مثل تعزيز التعلم الذاتي، وتسريع اكتساب المعرفة، وتوفير موارد مخصصة تلبي احتياجات المتعلم (Baidoo-Anu & Owusu Ansah, 2023; Dimpere et al., 2023; Firaina & Sulisworo, 2023, 2023; Huallpa n.d.; Baskara, 2023; Liu et al., 2023; Malinka et al., 2023; Saxena et al., 2023) غير أن هذه الدراسات تعاملت مع الجانب الأخلاقي في حدود التوصية بالاستخدام المسؤول، دون تقديم تحليل معمق لآليات مواجهة التحديات، أو اقتراح سياسات مؤسسية واضحة.

في المقابل، سلّطت مجموعة أخرى من الدراسات الضوء على المخاطر الأخلاقية التي يثيرها الاعتماد المتزايد على Gen AI، مثل قضايا الموثوقية، والشفافية، والخصوصية، والأمان (Baidoo-Anu & Owusu Ansah, 2023; Liu et al., 2023; Malinka et al., 2023) وهنا لم تعد التحديات مجرد نتائج عرضية، بل مؤشرات على غياب أطر واضحة للمساءلة والتنظيم؛

فالغش والانتحال وصعوبة التحقق من صحة المحتوى وتسريب البيانات الحساسة، بل وحتى استغلال النماذج في هجمات سيبرانية، تكشف أن الإشكالية تتجاوز حدود التعليم لتتصل مباشرة بالأمن السيبراني والسياسات العامة.

أما على مستوى الدراسات التطبيقية، فقد أظهرت دراسة (Ahmed et al., 2024) جانباً مهماً من التوتر القائم بين وعود الذكاء الاصطناعي التوليدي، ومخاطره في التعليم العالي. فبينما أبرزت نتائجها أن هذه النماذج قادرة على تعزيز التعلم الشخصي من خلال توفير مواد مخصصة، ودعم أعضاء هيئة التدريس في إعداد المحتوى والتقييم، إلا أنها كشفت بأن هذه الفوائد تقتصر بالعديد من المخاطر، أبرزها تهديد النزاهة الأكاديمية، وتقويض مهارات التفكير النقدي والإبداعي، إضافةً إلى انتهاك الخصوصية والأمان. وهذا التناقض يعكس معضلة أعمق، فالنماذج الذكية ذاتها التي يُنتظر أن تدعم الاستقلالية التعليمية قد تتحول إلى وسيلة لإضعافها، وهذا يثير نقاشاً حول قدرة المؤسسات الأكاديمية على ضبط التقنية بحيث تعزز القيم التي تسعى إلى ترسيخها.

وفي سياقٍ مقارن، أظهرت دراسة (Mironova et al., 2024) أن المواقف الأخلاقية تجاه استخدام ChatGPT ليست موحدة؛ ففي خمس دول (لاتفيا، وليتوانيا، وأوزبكستان، وأوكرانيا، وبلغاريا) اعتبر معظم الطلبة أن هذا الاستخدام أخلاقي، بينما رأت أقلية عكس ذلك. هذا التباين يعكس أن الحكم الأخلاقي على التقنية يتأثر بالقيم والمعايير المحلية، لا باعتبارات عالمية مطلقة. الأمر الذي تؤكدُه أيضاً دراسة (Huallpa et al., 2023)، حيث أشارت إلى أن فقدان التفاعل البشري والتحيز يشكلان أبرز التحديات، لكنها أوضحت أن الحلول لا تكمن في رفض التقنية، بل في وضع سياسات مؤسسية واضحة لحماية البيانات وتقليل التحيز، أي أن المشكلة تكمن في الإدارة لا في النماذج الذكية. ويكشف هذا التباين أن الحكم الأخلاقي على التقنية ليس مطلقاً، بل يتأثر بالقيم الأخلاقية، والمعايير الاجتماعية السائدة في كل بيئة. ومن هنا تبرز الحاجة إلى أهمية مراعاة اختلاف السياقات بدلاً من افتراض مقاربة موحدة صالحة لجميع البيئات الأكاديمية.

أما دراسة (de Fine Licht, 2024) فقد ذهبت في اتجاه أكثر تشدداً، حيث ناقشت مبررات تقييد أو حتى حظر استخدام أدوات الذكاء الاصطناعي التوليدي في الجامعات.

واستندت الدراسة إلى أن هذه النماذج الذكية قد تهدد النزاهة الأكاديمية، وتقلل من التفاعل مع المادة التعليمية، خاصة في ظل ضعف وعي الطلبة بحقوقهم الرقمية، ومخاطر الوصول إلى بياناتهم. وبالرغم من منطقية هذا الطرح إلا أن فعاليته تظل موضع تساؤل متى ما غابت البدائل القادرة على تحويل الحظر إلى إصلاح، فالحظر وحده لا يعالج المعضلة، بل قد يدفع الطلبة لاستخدام هذه النماذج الذكية بطرق غير رسمية، وغير خاضعة لأي تنظيم. وهذا ما يثير معضلة تربوية وأخلاقية تتمثل في السؤال: هل تكمن الحماية في تقييد الوصول، أم في بناء وعي نقدي يتيح للطلبة إدارة علاقتهم بهذه النماذج الذكية بوعي ومسؤولية؟

ويتضح من هذا المشهد أنه ليس مجرد اختلاف في النتائج، بل جدل أعمق حول معنى نماذج الذكاء الاصطناعي التوليدي في الفضاء الأكاديمي، فهل هي أداة محايدة قابلة للتسخير وفق ما نرسمه لها من سياسات، أم أنها قوة مولدة لأسئلة جديدة حول القيم، والحرية، والمساءلة؟ إن تجاوز هذا الجدل يتطلب مقاربات تربط بين التقنية والإنسان، لتجعل من الذكاء الاصطناعي التوليدي أفقاً للابتكار المسؤول لا مجرد أداة للجدل الأخلاقي.

التعليق على الدراسات السابقة وما يميز الدراسة الحالية.

- من خلال استعراض الدراسات السابقة يتضح أن الدراسات التي تناولت قبول واستخدام نماذج الذكاء الاصطناعي التوليدي في التعليم العالي انقسمت إلى مسارين رئيسيين:
- المسار الأول ركّز على تطبيق النظرية الموحدة لقبول واستخدام التكنولوجيا بإصدارها (UTAUT/UTAUT2)، حيث برزت العوامل التقنية مثل الأداء المتوقع (PE) والجهد المتوقع (EE) باعتبارها المحرك الرئيسي للنية السلوكية (Bazelaïs et al., 2024; Biloš & Budimir, 2024; Tian et al., 2024) غير أن بعض المحاولات المحدودة أضافت عوامل جديدة مثل القلق (Budhathoki et al., 2024) أو قيمة التعلم (Bahadur et al., 2024) أو مخاطر الخصوصية (Alzahrani, 2025) ورغم أهميتها، ظلت هذه الإضافات جزئية وغير مدمجة في إطار شامل يوازن بين التقنية والأخلاقيات.
 - المسار الثاني عالج القضايا الأخلاقية، لكنه اعتمد غالباً على مراجعة الدراسات السابقة (Baidoo-Anu & Owusu Ansah, 2023; Liu et al., 2023; Malinka et al., 2023)، دون تقديم نماذج تفسيرية أو اختبارات ميدانية. بينما حاولت دراسات أخرى

اقترح أطر تنظيمية (de Fine Licht, 2024; Qadhi et al., 2024) لكنها بقيت نظرية ولم تُختبر عملياً.

هذا التباين يعكس فجوة بحثية واضحة تتمثل في غياب رؤية متكاملة تربط بين العوامل التقنية التي تفسر التبني والاستخدام، والعوامل الأخلاقية التي تضبط الاستخدام المسؤول. ومن هنا تأتي أهمية الدراسة الحالية، التي تسعى إلى سد هذه الفجوة من خلال: إضافة أربعة عوامل أخلاقية جديدة إلى UTAUT2 وهي: (المسؤولية الأخلاقية، والخصوصية المدركة، وأمن البيانات، والوعي الأخلاقي). وبناء نموذج متكامل يجمع بين البعدين التقني والأخلاقي. ومقارنة السياق العربي بالسياقات العالمية لتوضيح الفروق الثقافية والمؤسسية. ويوضح الجدول (١) الفرق بين الدراسات السابقة في السياقين العربي والدولي.

جدول (١): مقارنة بين الدراسات السابقة في السياقين العربي والدولي

السياق العربي (المملكة العربية السعودية، والشرق الأوسط)	السياقات الدولية (أوروبا، وآسيا)	جانب المقارنة	الفجوة البحثية
ركزت على العوامل التقنية الأساسية في امتداد النظرية الموحدة لقبول واستخدام UTAUT التكنولوجي (PE) ، مثل: الأداء المتوقع (EE) ، والجهد المتوقع والتأثير الاجتماعي (SI) (Alzahrani, 2025; Elshaer et al., 2024; Sobaih et al., 2024).	ركزت على العوامل التقنية مثل الأداء المتوقع ، والجهد المتوقع (PE) وتم اعتبارها المحرك (EE) الرئيسي للنية السلوكية (Bazelaï et al., 2024; Biloš & Budimir, 2024; Tian et al., 2024).	التركيز	اقتصرت على العوامل التقنية فقط، وهذا يمثل فجوة في تحليل الجانب الأخلاقي لاستخدام نماذج الذكاء الاصطناعي في Gen AI التوليدي التعليم العالي.
إضافة عامل أخلاقي واحد وهو عامل مخاطر	إضافة بعض العوامل	إضافة عوامل	لا يوجد إطار متكامل يجمع ما بين البعدين

إضافية	مثل عامل القلق أو قيمة التعلم (Bahadur et al., 2024; Budhathoki et al., 2024)، واقتراح أطر تنظيمية لم يتم اختبارها أو تطبيقها (de Fine Licht, 2024; Qadhi et al., 2024).	الخصوصية (Privacy Risk) (Alzaharani, 2025).	التقني، والأخلاقي في السياقين العربي والدولي.
المنهجية	مراجعات للدراسات السابقة تبرز الفرص والمزايا لنماذج الذكاء الاصطناعي التوليدي (Baidoo-Anu & Owusu Ansah, 2023; Liu et al., 2023).	المنهج الوصفي مثل Elshaer et al., 2024; Sobaih et al., 2024) الاستبانة كأداة رئيسية لجمع البيانات، مما يجعلها مشابهة في منهجيتها لعدد كبير من الدراسات الدولية إلا أنها لم تختبر الأبعاد الأخلاقية بشكل عملي.	الحاجة إلى دراسات عملية وتطبيقية، تربط بين العوامل التقنية والأخلاقية.

الإطار النظري

الذكاء الاصطناعي التوليدي (Generative Artificial Intelligence (GAI)

تعود أصول الذكاء الاصطناعي إلى الخمسينيات عندما تساءل Alan Turing إذا كان من الممكن لبرنامج كمبيوتر أن يتحدث مع عدة أشخاص دون أن يدركوا أن تفاعلهم صناعي،

وبعدها تم اعتبار هذا السؤال بمثابة الفكرة المبتكرة لتطبيقات الدردشة القائمة على الذكاء الاصطناعي التوليدي (Hasal et al., 2021; Limna et al., 2023; Qammar et al., 2023; Shahriar & Hayawi, 2023)

وتُعرف تطبيقات الذكاء الاصطناعي التوليدي بأنها أنظمة الحوار بين الإنسان والحاسوب عبر الإنترنت باللغة الطبيعية؛ وهي أنظمة محادثة ذكية يمكنها التواصل مع البشر باستخدام لغتهم الطبيعية، وفي الوقت الفعلي حيث تقوم برامج الدردشة الآلية بمعالجة مدخلات المستخدم، ومن ثم تقديم المخرجات ذات الصلة بالمدخلات في شكل نص باللغة الطبيعية (Hasal et al., 2021; Qammar et al., 2023)، ويمكن للذكاء الاصطناعي التوليدي التنبؤ بالنتيجة (إعطاء إجابة) بناءً على المحادثات التاريخية باستخدام نصوص سابقة مماثلة، ومجموعة بيانات مناسبة لنمذجة حوارات طويلة ومفتوحة المجال قريبة من اللغة البشرية المنطوقة، وكلما زاد عدد بيانات المحادثة المتوفرة حول موضوع معين فإن تطبيق الدردشة القائم على الذكاء الاصطناعي يقدم أداءً أفضل (Hasal et al., 2021)

ويعد "الذكاء الاصطناعي التوليدي" Generative AI أحد فروع الذكاء الاصطناعي AI (Feuerriegel et al., 2024)، ويشير مصطلح GenAI إلى خوارزميات التعلم الآلي المصممة للسماح للمستخدمين بإنشاء محتوى جديد من خلال طرح الأسئلة (المحفزات) ويمكن لهذه النماذج توفير معلومات واقعية، والإجابة على الأسئلة، وتعديل النصوص الموجودة (Biton & Segal, 2025)، وإنشاء بيانات أو محتوى جديد، توليد محتوى جديد وذو معنى مثل النصوص والصور ومقاطع الفيديو والصوت، وذلك من خلال البيانات التي تم تدريبه عليها (Biton & Segal, 2025; Feuerriegel et al., 2024).

وتقوم نماذج الذكاء الاصطناعي التوليدي GAI بإنتاج مخرجات جديدة بدلاً من مجرد إعادة إنتاج المدخلات وتصنيفها ومعالجتها وتحليلها (Mcminn, 2024) وهو ما تقوم به النظم الخبيرة Expert Systems (Feuerriegel et al., 2024). وأحدثت هذه النماذج انتشاراً كبيراً وثورة في طريقة عملنا وتواصلنا، ومن أمثلتها Dall-E2 وGPT-4 وCopilot (Venter et al., 2025).

ويرى الباحثون أن اعتماد تطبيقات الذكاء الاصطناعي التوليدي في مؤسسات التعليم العالي له تأثير إيجابي على تعزيز إنتاجية الطلبة، والكفاءة في العملية التعليمية (Dempere et

(al., 2023; Firaina & Sulisworo, 2023) حيث أنه يدعم طلبه الدراسات العليا والباحثين في البحث عن المعلومات والأفكار، والترجمة، وتحليل النصوص الأكاديمية وإنشاء الملخصات، ومساعدتهم في اكتساب رؤى واكتشافات جديدة (Febriyani et al., 2023; Ray, 2023). وبالرغم من هذه الإمكانيات، إلا أن الدراسات تؤكد أن استخدام هذه النماذج الذكية في التعليم العالي لا يزال في مراحله الأولى، ويصاحبه تحديات متعلقة بالمصادقية والنزاهة الأكاديمية وحماية البيانات، وهو ما يتطلب سياسات تنظيمية وأطر مؤسسية تضمن الاستخدام المسؤول والمتوافق مع القيم الأكاديمية (de Fine Licht, 2024; Huallpa et al., 2023; UNESCO, 2023).

النظرية الموحدة لقبول واستخدام التكنولوجيا (Unified Theory of Acceptance and Use of Technology (UTAUT))

قام (Venkatesh et al., 2003) بتطوير نموذج للنظرية الموحدة لقبول واستخدام التكنولوجيا (UTAUT1) وذلك بهدف قياس قبول واستخدام المستخدمين للتقنية، وتعد هذه النظرية تطوير لنماذج سابقة هدفت إلى فهم كيف ولماذا يتقبل الأفراد التكنولوجيا ويعتمدونها؛ ويتكون هذا النموذج من مجموعة من النظريات المستخدمة في دراسة سلوك المستخدم وقبوله للتكنولوجيا مثل: نظرية الفعل المبرر (Theory of Reasoned Action (TRA)، ونموذج قبول التكنولوجيا (Technology Acceptance Model (TAM)، ونظرية السلوك المخطط (TPB) Theory of Planned Behaviour، ودمج نموذج قبول التكنولوجيا ونظرية السلوك المخطط (MPCU) Computer Model of Personal (TAM, TPB)، ونموذج استخدام الكمبيوتر (Utilization، والنموذج التحفيزي (Motivational Model (MM)، والنظرية المعرفية الاجتماعية (Social Cognitive Theory (SCT).

وفي عام ٢٠١٢م طور (Venkatesh et al., 2012) امتداد للنظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT2 بالنسبة للمستهلكين، حيث تم تطوير نموذج UTAUT1 في الأصل لتحليل استخدام وقبول التكنولوجيا للموظفين، ولم يخصص للتقنيات التي يتعامل معها المستهلك، وتم إضافة عدداً من العوامل للنموذج الأول من النظرية UTAUT1، وهي: دافع المتعة Hedonic Motivation، وقيمة السعر Price Value، والعادة Habit، بالإضافة إلى حذف متغير طوعية الاستخدام من النظرية وذلك بهدف جعل النموذج أكثر اتساقاً مع البيئة التي تُقدم فيها

التكنولوجيا للمستهلكين. وتهدف النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT إلى شرح نية المستخدم في اعتماد واستخدام نظام معلومات أو تكنولوجيا معينة (Venkatesh et al., 2003, 2012).

وتعتمد الدراسة الحالية على نموذج يجمع ما بين عوامل امتداد النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT2 مثل: الأداء المتوقع، والتأثير الاجتماعي، والظروف الميسرة، والعادة؛ وعوامل أخلاقية مثل: المسؤولية الأخلاقية، والخصوصية المدركة، وأمن البيانات، والوعي الأخلاقي لتحقيق أهداف الدراسة. ويمثل الشكل (١) النموذج الذي تعتمد عليه الدراسة الحالية.

■ الأداء المتوقع Performance Expectancy

مدى اعتقاد الأفراد بأن استخدام التكنولوجيا سيحسن من أدائهم في المهام التي يقومون بها (Venkatesh et al., 2003).

■ التأثير الاجتماعي Social Influence

مدى تأثر الأفراد بأراء أقرانهم مثل الزملاء عند تبني التكنولوجيا الجديدة (Venkatesh et al., 2012).

■ الظروف الميسرة Facilitating Conditions

مدى توفر البنية التحتية، والدعم الفني، والموارد اللازمة لاستخدام التكنولوجيا بفعالية (Venkatesh et al., 2003).

■ العادة Habit

درجة اعتياد الأفراد على استخدامهم التكنولوجيا بشكل تلقائي بناءً على خبراتهم السابقة (Venkatesh et al., 2012).

■ عامل المسؤولية الأخلاقية Ethical Responsibility

ويستند هذا العامل إلى نظرية المسؤولية الأخلاقية لـ (P.F. Strawson (Todd, 2016، والتي ترى أن المسؤولية الأخلاقية تنبع من التفاعلات الاجتماعية وتوقعات المجتمع، وليس مجرد خاصية فردية. فالأفراد يُعدّون مسؤولين أخلاقياً لأن المجتمع ينظر إليهم كذلك، مما يعكس تحولاً نحو الفهم الاجتماعي للعلاقات الأخلاقية. وقد أضيف هذا العامل إلى UTAUT2 بهدف اعتبار أن استخدام نماذج الذكاء الاصطناعي التوليدي

مقبول بشكل أخلاقي. وتبرز أهمية هذا العامل في تفسير تأثير المعايير الاجتماعية والاستجابات العاطفية على تبني هذه النماذج وعلى إدراك الطلبة لمفهوم المسؤولية الأخلاقية (Todd, 2016).

عامل الخصوصية المدركة Perceived Privacy

ويستند هذا العامل إلى نظرية Altman's Conceptualization of Privacy التي تفسر الخصوصية بوصفها عملية ديناميكية لإدارة الحدود الشخصية والتحكم في مستوى التفاعل وكشف المعلومات (Knijnenburg, Xinruupagee, et al., 2022) ووفقاً لهذه النظرية فإن الأفراد مسؤولون عن قراراتهم المتعلقة بكيفية مشاركة معلوماتهم وتحديد الجهات التي يثقون بها. وقد تمت إضافة هذا العامل إلى امتداد نموذج UTAUT2 لقياس إدراك طلبة الدكتوراة في جامعة الملك سعود لمخاطر الخصوصية المرتبطة باستخدام نماذج الذكاء الاصطناعي التوليدي، وهو ما ينعكس على مستوى قبولهم لهذه النماذج وتبنيها في البيئة الأكاديمية.

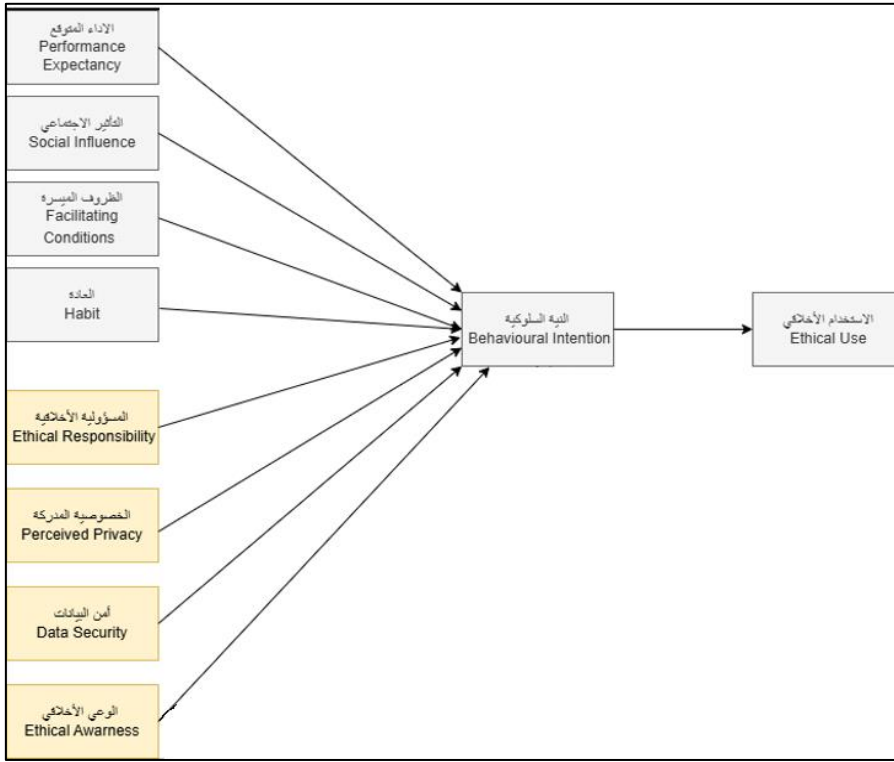
■ عامل أمن البيانات Data Security

ويستند هذا العامل إلى اللائحة العامة لحماية البيانات (GDPR) التي أقرها البرلمان الأوروبي عام ٢٠١٦ م، والتي وأصبحت إلزامية من مايو عام ٢٠١٨ م، بوصفها معيار عالمي لحماية البيانات الشخصية (GDPR.EU, n.d.)، وتم إضافة هذا العامل إلى نموذج UTAUT2 لفهم أثر الإجراءات الأمنية المتخذة لحماية بيانات الأفراد في تعزيز ثقتهم باستخدام نماذج GAI. حيث أن إدراك الطلبة لفاعلية تلك الإجراءات ينعكس على نواياهم السلوكية تجاه الاستخدام الأخلاقي لهذه النماذج واستعدادهم لاعتمادها في بيئتهم الأكاديمية.

■ عامل الوعي الأخلاقي Ethical Awareness

ويستند هذا العامل إلى نظرية انتشار الابتكارات (Diffusion of Innovations Theory) لـ Rogers (1962)، التي تفسر كيفية انتشار الأفكار والتقنيات الجديدة عبر مراحل تبدأ بالمعرفة وتنتهي بالتأكد، مع تحديد سمات الابتكار مثل الميزة النسبية، والتوافق، والتعقيد، وقابلية التجربة، وقابلية الملاحظة (Malouf, 2023). ويعرف الوعي الأخلاقي بأنه مدى إدراك الطلبة للمبادئ الأخلاقية والمخاطر المرتبطة باستخدام نماذج الذكاء الاصطناعي التوليدي. وبالاستناد

إلى نظرية انتشار الابتكارات، يفترض أن ارتفاع مستوى الوعي الأخلاقي يزيد من احتمالية تبني الاستخدام الأخلاقي لهذه النماذج الذكية في البيئة الأكاديمية.



شكل (١): النموذج الذي تعتمد عليه الدراسة

الأطر الفلسفية للعوامل الأخلاقية

لا يقتصر تحليل العوامل الأخلاقية (المسؤولية الأخلاقية، والخصوصية المدركة، وأمن البيانات، والوعي الأخلاقي) التي تم إضافتها إلى امتداد النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT2 على الأطر النظرية فقط، بل يمكن تحليل هذه العوامل من خلال إطار فلسفي، وربطها بالمدارس الفلسفية الكبرى في الأخلاق. حيث يتيح هذا الربط تقديم فهم أكثر عمق وشمول لنوايا طلبة الدراسات العليا السلوكية، واستخدامهم الأخلاقي لنماذج الذكاء الاصطناعي التوليدي في البيئة الأكاديمية، وهذا يتيح التعرف على التوجه السائد في الممارسات الأكاديمية، والتعرف على الحالات الاستثنائية التي قد تظهر نتيجة للضغوط المؤسسية، أو الدوافع الشخصية للطلبة.

وفيما يلي يتم عرض الأساس الفلسفي لكل عامل، من خلال مقارنته بالمبادئ الأخلاقية مثل الفلسفة الأخلاقية لكانط (Kant's Moral Philosophy (Kant, 2022)، ونظرية النفعية (Utilitarianism (Stanford Encyclopedia of Philosophy, 2025)، ونظرية الفضيلة (Virtue Ethics (Stanford Encyclopedia of Philosophy, 2022) لتعزيز العمق الفلسفي للدراسة الحالية.

عامل المسؤولية الأخلاقية: يرتبط هذا العامل ارتباطاً وثيقاً بالفلسفة الأخلاقية لكانط. وتقوم هذه الفلسفة على أن الأفراد ملزمون باحترام قواعد أخلاقية عامة تحكم السلوك الانساني. ويرى كانط أن الفعل لا يكون أخلاقياً إلا إذا استند إلى الواجب، أي إذ أمكن تعميمه كقانون صالح للجميع، بغض النظر عن النتائج أو المصالح، ويسمى بالأمر القاطع (Categorical Imperative). ووفقاً لهذه الفلسفة، فإن الفعل يكون أخلاقياً إذا أصبح قاعدة عامة قابلة للتطبيق، وهذا يعني أن جميع الأفعال غير الأخلاقية غير عقلانية لأنها تنافي الأمر القاطع (Kant, 2022).

وبناءً على ذلك، فإن التزام الطلبة بالنزاهة الأكاديمية وعدم استخدام نماذج الذكاء الاصطناعي التوليدي لأغراض غير أخلاقية مثل الغش أو الانتحال يمثل تطبيق عملي لمبدأ الواجب عند كانط. وقد أظهرت بعض الدراسات في التعليم العالي أن بعض الطلبة يعبرون عن التزامهم بالنزاهة الأكاديمية بوصفها قيمة أخلاقية داخلية، وليس مجرد التزام لتجنب العقوبات (Brown et al., 2020; Lau, 2021)، كما أن استخدام تقنيات الذكاء الاصطناعي يجب أن يكون بشكل متوازن ما بين تحقيق الفائدة، والالتزام بالمعايير الأخلاقية للمؤسسات التعليمية (Balalle & Pannilage, 2025).

عامل الخصوصية المدركة: يرتبط عامل الخصوصية المدركة بالطريقة التي يوازن بها الطلبة بين الفوائد التعليمية التي يحققها استخدام نماذج الذكاء الاصطناعي التوليدي Gen AI في العمليات الأكاديمية، مثل التلخيص، وتوليد الأفكار، والترجمة، وحل الواجبات، وإعداد العروض وغيرها؛ وبين المخاطر المرتبطة بكشف البيانات الشخصية أو إساءة استخدامها (Chan & Hu, 2023). ويعكس هذا التوازن نظرية النفعية (Utilitarianism)، والتي ترى أن الجودة الأخلاقية لأي فعل أو سياسة تعتمد بشكل كلي على نتائجه، أو القيمة الناتجة عنه.

ويعد الفعل مقبول أخلاقياً إذا كانت منافعه تفوق أضراره (Stanford Encyclopedia of Philosophy, 2025). لذا فإن الطلبة الذين يقبلون مشاركة بياناتهم الشخصية مع النماذج الذكية، يتبنون نهجاً قائماً على تعظيم المنفعة التعليمية، وتقليل الضرر المتعلق بالخصوصية. إضافةً إلى أن عامل الخصوصية المدركة لا يقتصر فقط على النظرية النفعية Utilitarianism، بل يرتبط أيضاً بالأمر الأخلاقي Categorical Imperative عند كانط (Kant, 2022)، والذي يؤكد على أن حماية خصوصية الأفراد واجب أخلاقي في ذاته، وليس مجرد وسيلة لتحقيق منفعة تعليمية. فالاستخدام غير المنضبط للبيانات الشخصية أو استغلالها قد يُعد انتهاكاً لمبدأ التعامل مع الإنسان كغاية لا كوسيلة (Bankins & Formosa, 2023; Hanna & Kazim, 2021).

عامل أمن البيانات: يرتبط عامل أمن البيانات بثقة الطلبة في نماذج الذكاء الاصطناعي التوليدي لتأمين بياناتهم الشخصية والأكاديمية. فمن منظور النفعية Utilitarianism يُعد تعزيز أمن البيانات أساس عند استخدام النماذج الذكية، وذلك لدوره في تحقيق أكبر قدر من المنفعة للطلبة، وذلك من خلال التقليل من المخاطر مثل تسريب البيانات، أو إساءة استخدامها (Stanford Encyclopedia of Philosophy, 2025).

ويمكن النظر إلى عامل أمن البيانات أيضاً من خلال الفلسفة الكانطية؛ حيث يفرض على المؤسسات واجب أخلاقي يتمثل في احترام الأشخاص كغايات في ذاتهم، وذلك من خلال حماية بياناتهم من الانتهاك أو الاستغلال (Kant, 2022)، فالتهاون في حماية أمن البيانات حسب فلسفة كانط يعتبر انتهاك لواجب أخلاقي أساسي يهدد كرامة الأفراد وحقوقهم، حتى لو لم يترتب عليه ضرر مباشر في بعض الحالات. كما أن المؤسسات الأكاديمية التي تضع سياسات صارمة لحوكمة البيانات، وتعزيز الشفافية في كيفية جمعها واستخدامها، تعزز ثقة الطلبة في تبني نماذج الذكاء الاصطناعي التوليدي GAI بشكل آمن ومسؤول (Rana et al., 2024).

عامل الوعي الأخلاقي: يرتبط عامل الوعي الأخلاقي بإدراك الطلبة للمخاطر، والقيم المرتبطة باستخدام نماذج الذكاء الاصطناعي التوليدي GAI في العمليات الأكاديمية، واتخاذ القرارات التي تعكس الفضيلة مثل النزاهة، والصدق، والمسؤولية. وعامل الوعي الأخلاقي يرتبط بنظرية الفضيلة Virtue Ethics (Stanford Encyclopedia of Philosophy, 2022)، والتي تركز على

تكوين شخصية فاضلة قادرة على اتخاذ قرارات حكيمة في المواقف الأخلاقية المعقدة (Stanford Encyclopedia of Philosophy, 2022)، ووفقاً لنظرية الفضيلة فإن عامل الوعي الأخلاقي يركز على تنمية الفضيلة في الطلبة لتمكينهم من اتخاذ قرارات مسؤولة. ويمثل الجدول (٢) الأطر الفلسفية للعوامل الأخلاقية التي تم اضافتها لامتداد النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT2

جدول (٢): الأطر الفلسفية للعوامل الأخلاقية

العامل الأخلاقي	النظرية الفلسفية/المدرسة	الربط بالإطار الفلسفي
المسؤولية الأخلاقية	الفلسفة الأخلاقية لكانط Kant's Moral Philosophy	الفعل الأخلاقي قائم على الواجب غير المشروط وفق مبدأ الأمر القاطع؛ حيث أن التزام الطلبة بالنزاهة الأكاديمية في الغش يمثل GenAI وعدم استخدام Kant, 2022). تطبيق لمبدأ الواجب (
الخصوصية المدركة	نظرية النفعية Utilitarianism، والفلسفة الأخلاقية لكانط Kant's Moral Philosophy	مشاركة البيانات مبررة إذا كانت المنافع التعليمية أكبر من الأضرار (Chan & Hu, 2023b). الفلسفة الأخلاقية: حماية الخصوصية واجب أخلاقي قائم بذاته، واستغلال البيانات يعد انتهاكاً لكرامة الفرد. (Kant, 2022).
أمن البيانات	نظرية النفعية Utilitarianism ، والفلسفة الأخلاقية لكانط	تعزيز أمن البيانات يحقق منفعة جماعية من خلال تقليل المخاطر. الفلسفة الأخلاقية: حماية البيانات واجب أخلاقي لحفظ الكرامة الإنسانية (Kant, 2022).
الوعي الأخلاقي	Virtue Ethics نظرية الفضيلة	تركز على تنمية الأخلاق الفاضلة مثل النزاهة، والمسؤولية لبناء شخصية أكاديمية قادرة على اتخاذ قرارات مسؤولة (Hagendorff, 2022).

المنهجية وتصميم أداة الدراسة

اعتمدت الدراسة على المنهج الوصفي باعتباره من المناهج الأكثر ملاءمة لتحليل الظواهر الاجتماعية والسلوكية وفهم العوامل المؤثرة فيها، حيث أن الدراسات الوصفية تُعد من أكثر التصاميم الكمية شيوعاً لوصف خصائص المجموعات البشرية، والسلوكيات المرتبطة بها (Creswell, 2014).

ولبناء أداة الدراسة تم الاسترشاد بدراسات حديثة مثل (Menon & Shilpa, 2023; Tian et al., 2024) التي اعتمدت فقرات لقياس الأداء المتوقع، والظروف الميسرة. ودراسة (Huallpa et al., 2023) التي تناولت الاعتبارات الأخلاقية في دمج الذكاء الاصطناعي بالتعليم. بالإضافة إلى دراسة (Sebastian, 2023) التي ركزت على صياغة مؤشرات للوعي الأخلاقي. وبناءً على ذلك، تم تصميم الاستبانة بحيث تغطي:

- العوامل التقنية: الأداء المتوقع، والتأثير الاجتماعي، والظروف الميسرة، والعادة.
- العوامل الأخلاقية: المسؤولية الأخلاقية، والخصوصية المدركة، وأمن البيانات، والوعي الأخلاقي.

وصيغت الفقرات بالاستناد إلى مقاييس معتمدة في دراسات سابقة، حيث تم تكييف الفقرات التقنية (Bouteraa et al., 2024; Venkatesh et al., 2003; Yilmaz et al., 2023)، بينما صيغت الفقرات الأخلاقية بالاسترشاد بـ (Huallpa et al., 2023; Sebastian, 2023).

وتم استخدام مقياس ليكرت الخماسي (five-point Likert scale) لقياس درجة استجابة الطلبة، وهو المقياس الأكثر شيوعاً في الدراسات الاجتماعية والكمية المماثلة (Sobaih et al., 2024). وتم اختياره لكونه أداة معيارية وموثوقة على نطاق واسع في البحوث الاجتماعية؛ إذ يمكن من تحويل التصورات الذاتية إلى بيانات كمية قابلة للتحليل، مما يعزز دقة المقارنات الإحصائية واستخلاص النتائج (Koo & Yang, 2025).

ولضمان صدق الأداة، عُرضت الاستبانة على مجموعة من الأساتذة المحكمين وتم تعديلها وفق ملاحظاتهم (الصدق الظاهري)، كما تم التحقق من الصدق الداخلي من خلال معاملات الارتباط بين الفقرات والمحاور. وقد تم التحقق من ثبات الأداة باستخدام معامل

كرونباخ ألفا، حيث بلغت قيمته (٠,٩٥)، وهو ما يُعد مؤشراً مرتفعاً على الاتساق الداخلي (Tavakol & Dennick, 2011).

عينة الدراسة وجمع البيانات

تم استخدام العينة القصدية لاختيار طلبة الدراسات العليا (مرحلة الدكتوراه) في قسم علم المعلومات بجامعة الملك سعود لارتباط تخصصهم بنظم المعلومات والتقنيات وهذا يزيد من احتمال تعاملهم الفعلي مع تطبيقات الذكاء الاصطناعي التوليدي في سياق البحث والدراسة. حيث تدعم بعض الدراسات السابقة مثل (Alzahrani, 2025; Menon & Shilpa, 2023) اختيار طلبة التخصصات المعلوماتية، نظراً لارتفاع معدلات استخدامهم لتقنيات الذكاء الاصطناعي في التعليم والبحث، وهذا يعزز من مناسبة اختيارهم لهذه الدراسة.

واعتمدت الدراسة على عينة متجانسة مكونة من ١٥ طالب وطالبة في مرحلة الدكتوراه من قسم علم المعلومات بجامعة الملك سعود، وذلك بهدف التركيز على فئة محددة تشترك في خصائص أكاديمية موحدة، وهذا يتوافق مع ما أشار إليه (Palinkas et al., 2015) من أن العينات المتجانسة تعتبر مناسبة في الدراسات الاستكشافية التي تهدف إلى فهم سلوكيات فئة معينة بعمق أكبر. حيث يساهم هذا النوع من العينات في تقليل التباين، مما يسمح باستكشاف أكثر تركيز وتفصيل للظاهرة محل الدراسة.

وتم الحصول على موافقة لجنة أخلاقيات البحث العلمي في جامعة الملك سعود، وإرسال الاستبانة بشكل إلكتروني للطلبة، وتم إعلامهم بأن المشاركة طوعية وتستخدم البيانات لأغراض بحثية فقط، وقد بلغ عدد الاستجابات المقبولة ١٥ استجابة مكتملة تمثل مجتمع الدراسة المستهدف.

التحليل الإحصائي

تم استخدام الحزمة الإحصائية للعلوم الاجتماعية SPSS (Statistical Package for Social Sciences) في تحليل بيانات الاستبانة، حيث تم إجراء التحليلات الإحصائية الوصفية والاستنتاجية المناسبة لطبيعة أسئلة الدراسة. وشملت التحليلات المستخدمة:

- المتوسطات الحسابية والانحرافات المعيارية لتحليل استجابات المشاركين على محاور الاستبانة.
- الانحدار الخطي البسيط لقياس أثر العوامل التقنية، والعوامل الأخلاقية ككل على النية السلوكية.
- الانحدار الخطي المتعدد لفحص دور كل عامل تقني على حدة (الأداء المتوقع، والتأثير الاجتماعي، والظروف الميسرة، والعادة)، وفحص دور كل عامل أخلاقي على حدة (المسؤولية الأخلاقية، والخصوصية المدركة، وأمن البيانات، والوعي الأخلاقي) وتأثير تلك العوامل على النية السلوكية.

نتائج الدراسة

أولاً- صدق أداة الدراسة

• صدق الاتساق الداخلي (Internal consistency Validity):

بعد التأكد من الصدق الظاهري لأداة الدراسة تم تطبيقها على عينة الدراسة، ومن ثم التحقق من صدق المقياس عن طريق حساب معامل ارتباط بيرسون لمعرفة الصدق الداخلي للاستبانة حيث تم حساب معامل الارتباط بين درجة كل عبارة من عبارات الاستبانة بالدرجة الكلية للمحور الذي تنتمي إليه العبارة كما توضح ذلك الجداول التالية:

جدول (٣): معاملات ارتباط بيرسون لعبارات محاور العوامل التقنية بالدرجة الكلية للمحور

المحور الأول		المحور الثاني		المحور الثالث		المحور الرابع	
رقم	معامل الارتباط	رقم	معامل الارتباط	رقم	معامل الارتباط	رقم	معامل الارتباط
١	**٠,٩٤٠	١	**٠,٦٠٨	١	**٠,٧٠٦	١	**٠,٩٣٣
٢	**٠,٩٠٨	٢	**٠,٨٣٥	٢	**٠,٧٠٦	٢	**٠,٩٦٥
٣	**٠,٩٦١	٣	**٠,٦٧٦	٣	*٠,٦٠١	٣	**٠,٩٤٩

** دال عند مستوى الدلالة ٠,٠١ فأقل

يتضح من الجدول رقم (٣) أن قيم معامل ارتباط كل عبارة من العبارات مع المحور الذي تنتمي له موجبة ودالة إحصائياً عند مستوى الدلالة (٠,٠٥) أو (٠,٠١) فأقل مما يدل على صدق اتساقها مع المحور.

جدول (٤): معاملات ارتباط بيرسون لعبارات محور العوامل الأخلاقية بالدرجة الكلية للمحور

المحور الأول		المحور الثاني		المحور الثالث		المحور الرابع	
رقم	معامل الارتباط	رقم	معامل الارتباط	رقم	معامل الارتباط	رقم	معامل الارتباط
١	**٠,٩٤١	١	**٠,٨٣٦	١	**٠,٩١٨	١	**٠,٨٧٧
٢	**٠,٧٦٦	٢	*٠,٦١٢	٢	**٠,٩٠٧	٢	**٠,٨٩١
٣	**٠,٨٠٦	٣	**٠,٩١٣	٣	**٠,٩٠١	٣	**٠,٨٨٩

** دال عند مستوى الدلالة ٠,٠١ فأقل

يتضح من جدول (٤) أن قيم معامل ارتباط كل عبارة من العبارات مع المحور الذي تنتمي له موجبة ودالة إحصائية عند مستوى الدلالة (٠,٠٥) أو (٠,٠١) فأقل مما يدل على صدق اتساقها مع المحور.

جدول (٥): معاملات ارتباط بيرسون لعبارات محور تأثير عامل النية السلوكية (Behavioral Intention) على استخدام طلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي بالدرجة الكلية للمحور

رقم العبارة	معامل الارتباط
١	**٠,٩٠٥
٢	**٠,٩٥٨
٣	**٠,٩٢٥

** دال عند مستوى الدلالة ٠,٠١ فأقل

يتضح من الجدول رقم (٥) أن قيم معامل ارتباط كل عبارة من العبارات مع المحور الذي تنتمي له موجبة ودالة إحصائية عند مستوى الدلالة (٠,٠١) فأقل مما يدل على صدق اتساقها مع المحور.

جدول (٦): معاملات ارتباط بيرسون لعبارات محور تأثير عامل الاستخدام الأخلاقي (Ethical Use) على استخدام طلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي بالدرجة الكلية للمحور

رقم العبارة	معامل الارتباط
١	*.٠٠٦٢٠
٢	**٠.٠٨٩٣
٣	**٠.٠٦٧٧

** دال عند مستوى الدلالة ٠.٠١ فأقل

يتضح من جدول (٦) أن قيم معامل ارتباط كل عبارة من العبارات مع المحور الذي تنتمي له موجبة ودالة إحصائية عند مستوى الدلالة (٠.٠٥) أو (٠.٠١) فأقل مما يدل على صدق اتساقها مع المحور.

جدول (٧): معاملات ارتباط بيرسون لمحاور الاستبانة بالدرجة الكلية لأداة الدراسة

المحور	معامل الارتباط
العوامل التقنية	**٠.٠٦٦٤
العوامل الأخلاقية	**٠.٠٧٦١
تأثير عامل النية السلوكية (Behavioral Intention) على استخدام طلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي	**٠.٠٧٠٨
تأثير عامل الاستخدام الأخلاقي (Ethical Use) على استخدام طلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي	**٠.٠٦٤٩

** دال عند مستوى الدلالة ٠.٠١ فأقل

من جدول (٧) نجد أن جميع معاملات ارتباط بيرسون بين المحاور والدرجة الكلية للاستبانة موجبة ودالة إحصائياً عند مستوى الدلالة (٠,٠١)، حيث كان الحد الأدنى (٠,٦٤٩) لمعاملات الارتباط، فيما كان الحد الأعلى (٠,٧٦١)، وعليه فإن جميع المحاور متسقة داخلياً مع الدرجة الكلية، مما يثبت صدق الاتساق الداخلي للاستبانة.

ثبات أداة الدراسة (Reliability):

تم التحقق من ثبات أداة الدراسة عن باستخدام معادلة ألفا كرونباخ (cronbach,s) (Alpha(α))، وتعتبر من أشهر المقاييس المستخدمة لقياس الثبات الداخلي (Tavakol & Dennick, 2011) عن طريق حساب درجة ثبات كل محور من محاور الدراسة، وكذلك حساب قيمة الثبات الكلي لأداة الدراسة.

جدول (٨): معامل ألفا كرونباخ لقياس ثبات أداة الدراسة

القياسات	عدد العبارات	محاور الاستبانة
٠,٧٢٨	١٢	العوامل التقنية
٠,٨٩٠	١٢	العوامل الأخلاقية
٠,٩٢٠	٣	تأثير عامل النية السلوكية (Behavioral Intention) على استخدام طلبية الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي
٠,٨٦٦	٣	تأثير عامل الاستخدام الأخلاقي (Ethical Use) على استخدام طلبية الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي
٠,٨٥٨	٣٠	الثبات العام

يتضح من جدول (٨) أن معاملات الثبات لمحاور الاستبانة تراوحت بين (٠,٧٣ - ٠,٩٢)، وأن معامل الثبات الكلي بلغ (٠,٩٥)، وذلك يدل على مستوى عالي من موثوقية الاتساق الداخلي للمقياس التدريجي الخاص بهذه العينة، وتعتبر القيم التي تزيد عن (٠,٧٠) مقبولة.

معييار الحكم على نتائج الدراسة:

جدول (٩): درجات فئات معيار نتائج الدراسة وحدودها وفقاً لمقياس ليكرت الخماسي

الدرجة	معييار الحكم على النتائج	فئة المتوسط	
		من	إلى
٥	أو افق بشدة	٤,٢١	٥
٤	أو افق	٣,٤١	٤,٢٠
٣	أو افق إلى حد ما	٢,٦١	٣,٤٠
٢	لا أو افق	١,٨١	٢,٦٠
١	لا أو افق بشدة	١	١,٨٠

الأساليب الإحصائية:

التكرارات والنسب المئوية (Percentage & Frequencies): المتوسط الحسابي الموزون (المرجح) (Weighted Mean): المتوسط الحسابي (Mean): (متوسط متوسطات العبارات)، الانحراف المعياري (Standard Deviation) معامل ارتباط بيرسون (Pearson): لقياس الاتساق الداخلي بين عبارات الأداة (الاستبانة) وكل محور تنتمي إليه، وكذلك لتوضيح العلاقة في فروض الدراسة، معامل الثبات ألفا كرونباخ (α) (cronbach,s Alpha): لحساب معامل ثبات أداة الدراسة، تحليل الانحدار المتعدد (Regression) للتعرف على العوامل المؤثرة على الاستخدام الأخلاقي لنماذج الذكاء الاصطناعي التوليدي في قسم علم المعلومات بجامعة الملك سعود باستخدام امتداد النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT2، تحليل التباين الأحادي (One Way ANOVA) لتوضيح دلالة الفروق في إجابات أفراد عينة الدراسة طبقاً إلى اختلاف متغيراتهم التي تنقسم إلى أكثر من فئتين، حيث تم استخدام هذا الطريقة لمعرفة دلالة الفروق بين المتوسطات الحسابية وفقاً لمتغير (الخبرة التقنية)، واختبار (أقل فرق معنوي) (scheffe) لتوضيح دلالة الفروق في إجابات أفراد عينة الدراسة بين فئات متغير الخبرة التقنية، في حال أظهر اختبار تحليل التباين وجود فروق بين فئات هذا المتغير.

ثانياً. النتائج المتعلقة بوصف محاور الدراسة

المحور الأول: "ما تأثير عامل الأداء المتوقع على استخدام طلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي؟"

جدول (١٠): إجابات أفراد عينة الدراسة على عبارات محور الأداء المتوقع مرتبه تنازلياً

حسب متوسطات الإجابة

م	العبارات	المتوسط الحسابي	الانحراف المعياري	درجة الاستجابة	الترتيب
١	أعتقد أن استخدام نماذج الذكاء الاصطناعي التوليدي تعزز من جودة إنجازي الأكاديمي.	٤,٣٣	٠,٦١٧	أوافق بشدة	٣
٢	تساعدني نماذج الذكاء الاصطناعي التوليدي على إتمام مهماتي الأكاديمية بكفاءة.	٤,٤٠	٠,٦٣٢	أوافق بشدة	٢
٣	أرى أن نماذج الذكاء الاصطناعي التوليدي توفر حلولاً ذكية تدعم مهماتي الأكاديمية.	٤,٤٠	٠,٥٠٧	أوافق بشدة	١
	المتوسط العام	٤,٣٨	٠,٥٤٧	أوافق بشدة	

أظهرت نتائج جدول (١٠) أن طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود موافقون بشدة على تأثير عامل الأداء المتوقع في استخدامهم الأخلاقي لنماذج الذكاء الاصطناعي التوليدي، بمتوسط عام (٤,٣٨ من ٥,٠٠) وانحراف معياري (٠,٥٤٧). وقد جاءت العبارات الثلاث جميعها ضمن الفئة العليا (أوافق بشدة)، حيث حلت عبارة "أرى أن نماذج الذكاء الاصطناعي التوليدي توفر حلولاً ذكية تدعم مهماتي الأكاديمية" في المرتبة الأولى بمتوسط ٤,٤٠، تلتها عبارة "تساعدني هذه النماذج على إتمام مهماتي الأكاديمية بكفاءة" بمتوسط ٤,٤٠، ثم عبارة "أعتقد أن استخدام نماذج الذكاء الاصطناعي التوليدي تعزز من جودة

إنجازي الأكاديمي" بمتوسط ٤,٣٣. وتشير هذه النتائج إلى تجانس مرتفع في استجابات أفراد العينة، وهذا يعكس إدراك إيجابي واضح لقيمة الأداء المتوقع في تعزيز الاستخدام الأخلاقي لهذه النماذج الذكية.

المحور الثاني: "ما تأثير عامل التأثير الاجتماعي على استخدام طلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي؟"
جدول (١١): إجابات أفراد عينة الدراسة على عبارات محور التأثير الاجتماعي مرتبة تنازلياً

حسب متوسطات الإجابة

م	العبارات	المتوسط الحسابي	الانحراف المعياري	درجة الاستجابة	الترتيب
١	يشجعني زملائي على استخدام نماذج الذكاء الاصطناعي التوليدي في مهام الأكاديمية.	٣,٨٠	٠,٨٦٢	أوافق	٣
٢	ألاحظ أن استخدام نماذج الذكاء الاصطناعي التوليدي منتشر على نطاق واسع في بيئي الأكاديمية.	٤,٠٠	١	أوافق	١
٣	أشعر أن هناك قبول اجتماعي لاستخدام نماذج الذكاء الاصطناعي التوليدي في جامعة الملك سعود.	٣,٨٠	٠,٧٧٥	أوافق	٢
	المتوسط العام	٣,٨٧	٠,٦٢٧	أوافق	

أظهرت نتائج جدول (١١) أن طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود موافقون على تأثير عامل التأثير الاجتماعي في استخدامهم الأخلاقي لنماذج الذكاء الاصطناعي التوليدي، بمتوسط عام (٣,٨٧ من ٥,٠٠) وانحراف معياري (٠,٦٢٧). وقد جاءت العبارات الثلاث جميعها ضمن الفئة الرابعة (أوافق)، حيث حلت عبارة "ألاحظ أن استخدام

نماذج الذكاء الاصطناعي التوليدي منتشر على نطاق واسع في بيئي الأكاديمية "في المرتبة الأولى بمتوسط (٤,٠٠)، تلتها عبارة "أشعر أن هناك قبول اجتماعي لاستخدام نماذج الذكاء الاصطناعي التوليدي في جامعة الملك سعود" بمتوسط (٣,٨٠)، ثم عبارة "يشجعني زملائي على استخدام نماذج الذكاء الاصطناعي التوليدي في مهام الأكاديمية" بمتوسط (٣,٨٠). وتشير هذه النتائج إلى وجود درجة متقاربة من الموافقة بين أفراد العينة، وهو ما يعكس دوراً ملحوظاً للتأثير الاجتماعي في تشكيل استخدامهم الأخلاقي لهذه النماذج.

المحور الثالث: "ما تأثير عامل الظروف الميسرة على استخدام طلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي؟"
جدول (١٢): إجابات أفراد عينة الدراسة على عبارات محور الظروف الميسرة مرتبة تنازلياً

حسب متوسطات الإجابة

م	العبارات	المتوسط الحسابي	الانحراف المعياري	درجة الاستجابة	الترتيب
١	أمتلك المهارات الأساسية التي تؤهلني لاستخدام نماذج الذكاء الاصطناعي التوليدي.	٤,٤٠	٠,٦٣٢	أوافق بشدة	١
٢	أستطيع الوصول بسهولة إلى نماذج الذكاء الاصطناعي التوليدي عند حاجتي إليها.	٤,٤٠	٠,٦٣٢	أوافق بشدة	١
٣	تتوفر بنية تقنية مناسبة في جامعة الملك سعود تدعم استخدام نماذج الذكاء الاصطناعي التوليدي.	٣,٤٠	١,١٢١	محايد	٢
	المتوسط العام	٤,٠٧	٠,٥٢٣	أوافق	

أظهرت نتائج جدول (١٢) أن طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود موافقون على تأثير عامل الظروف الميسرة في استخدامهم الأخلاقي لنماذج الذكاء الاصطناعي التوليدي، بمتوسط عام (٤,٠٧ من ٥,٠٠) وانحراف معياري (٠,٥٢٣). وقد توزعت استجابات العينة بين الفئة العليا (أوافق بشدة) والفئة الوسطى (محايد)، حيث جاءت عبارتا "أمتلك المهارات الأساسية التي تؤهلني لاستخدام نماذج الذكاء الاصطناعي التوليدي" و "أستطيع الوصول بسهولة إلى نماذج الذكاء الاصطناعي التوليدي عند حاجتي إليها" في المرتبة الأولى بمتوسط (٤,٤٠)، بينما جاءت عبارة "تتوفر بنية تقنية مناسبة في جامعة الملك سعود تدعم استخدام نماذج الذكاء الاصطناعي التوليدي" في المرتبة الثانية بمتوسط (٣,٤٠). وتشير هذه النتائج إلى أن توفر المهارات الفردية وسهولة الوصول يسهمان بقوة في تيسير الاستخدام، في حين يظهر عامل البنية التقنية مستوى حياد نسبي بين أفراد العينة.

المحور الرابع: "ما تأثير عامل العادة على استخدام طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي؟"

جدول (١٣): إجابات أفراد عينة الدراسة على عبارات محور العادة مرتبة تنازلياً حسب

متوسطات الإجابة

م	العبارات	المتوسط الحسابي	الانحراف المعياري	درجة الاستجابة	الترتيب
١	جرت العادة أن استخدم نماذج الذكاء الاصطناعي التوليدي في مهام الأكاديمية.	٣,٤٧	٠,٩٩	أوافق	١
٢	أستخدم نماذج الذكاء الاصطناعي التوليدي بشكل تلقائي عند أداء مهام الأكاديمية.	٣,٣٣	١,١١٣	محايد	٢
٣	يعد استخدام نماذج الذكاء الاصطناعي التوليدي جزء من	٣,١٣	١,٣٥٦	محايد	٣

				روتين التعامل مع مهام الأكاديمية.	
	محايد	١,٠٩٤	٣,٣١	المتوسط العام	

أظهرت نتائج جدول (١٣) أن طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود جاءت استجاباتهم في الفئة المحايدة تجاه تأثير عامل العادة في استخدامهم الأخلاقي لنماذج الذكاء الاصطناعي التوليدي، بمتوسط عام (٣,٣١ من ٥,٠٠) وانحراف معياري (١,٠٩٤). وقد توزعت العبارات بين الموافقة والحياد، حيث جاءت عبارة "جرت العادة أن استخدم نماذج الذكاء الاصطناعي التوليدي في مهام الأكاديمية" في المرتبة الأولى بمتوسط (٣,٤٧) ضمن فئة (أوافق)، بينما جاءت عبارتنا "أستخدم نماذج الذكاء الاصطناعي التوليدي بشكل تلقائي عند أداء مهام الأكاديمية" و"يعد استخدام نماذج الذكاء الاصطناعي التوليدي جزءاً من روتين التعامل مع مهام الأكاديمية" في المرتبتين الثانية والثالثة بمتوسط (٣,٣٣) و(٣,١٣) ضمن فئة (محايد). وتشير هذه النتائج إلى أن عامل العادة لم يتبلور بقوة لدى أفراد العينة، مما يعكس أن الاستخدام الأخلاقي لهذه النماذج لا يزال في طور الممارسة غير الروتينية. المحور الخامس: "ما تأثير عامل المسؤولية الأخلاقية على استخدام طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي؟"

جدول (١٤): إجابات أفراد عينة الدراسة على عبارات محور المسؤولية الأخلاقية مرتبه

تنازلياً حسب متوسطات الإجابة

م	العبارات	المتوسط	الانحراف	درجة	الترتيب
---	----------	---------	----------	------	---------

	الحسابي	المعياري	الاستجابة	
١	٤,٨٠	٠,٤١٤	أوافق بشدة	٢
٢	٤,٨٠	٠,٤١٤	أوافق بشدة	٢
٣	٤,٨٧	٠,٣٥٢	أوافق بشدة	١
	٤,٨٢	٠,٣٣٠	أوافق بشدة	

أظهرت نتائج جدول (١٤) أن طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود موافقون بشدة على تأثير عامل المسؤولية الأخلاقية في استخدامهم لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي، بمتوسط عام (٤,٨٢ من ٥,٠٠) وانحراف معياري (٠,٣٣٠). وقد جاءت جميع العبارات ضمن الفئة العليا (أوافق بشدة)، حيث حلت عبارة "أحرص على استخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي" في المرتبة الأولى بمتوسط (٤,٨٧)، تلتها عبارتاً "ألتزم بعدم استخدام نماذج الذكاء الاصطناعي التوليدي فيما يخالف النزاهة الأكاديمية" و "أنحقق من صحة المعلومات التي تقدمها لي نماذج الذكاء الاصطناعي التوليدي" في المرتبة الثانية بمتوسط (٤,٨٠). وتشير هذه النتائج إلى وعي عالٍ بالمسؤولية الأخلاقية لدى أفراد العينة، ويعكس التزاماً واضحاً باستخدام هذه النماذج الذكية بما يتوافق مع قيم النزاهة الأكاديمية.

المحور السادس: "ما تأثير عامل الخصوصية المدركة على استخدام طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي؟"

جدول (١٥): إجابات أفراد عينة الدراسة على عبارات محور الخصوصية المدركة مرتبه تنازلياً حسب متوسطات الإجابة

م	العبارات	المتوسط الحسابي	الانحراف المعياري	درجة الاستجابة	الترتيب
١	أشعر بالقلق من احتمال كشف معلوماتي الشخصية عند استخدام نماذج الذكاء الاصطناعي التوليدي.	٣,٨٠	١,٢٠٧	أوافق	٣
٢	أحرص على تجنب إدخال بيانات حساسة عند استخدام نماذج الذكاء الاصطناعي التوليدي.	٤,٤٧	٠,٧٤٣	أوافق بشدة	١
٣	لا أثق بشكل كامل في كيفية تعامل نماذج الذكاء الاصطناعي التوليدي مع معلوماتي الشخصية.	٤,٤٧	١,١٢٥	أوافق بشدة	٢
	المتوسط العام	٤,٢٤	٠,٨٣١	أوافق بشدة	

أظهرت نتائج جدول (١٥) أن طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود موافقون بشدة على تأثير عامل الخصوصية المدركة في استخدامهم الأخلاقي لنماذج الذكاء الاصطناعي التوليدي، بمتوسط عام (٤,٢٤ من ٥,٠٠) وانحراف معياري (٠,٨٣١). وقد تنوعت العبارات بين الموافقة والموافقة الشديدة، حيث جاءت عبارة "أحرص على تجنب إدخال بيانات حساسة عند استخدام نماذج الذكاء الاصطناعي التوليدي" في المرتبة الأولى بمتوسط (٤,٤٧)، تلتها عبارة "لا أثق بشكل كامل في كيفية تعامل نماذج الذكاء الاصطناعي التوليدي مع معلوماتي الشخصية" في المرتبة الثانية بمتوسط (٤,٤٧)، ثم عبارة "أشعر بالقلق من احتمال كشف معلوماتي الشخصية عند استخدام نماذج الذكاء الاصطناعي التوليدي" في المرتبة

الثالثة بمتوسط (٣,٨٠). وتشير هذه النتائج إلى إدراك مرتفع لدى أفراد العينة لمخاطر الخصوصية، وانعكاس ذلك على سلوكهم الحذر في التعامل مع هذه النماذج. المحور السابع: "ما تأثير عامل أمن البيانات على استخدام طلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي؟" جدول (١٦): إجابات أفراد عينة الدراسة على عبارات محور أمن البيانات مرتبة تنازلياً

حسب متوسطات الإجابة

م	العبارات	المتوسط الحسابي	الانحراف المعياري	درجة الاستجابة	الترتيب
١	أشعر بعدم الثقة في كفاءة نماذج الذكاء الاصطناعي التوليدي في تأمين بياناتي الشخصية والأكاديمية الخاصة بي.	٤,٢٠	١,٢٠٧	أوافق	٣
٢	أعتقد أن إجراءات حماية البيانات في نماذج الذكاء الاصطناعي التوليدي مازالت غير كافية وتحتاج إلى تطوير.	٤,٥٣	٠,٦٤	أوافق بشدة	٢
٣	أرى أن استخدام نماذج الذكاء الاصطناعي التوليدي قد ينطوي عليها مخاطر تتعلق بأمن بياناتي الشخصية والأكاديمية الخاصة بي.	٤,٦٠	٠,٦٣٢	أوافق بشدة	١
	المتوسط العام	٤,٤٤	٠,٧٥٢	أوافق بشدة	

أظهرت نتائج جدول (١٦) أن طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود موافقون بشدة على تأثير عامل أمن البيانات في استخدامهم الأخلاقي لنماذج

الذكاء الاصطناعي التوليدي، بمتوسط عام (٤,٤٤ من ٥,٠٠) وانحراف معياري (٠,٧٥٢). وقد جاءت جميع العبارات ضمن الفئة العليا، حيث حلت عبارة "أرى أن استخدام نماذج الذكاء الاصطناعي التوليدي قد ينطوي عليها مخاطر تتعلق بأمن بياناتي الشخصية والأكاديمية الخاصة بي" في المرتبة الأولى بمتوسط (٤,٦٠)، تلتها عبارة "أعتقد أن إجراءات حماية البيانات في نماذج الذكاء الاصطناعي التوليدي مازالت غير كافية وتحتاج إلى تطوير" بمتوسط (٤,٥٣)، ثم عبارة "أشعر بعدم الثقة في كفاءة نماذج الذكاء الاصطناعي التوليدي في تأمين بياناتي الشخصية والأكاديمية الخاصة بي" بمتوسط (٤,٢٠). وتشير هذه النتائج إلى وجود وعي مرتفع لدى أفراد العينة بالمخاطر الأمنية، مما يعزز إدراكهم لأهمية حماية البيانات في تبني الاستخدام الأخلاقي لهذه النماذج الذكية.

المحور الثامن: "ما تأثير عامل الوعي الأخلاقي على استخدام طلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي؟"

جدول (١٧): إجابات أفراد عينة الدراسة على عبارات محور الوعي الأخلاقي مرتبة تنازلياً

حسب متوسطات الإجابة

م	العبارات	المتوسط الحسابي	الانحراف المعياري	درجة الاستجابة	الترتيب
١	أنا على دراية بالمخاطر الأخلاقية المرتبطة باستخدام نماذج الذكاء الاصطناعي التوليدي في التعليم.	٤,٤٠	٠,٧٣٧	أوافق بشدة	٢
٢	أدرك بأن استخدام مخرجات نماذج الذكاء الاصطناعي التوليدي دون توثيق أو مراجعة يُعتبر سرقة أدبية ويعد انتهاكاً للأمانة العلمية.	٤,٦٧	٠,٦١٧	أوافق بشدة	١
٣	اطلعت على أخلاقيات استخدام نماذج الذكاء الاصطناعي التوليدي	٤,١٣	٠,٩٩	أوافق	٣

	في المجال الأكاديمي.			
	المتوسط العام	٤,٤٠	٠,٦٩٢	أو افق بشدة

أظهرت نتائج جدول (١٧) أن طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود موافقون بشدة على تأثير عامل الوعي الأخلاقي في استخدامهم لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي، بمتوسط عام (٤,٤٠ من ٥,٠٠) وانحراف معياري (٠,٦٩٢). وقد جاءت معظم العبارات ضمن الفئة العليا، حيث حلت عبارة "أدرك بأن استخدام مخرجات نماذج الذكاء الاصطناعي التوليدي دون توثيق أو مراجعة يُعتبر سرقة أدبية ويعد انتهاكاً للأمانة العلمية" في المرتبة الأولى بمتوسط (٤,٦٧)، تلتها عبارة "أنا على دراية بالمخاطر الأخلاقية المرتبطة باستخدام نماذج الذكاء الاصطناعي التوليدي في التعليم" بمتوسط (٤,٤٠)، ثم عبارة "اطلعت على أخلاقيات استخدام نماذج الذكاء الاصطناعي التوليدي في المجال الأكاديمي" في المرتبة الثالثة بمتوسط (٤,١٣). وتشير هذه النتائج إلى ارتفاع مستوى الوعي الأخلاقي لدى أفراد العينة، بما يعزز تبنيهم المسؤول لهذه النماذج في بيئتهم الأكاديمية.

المحور التاسع: "ما تأثير عامل النية السلوكية على استخدام طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي؟"

جدول (١٨): إجابات أفراد عينة الدراسة على عبارات محور النية السلوكية مرتبه تنازلياً

حسب متوسطات الإجابة

م	العبارات	المتوسط الحسابي	الانحراف المعياري	درجة الاستجابة	الترتيب
١	أنوي الاستمرار في استخدام	٤,٠٧	٠,٧٠٤	أوافق	٣

				نماذج الذكاء الاصطناعي التوليدي لدعم عملياتي الأكاديمية.	
٢	أوافق بشدة	٠,٧٢٤	٤,٣٣	أنوي استخدام نماذج الذكاء الاصطناعي التوليدي بطريقة تتوافق مع أخلاقيات البحث والتعليم.	٢
١	أوافق بشدة	٠,٦٤	٤,٤٧	أخطط لاستخدام نماذج الذكاء الاصطناعي التوليدي في أعمالي الأكاديمية مستقبلاً بشكل مسؤول.	٣
	أوافق بشدة	٠,٦٤١	٤,٢٩	المتوسط العام	

أظهرت نتائج جدول (١٨) أن طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود موافقون بشدة على تأثير عامل النية السلوكية في استخدامهم الأخلاقي لنماذج الذكاء الاصطناعي التوليدي، بمتوسط عام (٤,٢٩ من ٥,٠٠) وانحراف معياري (٠,٦٤١). وقد توزعت استجاباتهم بين الموافقة والموافقة الشديدة، حيث حلت عبارة "أخطط لاستخدام نماذج الذكاء الاصطناعي التوليدي في أعمالي الأكاديمية مستقبلاً بشكل مسؤول" في المرتبة الأولى بمتوسط (٤,٤٧)، تلتها عبارة "أنوي استخدام نماذج الذكاء الاصطناعي التوليدي بطريقة تتوافق مع أخلاقيات البحث والتعليم" بمتوسط (٤,٣٣)، ثم عبارة "أنوي الاستمرار في استخدام نماذج الذكاء الاصطناعي التوليدي لدعم عملياتي الأكاديمية" في المرتبة الثالثة بمتوسط (٤,٠٧). وتشير هذه النتائج إلى أن لدى أفراد العينة توجه إيجابي قوي للاستمرار في الاستخدام المسؤول والمتوافق مع القيم الأخلاقية للنماذج الذكية.

المحور العاشر: "ما تأثير عامل الاستخدام الأخلاقي على استخدام طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي؟"
جدول (١٩): إجابات أفراد عينة الدراسة على عبارات محور الاستخدام الأخلاقي مرتبه تنازلياً حسب متوسطات الإجابة

م	العبارات	المتوسط الحسابي	الانحراف المعياري	درجة الاستجابة	الترتيب
١	أستخدم نماذج الذكاء الاصطناعي التوليدي بطريقة تحترم المبادئ الأخلاقية والمعايير الأكاديمية.	٤,٦٠	٠,٥٠٧	أوافق بشدة	١
٢	أتجنب الاعتماد الكامل على نماذج الذكاء الاصطناعي التوليدي في أداء مهامي الأكاديمية.	٤,٣٣	١,٠٤٧	أوافق بشدة	٣
٣	أحرص على توثيق أي محتوى تم إنشاؤه باستخدام نماذج الذكاء الاصطناعي التوليدي عند استخدامي له في أبحاثي أو دراستي.	٤,٤٧	٠,٦٤	أوافق بشدة	٢
	المتوسط العام	٤,٤٧	٠,٥٦١	أوافق بشدة	

أظهرت نتائج جدول (١٩) أن طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود موافقون بشدة على تأثير عامل الاستخدام الأخلاقي في تعاملهم مع نماذج الذكاء الاصطناعي التوليدي، بمتوسط عام (٤,٤٧ من ٥,٠٠) وانحراف معياري (٠,٥٦١). وقد جاءت جميع العبارات ضمن الفئة العليا، حيث حلت عبارة "أستخدم نماذج الذكاء الاصطناعي التوليدي بطريقة تحترم المبادئ الأخلاقية والمعايير الأكاديمية" في المرتبة الأولى بمتوسط (٤,٦٠)، تلتها عبارة "أحرص على توثيق أي محتوى تم إنشاؤه باستخدام نماذج الذكاء

الاصطناعي التوليدي عند استخدامي له في أبحاثي أو دراساتي "في المرتبة الثانية بمتوسط (٤,٤٧)، ثم عبارة "أتجنب الاعتماد الكامل على نماذج الذكاء الاصطناعي التوليدي في أداء مهامي الأكاديمية" في المرتبة الثالثة بمتوسط (٤,٣٣). وتشير هذه النتائج إلى التزام عالٍ لدى أفراد العينة بالممارسات الأخلاقية، بما يعزز الاستخدام المسؤول لهذه النماذج الذكية في بيئتهم الأكاديمية.

ثانياً: النتائج المتعلقة بأسئلة الدراسة

السؤال الأول: ما تأثير العوامل التقنية (الأداء المتوقع، التأثير الاجتماعي، الظروف الميسرة، العادة) على النية السلوكية لطلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي.

أولاً- الانحدار البسيط

جدول (٢٠): نتائج تحليل التباين للانحدار (Analysis of variance) للتعرف على تأثير العوامل التقنية ككل على النية السلوكية لطلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي

المصدر	قيمة R ² معامل التحديد	مجموع المربعات	درجات الحرية	متوسط المربعات	قيمة (ف) المحسوبة	مستوى دلالة (ف)
الانحدار	٠,٣١٥	١,٨١٢	١	١,٨١٢	٥,٩٨٣	٠,٠٢٩ *
الخطأ		٣,٩٣٦	١٣	٠,٣٠٣		
المجموع		٥,٧٤٨	١٤			

* ذات دلالة إحصائية عند مستوى $(\alpha \leq 0,05)$.

يتضح من الجدول (٢٠) أن معامل التحديد (R^2) بلغ (٠,٣١٥)، أي أن العوامل التقنية مجتمعة تفسر ما نسبته (٣١,٥٪) من التباين في النية السلوكية لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي. كما أظهرت النتائج أن النموذج دال إحصائياً عند مستوى (٠,٠٥)، حيث بلغت قيمة (ف) المحسوبة (٥,٩٨٣) بدلالة إحصائية (٠,٠٢٩). وتشير هذه النتيجة إلى وجود تأثير معنوي للعوامل التقنية مجتمعة على النية السلوكية، وهو ما

يعكس أن تعزيز الأداء المتوقع، وتوفير الدعم التقني والاجتماعي يسهم في زيادة توجه طلبة الدكتوراه نحو الاستخدام الأخلاقي لهذه النماذج الذكية.

وتتفق هذه النتيجة مع ما توصلت إليه دراسة (Alzahrani, 2025)، والتي أثبتت أن الأداء المتوقع، والظروف الميسرة لهما تأثير إيجابي على النية السلوكية، وكذلك مع دراسة (Sobaih et al., 2024) التي وجدت أن الأداء المتوقع، والتأثير الاجتماعي يسهمان في تعزيز نية استخدام ChatGPT بين الطلبة. وهذا يعزز أهمية دمج العوامل التقنية في تفسير سلوكيات الاستخدام الأخلاقي، وبالأخص في بيئات التعليم العالي التي تسعى للإفادة من نماذج الذكاء الاصطناعي التوليدي Gen AI.

جدول (٢١): نتائج تحليل الانحدار البسيط للتعرف على تأثير العوامل التقنية ككل على النية السلوكية لطلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي

المتغيرات المستقلة	B	الخطأ المعياري	Beta	قيمة (ت)	الدلالة الإحصائية
الثابت	١,١٣٤	١,٢٩٨		٠,٨٧٤	٠,٣٩٨
العوامل التقنية ككل	٠,٨٠٨	٠,٣٣	٠,٥٦١	٢,٤٤٦	٠,٠٢٩

يتضح من الجدول (٢١) أن قيمة الثابت (١,١٣٤) لم تكن ذات دلالة إحصائية (Sig = 0.398 > 0.05) مما يعني أن وجوده لا يضيف تفسير للنية السلوكية عند غياب تأثير العوامل التقنية. في المقابل، ظهر أن العوامل التقنية ككل تؤثر بشكل دال إحصائياً على النية السلوكية، حيث بلغت قيمة معامل بيتا ($\beta = 0.561$) عند مستوى دلالة $\text{Sig} = 0.029 < 0.05$). وهذا يدل على أن ارتفاع مستوى العوامل التقنية بما تشمله من (أداء متوقع، وتأثير اجتماعي، وظروف ميسرة، وعادة) يعزز من النية السلوكية لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي لدى طلبة الدكتوراه.

وتتفق هذه النتيجة مع دراسة (Alzahrani, 2025) التي أكدت على أن العوامل التقنية لها دور أساسي في نية استخدام النماذج الذكية، كما تتفق مع دراسة (Sobaih et al.,

(2024) في أثر الأداء المتوقع، والتأثير الاجتماعي. بينما تختلف النتائج عن بعض الدراسات مثل (Bazelaï et al., 2024) التي لم تثبت وجود تأثير للعوامل التقنية ككل على الاستخدام الفعلي، مما يشير إلى أن تأثير هذه العوامل قد يختلف باختلاف البيئة الأكاديمية والسياق الثقافي.

ثانياً: الانحدار المتعدد

جدول (٢٢): نتائج تحليل التباين للانحدار (Analysis of variance) للتعرف على تأثير العوامل التقنية (الأداء المتوقع، التأثير الاجتماعي، الظروف الميسرة، العادة) على النية السلوكية لطلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي

المصدر	قيمة R^2 معامل التحديد	مجموع المربعات	درجات الحرية	متوسط المربعات	قيمة (ف) المحسوبة	مستوى دلالة (ف)
الانحدار	٠,٣٤٣	١,٩٧	٤	٠,٤٩٢	١,٣٠٣	٠,٣٣٣
الخطأ		٣,٧٧٩	١٠	٠,٣٧٨		
المجموع		٥,٧٤٨	١٤			

* ذات دلالة إحصائية عند مستوى $(\alpha \leq 0,05)$.

يوضح جدول (٢٢) أن قيمة معامل التحديد (R^2) بلغت (٠,٣٤٣)، أي أن العوامل التقنية المتمثلة في الأداء المتوقع، التأثير الاجتماعي، الظروف الميسرة، والعادة تفسر ما نسبته (٣٤,٣٪) من التباين في النية السلوكية لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي. إلا أن النموذج لم يكن دالاً إحصائياً، حيث بلغت قيمة (ف) (١,٣٠٣) عند مستوى دلالة (٠,٣٣٣) (> 0.05)، مما يدل على أن هذه العوامل مجتمعة لم تُظهر تأثيراً معنوياً على النية السلوكية لدى طلبة الدكتوراة.

وتباين هذه النتيجة مع ما توصلت إليه دراسة (Elshaer et al., 2024; Sobaih et al., 2024) من أن الأداء المتوقع، والتأثير الاجتماعي كانا من أبرز محددات النية السلوكية لاستخدام ChatGPT بين الطلبة. كما تختلف عن دراسة (Khlaif et al., 2024) التي أكدت على أن الأداء

المتوقع، والعادة لهما تأثير إيجابي على النية السلوكية. وتتقارب هذه النتيجة مع ما توصلت إليه دراسة (Yilmaz et al., 2023) في الجامعات التركية، والتي بينت أن بعض العوامل التقنية مثل الظروف الميسرة، والتأثير الاجتماعي لم يكن لها تأثير دال على النية السلوكية. ويشير ذلك إلى أن قوة تأثير هذه العوامل قد تختلف باختلاف السياق الأكاديمي والثقافي.

جدول (٢٣): نتائج تحليل الانحدار المتعدد للتعرف على تأثير العوامل التقنية (الأداء المتوقع، التأثير الاجتماعي، الظروف الميسرة، العادة) على النية السلوكية لطلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي

المتغيرات المستقلة	B	الخطأ المعياري	Beta	قيمة (ت)	الدلالة الإحصائية
الثابت	١,١٧١	١,٩٠٧		٠,٦١٤	٠,٥٥٣
الأداء المتوقع	٠,١٥٣	٠,٣٣٩	٠,١٣	٠,٤٥١	٠,٦٦٢
التأثير الاجتماعي	٠,٠٧٦	٠,٢٦٦	٠,٠٧٥	٠,٢٨٧	٠,٧٨
الظروف الميسرة	٠,٣٤٣	٠,٣١٩	٠,٢٨	١,٠٧٧	٠,٣٠٧
العادة	٠,٢٢٩	٠,١٦٩	٠,٣٩٢	١,٣٥٩	٠,٢٠٤

* ذات دلالة إحصائية عند مستوى $(\alpha \leq 0,05)$.

يتضح من الجدول (٢٣) إلى أن الثابت لم يكن ذا دلالة إحصائية ($\text{Sig} = 0.553 > 0.05$)، مما يعني أن قيمته لا تضيف تفسيراً جوهرياً للنية السلوكية. كما أظهرت معاملات الانحدار أن جميع العوامل التقنية (الأداء المتوقع، والتأثير الاجتماعي، والظروف الميسرة، والعادة) لم تكن دالة إحصائية، إذ تجاوزت قيم الدلالة لجميعها مستوى $(0,05)$. وهذا يشير إلى أن أي عامل منفصل من هذه العوامل لا يفسر النية السلوكية بشكل معنوي لدى طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود.

وتتعارض هذه النتيجة مع ما توصلت إليه دراسات مثل (Elshaer et al., 2024; Sobaih et al., 2024) والتي أكدت أن عامل الأداء المتوقع، وعامل التأثير الاجتماعي من العوامل المؤثرة على

النية السلوكية. كما تختلف مع دراسة (Khlaif et al., 2024) والتي توصلت إلى أن العادة لها تأثير إيجابي على نية الاستخدام. وتتفق هذه النتيجة مع ما أشار إليه (Bazelais et al., 2024; Yilmaz et al., 2023) من أن بعض العوامل التقنية مثل الظروف الميسرة، والتأثير الاجتماعي لم يكن لها أثر دال إحصائياً على النية السلوكية، وهو ما يعكس تباين أثر هذه العوامل تبعاً لاختلاف السياق الأكاديمي والثقافي.

السؤال الثاني: ما تأثير العوامل الأخلاقية (المسؤولية الأخلاقية، الخصوصية المدركة، أمن البيانات، الوعي الأخلاقي) على النية السلوكية لطلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي أولاً: الانحدار البسيط:

جدول (٢٤): نتائج تحليل التباين للانحدار (Analysis of variance) للتعرف على تأثير العوامل الأخلاقية ككل على النية السلوكية لطلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي

المصدر	قيمة R^2 معامل التحديد	مجموع المربعات	درجات الحرية	متوسط المربعات	قيمة (ف) المحسوبة	مستوى دلالة (ف)
الانحدار	٠,١٣٣	٠,٧٦٦	١	٠,٧٦٦	١,٩٩٨	٠,١٨١
الخطأ		٤,٩٨٢	١٣	٠,٣٨٣		
المجموع		٥,٧٤٨	١٤			

* ذات دلالة إحصائية عند مستوى $(\alpha \leq 0,05)$.

يتضح من الجدول (٢٤) أن معامل التحديد (R^2) بلغ (٠,١٣٣)، أي أن العوامل الأخلاقية ككل تفسر ما نسبته (١٣,٣٪) فقط من التباين في النية السلوكية. إلا أن النموذج لم يكن دالاً إحصائياً، حيث بلغت قيمة (ف) (١,٩٩٨) عند مستوى دلالة (٠,١٨١) (> 0.05)، مما يعني أن العوامل الأخلاقية مجتمعة لم تُظهر تأثيراً معنوياً على النية السلوكية لطلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود.

وتختلف هذه النتيجة مع دراسة (Alzahrani, 2025) التي أثبتت أن لمخاطر الخصوصية تأثيراً إيجابياً على نية استخدام أعضاء هيئة التدريس لـ ChatGPT، وكذلك مع دراسة (Acosta-Enriquez et al., 2024) التي بينت أن الخصوصية المدركة، والأخلاقيات المدركة تعد من العوامل المؤثرة في النية السلوكية. كما تختلف عن دراسة (Huallpa et al., 2023) التي أشارت إلى أهمية دمج المعايير الأخلاقية المؤسسية في تعزيز نوايا الاستخدام. وفي المقابل، يمكن تفسير هذه النتيجة بخصوصية عينة الدراسة (طلبة الدكتوراه) التي قد ترى أن التحديات الأخلاقية وحدها غير كافية لتحديد نية الاستخدام ما لم تقترن بعوامل تقنية أو تنظيمية داعمة.

جدول (٢٥): نتائج تحليل الانحدار البسيط للتعرف على تأثير العوامل الأخلاقية ككل على النية السلوكية لطلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي

المتغيرات المستقلة	B	الخطأ المعياري	Beta	قيمة (ت)	الدلالة الإحصائية
الثابت	٢,٣٧٤	١,٣٦٤		١,٧٤	٠,١٠٥
العوامل الأخلاقية ككل	٠,٤٢٨	٠,٣٠٣	٠,٣٦٥	١,٤١٤	٠,١٨١

* ذات دلالة إحصائية عند مستوى $(\alpha \leq 0,05)$.

يوضح جدول (٢٥) أن قيمة الثابت (٢,٣٧٤) لم تكن ذات دلالة إحصائية (Sig = 0.105 > 0.05). وهذا يشير إلى أن وجود الثابت لا يسهم في تفسير النية السلوكية. كما بينت النتائج أن تأثير العوامل الأخلاقية ككل لم يكن دالاً إحصائياً (Beta = 0.365, Sig = 0.181 > 0.05)، مما يعني أنها لم تفسر النية السلوكية بشكل معنوي لدى طلبة الدكتوراه في قسم علم المعلومات بجامعة الملك سعود.

وتتعارض هذه النتيجة مع ما توصلت إليه دراسة (Acosta-Enriquez et al., 2024) التي أكدت على أن الخصوصية المدركة، والأخلاقيات المدركة تؤثر على النية السلوكية، وكذلك

مع دراسة (Mironova et al., 2024) التي أظهرت أن إدراك الطلبة للجوانب الأخلاقية يلعب دوراً في تحديد مواقفهم من استخدام النماذج الذكية. بينما تتقارب النتيجة مع بعض الدراسات مثل دراسة (Gallent-Torres et al., 2023) والتي أشارت إلى أن المخاوف الأخلاقية وحدها لا تكفي لضبط نية الاستخدام ما لم تصاحبها أطر مؤسسية أو تنظيمية واضحة.

ثانياً: الانحدار المتعدد

جدول (٢٦): نتائج تحليل التباين للانحدار (Analysis of variance) للتعرف على تأثير العوامل الأخلاقية (المسؤولية الأخلاقية، الخصوصية المدركة، أمن البيانات، الوعي الأخلاقي) على النية السلوكية لطلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي

المصدر	قيمة R^2 معامل التحديد	مجموع المربعات	درجات الحرية	متوسط المربعات	قيمة (ف) المحسوبة	مستوى دلالة (ف)
الانحدار	٠,٦٨٧	٣,٩٤٨	٤	٠,٩٨٧	٥,٤٨٣	٠,٠١٣ *
الخطأ		١,٨	١٠	٠,١٨		
المجموع		٥,٧٤٨	١٤			

* ذات دلالة إحصائية عند مستوى $(\alpha \leq ٠,٠٥)$.

يوضح جدول (٢٦) أن معامل التحديد (R^2) بلغ (٠,٦٨٧). وهذا يعني أن المسؤولية الأخلاقية، والخصوصية المدركة، وأمن البيانات، والوعي الأخلاقي تفسر مجتمعة ما نسبته (٦٨,٧٪) من التباين في النية السلوكية لطلبة الدكتوراه. كما أظهر النموذج معنوية إحصائية، حيث بلغت قيمة (ف) (٥,٤٨٣) عند مستوى دلالة (٠,٠١٣) < 0.05 ويؤكد ذلك صلاحية النموذج وقدرته على تفسير تأثير العوامل الأخلاقية بشكل ملحوظ.

وتتفق هذه النتيجة مع ما أشار إليه (Acosta-Enriquez et al., 2024; Alzahrani, 2025) في أن دمج البعد الأخلاقي مثل الخصوصية المدركة، والمخاطر الأخلاقية يعزز من قدرة النماذج على فهم نية الاستخدام. كما تدعمها دراسة (Huallpa et al., 2023) التي أوصت

بضرورة دمج المعايير الأخلاقية المؤسسية للحد من المخاطر، وهو ما ينعكس على استعداد الطلبة لاستخدام النماذج الذكية بشكل مسؤول.

جدول (٢٧): نتائج تحليل الانحدار المتعدد للتعرف على تأثير العوامل الأخلاقية (المسؤولية الأخلاقية، الخصوصية المدركة، أمن البيانات، الوعي الأخلاقي) على النية السلوكية لطلبة الدكتوراة في قسم علم المعلومات بجامعة الملك سعود لاستخدام نماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي

المتغيرات المستقلة	B	الخطأ المعياري	Beta	قيمة (ت)	الدلالة الإحصائية
الثابت	٠,٠٣٢	١,٨٢٤		٠,٠١٧	٠,٩٨٦
المسؤولية الأخلاقية	٠,٧٠٣	٠,٤٤٩	٠,٣٦٢	١,٥٦٦	٠,١٤٨
الخصوصية المدركة	٠,٦٢٨	٠,٣٠٩	٠,٨١٤	٢,٠٣١	٠,٠٧٠
أمن البيانات	١,٠١٧-	٠,٣٩٤	١,١٩٥-	٢,٥٨-	٠,٠٢٧*
الوعي الأخلاقي	٠,٦١٩	٠,٢٧٤	٠,٦٦٨	٢,٢٥٧	٠,٠٤٨*

* ذات دلالة إحصائية عند مستوى $(\alpha \leq 0.05)$.

يتضح من الجدول (٢٧) أن الثابت لم يكن دالاً إحصائياً ($\text{Sig} = 0.986 > 0.05$)، مما يعني أن النموذج يعتمد على المتغيرات المستقلة لتفسير النية السلوكية. وأظهرت النتائج أن المسؤولية الأخلاقية ($\beta = 0.362, \text{Sig} = 0.148$) والخصوصية المدركة ($\beta = 0.814, \text{Sig} = 0.070$) لم تكونا ذات دلالة إحصائية. في المقابل، ظهر لعامل أمن البيانات تأثير سلبي دال إحصائياً ($\beta = -1.195, \text{Sig} = 0.027$) مما يشير إلى أن تزايد المخاوف المرتبطة بأمن البيانات يقلل من النية السلوكية، بينما كان للوعي الأخلاقي تأثير إيجابي دال إحصائياً ($\beta = 0.668, \text{Sig} = 0.048$) وهذا يعني أن ارتفاع الوعي بالمخاطر والقضايا الأخلاقية يعزز من النية السلوكية لاستخدام هذه النماذج بشكل مسؤول.

وتتفق هذه النتائج مع دراسة (Mironova et al., 2024) التي أظهرت أن زيادة الوعي الأخلاقي يعزز من تبني الاستخدام المسؤول. بينما يختلف ذلك جزئياً مع دراسة (Alzaharani, 2025) التي أبرزت دور الخصوصية المدركة كمحفز أساسي، وهو ما لم يظهر هنا بشكل دال إحصائياً، وهذا يمكن تفسيره باختلاف طبيعة العينة وسياق الدراسة.

ثالثاً. مناقشة نتائج الدراسة

أظهرت نتائج دراسة العوامل المؤثرة على استخدام طلبة الدراسات العليا (مرحلة الدكتوراة) في قسم علم المعلومات بجامعة الملك سعود لنماذج الذكاء الاصطناعي التوليدي بشكل أخلاقي ما يلي:

أن عامل الأداء المتوقع (PE) هو الأكثر تأثيراً في النية السلوكية للطلبة لاستخدام نماذج الذكاء الاصطناعي التوليدي Gen AI، وهو ما يتفق مع ما توصلت إليه الدراسات السابقة في دول مختلفة مثل الصين، وأوروبا، وجنوب شرق آسيا (Bouteraa et al., 2024; Sobaih et al., 2024; Tian et al., 2024)، حيث تمثلت الفائدة التعليمية من هذه النماذج الذكية كمحرك رئيسي لاستخدامها في العمليات الأكاديمية. أما عامل التأثير الاجتماعي (SI) فلم يظهر دلالة إحصائية واضحة، وهذا يختلف مع ما أظهرته بعض الدراسات مثل (Budhathoki et al., 2024)، والتي أكدت على دور الأقران والتأثير الاجتماعي في تشكيل سلوك الاستخدام. وقد يعكس هذا التباين خصوصية البيئة الأكاديمية في جامعة الملك سعود، حيث يظهر أن قرارات الطلبة مرتبطة بالقيمة التعليمية

الفردية والأداء الشخصي أكثر من ارتباطها بتأثير الزملاء أو البيئة الاجتماعية. كما أن غياب الدلالة الإحصائية لعامل الظروف الميسرة (FC) قد يُفسّر بتطور البنية التحتية الرقمية في الجامعات السعودية، حيث لا يعتبر هذا العامل محدد رئيسي للاستخدام. وضعف تأثير عامل العادة (H) يمكن تفسيره بحدّة تعامل الطلبة مع نماذج Gen AI، حيث لم يتشكل لدى الطلبة نمط سلوكي روتيني في الاعتماد على هذه النماذج الذكية، وهذه النتيجة تختلف مع ما ظهر في بيانات أكثر تقدم في دمج هذه التقنيات (Bazelaïs et al., 2024). وتوضح النتائج أن الطلبة ينظرون إلى GenAI كأداة مساعدة لتحسين الإنتاجية، وإنجاز المهام الأكاديمية أكثر من كونه بديلاً للتفاعل البشري أو مصدراً للمعرفة الأخلاقية، وهو ما يتفق مع ما أشار إليه (Huallpa et al., 2023) من أن الفوائد المدركة غالباً تتجاوز المخاوف النظرية. كما أظهرت النتائج محدودية تأثير العوامل الأخلاقية، وهذا يكشف عن وجود فجوة بين القيمة والفعل (Value–Action Gap) (Hagendorff, 2022) حيث أن الطلبة يملكون الوعي الكافي بالقضايا الأخلاقية مثل الخصوصية، وأمن البيانات، إلا أنهم يفضلون الفائدة الأكاديمية على هذه الاعتبارات، وقد يكون ذلك نتيجة ضغوط لتحقيق الإنجاز الأكاديمي. وأوضحت اليونسكو (UNESCO, 2023) أن غياب الأطر والسياسات المؤسسية الواضحة يزيد من التحديات الأخلاقية لنماذج GAI، مع ضرورة وضع أطر تنظيمية لحماية الخصوصية، وضمان الاستخدام الأخلاقي لهذه النماذج الذكية في التعليم العالي.

التوصيات

بناءً على النتائج السابقة فإن الدراسة توصي بالآتي:

١. صياغة سياسات تنظيمية واضحة لتحديد ضوابط الاستخدام المسموح والمقيد لنماذج الذكاء الاصطناعي التوليدي في كليات الجامعة، وإبلاغ الطلبة من خلال الأدلة الإرشادية، وإتاحة قنوات للتواصل.
٢. إنشاء لجان دائمة لأخلاقيات الذكاء الاصطناعي في الكليات، تتركز مسؤوليتها في مراجعة تنفيذ السياسات، وتقديم الاستشارات، ورفع تقارير دورية لمجالس الجامعات.
٣. الاستثمار في تعزيز البنية التحتية للأمن السيبراني وحماية البيانات الشخصية من خلال تنفيذ حلول تقنية متقدمة. كما توصي الدراسة بإجراء اختبارات دورية للتحقق من

- مستوى الأمان. ويُفضل إلزام الطلبة باستخدام حسابات جامعية مؤمنة عند التعامل مع النماذج الذكية.
٤. تنظيم برامج توعوية تدريبية، وتكون إلزامية لطلبة الدراسات العليا حول الاستخدام الآمن والمسؤول لنماذج الذكاء الاصطناعي التوليدي، مع التركيز على قضايا النزاهة الأكاديمية، وحماية البيانات الشخصية، والتحيزات الخوارزمية.
٥. إدراج مقررات متخصصة ضمن برامج الدراسات العليا تركز على موضوعات مثل أخلاقيات الذكاء الاصطناعي، والخصوصية الرقمية، لتعزيز الوعي، وبناء المهارات النقدية للطلبة والباحثين.
٦. دمج أدوات الذكاء الاصطناعي التوليدي في العملية التعليمية، واعتبارها أدوات مساعدة بإشراف الأساتذة، لتحقيق التوازن بين الاستفادة التقنية، ودعم الدور التعليمي والبحثي للطلبة.
٧. اقتراح آلية لربط الجامعات بمؤشرات أداء رئيسية (KPIs) لقياس مدى التزامها بالسياسات الوطنية المتعلقة بالذكاء الاصطناعي التوليدي.
٨. دمج الإطار الأخلاقي للذكاء الاصطناعي التوليدي الصادر عن الهيئة السعودية للبيانات والذكاء الاصطناعي (سدايا) ضمن المقررات الجامعية. وذلك من خلال تخصيص وحدات أو مقررات دراسية لمعالجة قضايا أخلاقيات الذكاء الاصطناعي بشكل منهجي. وربط الممارسات الأكاديمية بالسياسات الوطنية المعتمدة.
٩. تطوير إطار عملي يقارن يمكن الجامعات العربية من الاستفادة من أفضل الممارسات العالمية، وبناء نموذج نقدي يعكس خصوصية السياقات المحلية.
١٠. إجراء دراسات طولية Longitudinal Studies لمتابعة تطور وعي الطلبة مع مرور الوقت في الاستخدام الأخلاقي للنماذج الذكية، وتزايد اعتمادهم عليها.
١١. إجراء دراسات مقارنة بين تخصصات أكاديمية متعددة مثل كليات العلوم الإنسانية والاجتماعية، والطب، والهندسة لفهم اختلافات الاستخدام الأخلاقي للنماذج الذكية حسب طبيعة المجالات العلمية.
١٢. إجراء دراسات نوعية Qualitative Research باستخدام المقابلات، ومجموعات التركيز، لفهم سلوكيات الطلبة لاستخدام نماذج الذكاء الاصطناعي التوليدي.

ختاماً، تُسهم هذه الدراسة في إثراء النقاش الأكاديمي حول الاستخدام الأخلاقي لنماذج الذكاء الاصطناعي التوليدي GAI في التعليم العالي، وذلك من خلال تقديم إطار يدمج بين العوامل التقنية، والعوامل الأخلاقية ضمن امتداد النظرية الموحدة لقبول واستخدام التكنولوجيا UTAUT2. حيث أظهرت النتائج أن عامل الأداء المتوقع هو العامل الأكثر تأثيراً على نية الطلبة السلوكية، في حين لم تُظهر عوامل مثل التأثير الاجتماعي، والظروف الميسرة، والعادة تأثيراً له دلالة إحصائية في جامعة الملك سعود. أما العوامل الأخلاقية، فلم تُحقق جميعها دلالة إحصائية، وهو ما يكشف عن وجود فجوة (value-action gap) بين وعي الطلبة بالمخاطر الأخلاقية مثل الخصوصية، وأمن البيانات، وبين نيّتهم السلوكية لاستخدام هذه النماذج الذكية بشكل أخلاقي، ويُحتمل أن يكون ذلك نتيجةً لتفضيل الطلبة الحصول على الفوائد الأكاديمية بشكل مباشر على الاعتبارات الأخلاقية طويلة المدى.

وتوضح هذه النتائج خصوصية الثقافة الأكاديمية المحلية، حيث يتخذ الطلبة قرار استخدام النماذج الذكية بشكل فردي أكثر من استجابتهم للتأثيرات الاجتماعية أو الاعتبارات الأخلاقية. كما أنها تتقاطع مع المبادئ التي أوصت بها اليونسكو (UNESCO, 2023) واللائحة الأوروبية للذكاء الاصطناعي (European Commission, 2024) والمتعلقة بالنزاهة الأكاديمية، وحماية الحقوق. وتؤكد نتائج الدراسة على ضرورة تكييف المبادئ العالمية لأخلاقيات الذكاء الاصطناعي التوليدي لضمان ملاءمتها للسياق المحلي، وتعزيز إمكانية تطبيقها بفاعلية في التعليم العالي لتحقيق الاستخدام الأخلاقي والمسؤول لنماذج الذكاء الاصطناعي التوليدي GAI.

المراجع

جامعة الملك سعود. (٢٠٢٥). مركز الدراسات المتقدمة في الذكاء الاصطناعي (ذكاء) تم

الاسترجاع في ١١ يناير ٢٠٢٥ ، من: [الصفحة الرئيسية | مركز الدراسات المتقدمة في](#)

[الذكاء الاصطناعي \(ذكاء\)](#)

الهيئة السعودية للبيانات والذكاء الاصطناعي (سدايا). (2023). الذكاء الاصطناعي

التوليدي في التعليم

<https://sdaia.gov.sa/ar/MediaCenter/KnowledgeCenter/Pages/SDAIAPublicat>

[ions.aspx](#)

Acosta-Enriquez, B. G., Arbulú Ballesteros, M. A., Huamaní Jordan, O., López Roca,

C., & Saavedra Tirado, K. (2024). Analysis of college students' attitudes

toward the use of ChatGPT in their academic activities: effect of intent to use, verification of information and responsible use. *BMC Psychology*, 12(1).

<https://doi.org/10.1186/s40359-024-01764-z>

Ahmed, Z., Shanto, S. S., Rime, M. H. K., Morol, M. K., Fahad, N., Hossen, M. J., &

Abdullah-Al-Jubair, M. (2024). The Generative AI Landscape in Education: Mapping the Terrain of Opportunities, Challenges and Student Perception. *IEEE Access*.

<https://doi.org/10.1109/ACCESS.2024.3461874>

Aljuaid, H. (2024). The Impact of Artificial Intelligence Tools on Academic Writing

Instruction in Higher Education: A Systematic Review. *Arab World English Journal*, 1(1), 26–55. <https://doi.org/10.24093/awej/chatgpt.2>

Alzahrani, -----Amal. (2025).

Understanding ChatGPT adoption in universities: the impact of faculty TPACK and

UTAUT2 Comprendiendo la adopción de ChatGPT en universidades: el impacto del

TPACK y UTAUT2 en los docentes. <https://doi.org/10.5944/ried.28.1.41498>

- Bahadur, S. G. C., Bhandari, P., Gurung, S. K., Srivastava, E., Ojha, D., & Dhungana, B. R. (2024). Examining the role of social influence, learning value and habit on students' intention to use ChatGPT: the moderating effect of information accuracy in the UTAUT2 model. *Cogent Education*, 11(1).
<https://doi.org/10.1080/2331186X.2024.2403287>
- Baidoo-Anu, D., & Owusu Ansah, L. (2023). Education in the Era of Generative Artificial Intelligence (AI): Understanding the Potential Benefits of ChatGPT in Promoting Teaching and Learning. In *Journal of AI* (Vol. 52, Issue 7).
- Balalle, H., & Pannilage, S. (2025). Reassessing academic integrity in the age of AI: A systematic literature review on AI and academic integrity. In *Social Sciences and Humanities Open* (Vol. 11). Elsevier Ltd.
<https://doi.org/10.1016/j.ssaho.2025.101299>
- Bankins, S., & Formosa, P. (2023). The Ethical Implications of Artificial Intelligence (AI) For Meaningful Work. *Journal of Business Ethics*, 185(4), 725–740.
<https://doi.org/10.1007/s10551-023-05339-7>
- Baskara, FX. R. (2023). The Promises and Pitfalls of Using Chat GPT for Self-Determined Learning in Higher Education: An Argumentative Review. *Prosiding Seminar Nasional Fakultas Tarbiyah Dan Ilmu Keguruan IAIM Sinjai*, 2, 95–101.
<https://doi.org/10.47435/sentikjar.v2i0.1825>
- Bazelais, P., Lemay, D. J., & Doleck, T. (2024). User acceptance and adoption dynamics of ChatGPT in educational settings. *Eurasia Journal of Mathematics, Science and Technology Education*, 20(2). <https://doi.org/10.29333/ejmste/14151>
- Biloš, A., & Budimir, B. (2024). Understanding the Adoption Dynamics of ChatGPT among Generation Z: Insights from a Modified UTAUT2 Model. *Journal of Theoretical and*

Applied Electronic Commerce Research, 19(2), 863–879.

<https://doi.org/10.3390/jtaer19020045>

Biton, Y., & Segal, R. (2025). Learning to Craft and Critically Evaluate Prompts: The Role of Generative AI (ChatGPT) in Enhancing Pre-service Mathematics Teachers' TPACK and Problem-Posing Skills. *International Journal of Education in Mathematics, Science and Technology*, 13(1), 202–223. <https://doi.org/10.46328/ijemst.4654>

Bouteraa, M., Bin-Nashwan, S. A., Al-Daihani, M., Dirie, K. A., Benlahcene, A., Sadallah, M., Zaki, H. O., Lada, S., Ansar, R., Fook, L. M., & Chekima, B. (2024). Understanding the diffusion of AI-generative (ChatGPT) in higher education: Does students' integrity matter? *Computers in Human Behavior Reports*, 14. <https://doi.org/10.1016/j.chbr.2024.100402>

Brown, T., Isbel, S., Logan, A., & Etherington, J. (2020). Predictors of academic integrity in undergraduate and graduate-entry masters occupational therapy students. *Hong Kong Journal of Occupational Therapy*, 33(2), 42–54. <https://doi.org/10.1177/1569186120968035>

Budhathoki, T., Zirar, A., Njoya, E. T., & Timsina, A. (2024). ChatGPT adoption and anxiety: a cross-country analysis utilising the unified theory of acceptance and use of technology (UTAUT). *Studies in Higher Education*, 49(5), 831–846. <https://doi.org/10.1080/03075079.2024.2333937>

Chan, C. K. Y., & Hu, W. (2023a). Students' voices on generative AI: perceptions, benefits, and challenges in higher education. *International Journal of Educational Technology in Higher Education*, 20(1). <https://doi.org/10.1186/s41239-023-00411-8>

Chan, C. K. Y., & Hu, W. (2023b). Students' voices on generative AI: perceptions, benefits, and challenges in higher education. *International Journal of Educational*

Technology in Higher Education, 20(1). <https://doi.org/10.1186/s41239-023-00411-8>

de Fine Licht, K. (2024). Generative Artificial Intelligence in Higher Education: Why the “Banning Approach” to Student use is Sometimes Morally Justified. *Philosophy and Technology*, 37(3). <https://doi.org/10.1007/s13347-024-00799-9>

Dempere, J., Modugu, K., Hesham, A., & Ramasamy, L. K. (2023). The impact of ChatGPT on higher education. *Frontiers in Education*, 8. <https://doi.org/10.3389/feduc.2023.1206936>

Elshaer, I. A., Hasanein, A. M., & Sobaih, A. E. E. (2024). The Moderating Effects of Gender and Study Discipline in the Relationship between University Students’ Acceptance and Use of ChatGPT. *European Journal of Investigation in Health, Psychology and Education*, 14(7), 1981–1995. <https://doi.org/10.3390/ejihpe14070132>

European Commission. (2024). *AI Act*. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>

Febriyani, W., Kusumasari, T., & Lubis, M. (2023). *Data Security: A Systematic Literature Review and Critical Analysis*.

Feuerriegel, S., Hartmann, J., Janiesch, C., & Zschech, P. (2024). Generative AI. *Business and Information Systems Engineering*, 66(1), 111–126. <https://doi.org/10.1007/s12599-023-00834-7>

Firaina, R., & Sulisworo, D. (2023a). Exploring the Usage of ChatGPT in Higher Education: Frequency and Impact on Productivity. *Buletin Edukasi Indonesia*, 2(01), 39–46. <https://doi.org/10.56741/bei.v2i01.310>

- Firaina, R., & Sulisworo, D. (2023b). Exploring the Usage of ChatGPT in Higher Education: Frequency and Impact on Productivity. *Buletin Edukasi Indonesia*, 2(01), 39–46.
<https://doi.org/10.56741/bei.v2i01.310>
- Gallent-Torres, C., Zapata-González, A., & Ortego-Hernando, J. L. (2023). The impact of Generative Artificial Intelligence in higher education: a focus on ethics and academic integrity. *RELIEVE - Revista Electronica de Investigacion y Evaluacion Educativa*, 29(2). <https://doi.org/10.30827/RELIEVE.V29I2.29134>
- Ganesamoorthy, M., & Selvakamal, P. (2024). *EMERGING TECHNOLOGIES AND TRENDS IN LIBRARY: A STUDY*.
- GDPR.EU. (n.d.). *What is GDPR, the EU's new data protection law?*
<https://Gdpr.Eu/What-Is-Gdpr/>.
- Hagendorff, T. (2022). A Virtue-Based Framework to Support Putting AI Ethics into Practice. *Philosophy and Technology*, 35(3). <https://doi.org/10.1007/s13347-022-00553-z>
- Hanna, R., & Kazim, E. (2021). Philosophical foundations for digital ethics and AI Ethics: a dignitarian approach. *AI and Ethics*, 1(4), 405–423.
<https://doi.org/10.1007/s43681-021-00040-9>
- Hasal, M., Nowaková, J., Ahmed Saghair, K., Abdulla, H., Snášel, V., & Ogiela, L. (2021). Chatbots: Security, privacy, data protection, and social aspects. *Concurrency and Computation: Practice and Experience*, 33(19). <https://doi.org/10.1002/cpe.6426>
- Helberger, N., & Diakopoulos, N. (2023). ChatGPT and the AI Act. *Internet Policy Review*, 12(1). <https://doi.org/10.14763/2023.1.1682>
- Herman, J. (2023, March 27). *Top UK Universities Ban Chat-GPT*.
<https://www.Redbrick.Me/Top-Uk-Universities-Ban-Chat-Gpt/>.

- Huallpa, J. J., Flores Arocutipa, J. P., Diaz Panduro, W., Huete, L. C., Antonio, F., Limo, F., Herrera, E. E., Arturo, R., Callacna, A., Andre, V., Flores, A., Ángel, M., Romero, M., Merino Quispe, I., & Hernández Hernández, A. (2023). Exploring the ethical considerations of using Chat GPT in university education. *Original Research*, 11(4), 105–115. <https://orcid.org/0000-0003-4067-2816>
- Kant, I. (2022). *Kant's Moral Philosophy*. <https://plato.stanford.edu/entries/kant-moral/>
- Khlaif, Z. N., Ayyoub, A., Hamamra, B., Bensalem, E., Mitwally, M. A. A., Ayyoub, A., Hattab, M. K., & Shadid, F. (2024). University Teachers' Views on the Adoption and Integration of Generative AI Tools for Student Assessment in Higher Education. *Education Sciences*, 14(10). <https://doi.org/10.3390/educsci14101090>
- Koo, M., & Yang, S.-W. (2025). Likert-Type Scale. *Encyclopedia*, 5(1), 18. <https://doi.org/10.3390/encyclopedia5010018>
- Lau, P. (2021). A Case Study on Research Postgraduate Students' Understanding of Academic Integrity at a Hong Kong University. *Frontiers in Education*, 6. <https://doi.org/10.3389/educ.2021.647626>
- Limna, P., Kraiwani, T., Jangjarat, K., Klayklung, P., & Chocksathaporn, P. (2023). The use of ChatGPT in the digital era: Perspectives on chatbot implementation. *Journal of Applied Learning and Teaching*, 6(1), 64–74. <https://doi.org/10.37074/jalt.2023.6.1.32>
- Liu, M., Ren, Y., Nyagoga, L. M., Stonier, F., Wu, Z., & Yu, L. (2023). Future of education in the era of generative artificial intelligence: Consensus among Chinese scholars on applications of ChatGPT in schools. *Future in Educational Research*, 1(1), 72–101. <https://doi.org/10.1002/fer3.10>
- Luk, C.-Y., Chung, H.-L., Yim, W.-K., & Leung, C.-W. (n.d.). *Regulating Generative AI: Ethical Considerations and Explainability Benchmarks*.

- Malinka, K., Perešini, M., Firc, A., Hujňák, O., & Januš, F. (2023). *On the Educational Impact of ChatGPT: Is Artificial Intelligence Ready to Obtain a University Degree?*
<https://doi.org/10.1145/3587102.3588827>
- Malouf, N. El. (2023). *Diffusion of Innovations*. <https://open.ncl.ac.uk>
- Mcdonald, N., Johri, A., Ali, A., & Hingle, A. (2025). *Generative Artificial Intelligence in Higher Education Generative Artificial Intelligence in Higher Education: Evidence from an Analysis of Institutional Policies and Guidelines*.
- Mcminn, S. (2024). *Generative AI in Higher Education; The ChatGPT Effect*.
- Menon, D., & Shilpa, K. (2023). "Chatting with ChatGPT": Analyzing the factors influencing users' intention to Use the Open AI's ChatGPT using the UTAUT model. *Heliyon*, 9(11). <https://doi.org/10.1016/j.heliyon.2023.e20962>
- Michel-Villarreal, R., Vilalta-Perdomo, E., Salinas-Navarro, D. E., Thierry-Aguilera, R., & Gerardou, F. S. (2023). Challenges and Opportunities of Generative AI for Higher Education as Explained by ChatGPT. *Education Sciences*, 13(9).
<https://doi.org/10.3390/educsci13090856>
- Mironova, J., Riashchenko, V., Kinderis, R., Djakona, V., & Dimitrova, S. I. (2024). Ethical concerns in using of Generative Tools in Higher Education: Cross - Country Study. *Vide. Tehnologija. Resursi - Environment, Technology, Resources*, 2, 444–447.
<https://doi.org/10.17770/etr2024vol2.8097>
- OECD. (2023). *Recommendation of the Council OECD Legal Instruments concerning Guidelines Governing the Protection of Privacy and Transborder Flows of Personal Data*. <http://legalinstruments.oecd.org>
- Palinkas, L. A., Horwitz, S. M., Green, C. A., Wisdom, J. P., Duan, N., & Hoagwood, K. (2015). Purposeful Sampling for Qualitative Data Collection and Analysis in Mixed Method Implementation Research. *Administration and Policy in Mental Health and*

Mental Health Services Research, 42(5), 533–544.

<https://doi.org/10.1007/s10488-013-0528-y>

Parveen, K., Phuc, T. Q. B., Alghamdi, A. A., Hajjej, F., Obidallah, W. J., Alduraywish, Y. A., & Shafiq, M. (2024). Unraveling the dynamics of ChatGPT adoption and utilization through Structural Equation Modeling. *Scientific Reports*, 14(1).

<https://doi.org/10.1038/s41598-024-74406-4>

Qadhi, S. M., Alduais, A., Chaaban, Y., & Khraisheh, M. (2024). Generative AI, Research Ethics, and Higher Education Research: Insights from a Scientometric Analysis. In *Information (Switzerland)* (Vol. 15, Issue 6). Multidisciplinary Digital Publishing Institute (MDPI). <https://doi.org/10.3390/info15060325>

Qammar, A., Wang, H., Ding, J., Naouri, A., Daneshmand, M., & Ning, H. (2023). *Chatbots to ChatGPT in a Cybersecurity Space: Evolution, Vulnerabilities, Attacks, Challenges, and Future Recommendations*. <http://arxiv.org/abs/2306.09255>

Rana, M. M., Siddiquee, M. S., Sakib, M. N., & Ahamed, M. R. (2024). Assessing AI adoption in developing country academia: A trust and privacy-augmented UTAUT framework. *Heliyon*, 10(18). <https://doi.org/10.1016/j.heliyon.2024.e37569>

Ray, P. P. (2023). ChatGPT: A comprehensive review on background, applications, key challenges, bias, ethics, limitations and future scope. In *Internet of Things and Cyber-Physical Systems* (Vol. 3, pp. 121–154). KeAi Communications Co. <https://doi.org/10.1016/j.iotcps.2023.04.003>

Saxena, D., Khandare, S., & Chaudhary, S. (2023). An Overview of ChatGPT: Impact on Academic Learning FMD Transactions on Sustainable Techno Learning An Overview of ChatGPT: Impact on Academic Learning. In *Transactions on Sustainable Techno Learning* (Vol. 1, Issue 1).

<https://www.researchgate.net/publication/370805477>

- Schreur, P. (2019). *Yewno: Transforming Data into Information, Transforming Information into Knowledge*. <http://creativecommons.org/licenses/by/4.0>
- Sebastian, G. (2023). *Privacy and Data Protection in ChatGPT and Other AI Chatbots: Strategies for Securing User Information*.
<https://doi.org/10.13140/RG.2.2.14633.57449>
- Shahriar, S., & Hayawi, K. (2023). Let's Have a Chat! A Conversation with ChatGPT: Technology, Applications, and Limitations. *Artificial Intelligence and Applications*.
<https://doi.org/10.47852/bonviewAIA3202939>
- Sobaih, A. E. E., Elshaer, I. A., & Hasanein, A. M. (2024). Examining Students' Acceptance and Use of ChatGPT in Saudi Arabian Higher Education. *European Journal of Investigation in Health, Psychology and Education*, 14(3), 709–721.
<https://doi.org/10.3390/ejihpe14030047>
- Stanford Encyclopedia of Philosophy. (2022, November 11). *Virtue Ethics*.
<https://Plato.Stanford.Edu/Entries/Ethics-Virtue/>.
<https://plato.stanford.edu/entries/ethics-virtue/>
- Stanford Encyclopedia of Philosophy. (2025, July 31). *The History of Utilitarianism*.
- Tan, L. (2015). *Self-Directed Learning: Learning in the 21st Century*.
<https://www.researchgate.net/publication/285591239>
- Tavakol, M., & Dennick, R. (2011). Making sense of Cronbach's alpha. In *International journal of medical education* (Vol. 2, pp. 53–55).
<https://doi.org/10.5116/ijme.4dfb.8dfd>
- The European Commission's. (2018). *A DEFINITION OF AI: MAIN CAPABILITIES AND SCIENTIFIC DISCIPLINES*. <https://ec.europa.eu/digital-single-market/en/high-level-expert-group-artificial-intelligence>

- Tian, W., Ge, J., Zhao, Y., & Zheng, X. (2024). AI Chatbots in Chinese higher education: adoption, perception, and influence among graduate students—an integrated analysis utilizing UTAUT and ECM models. *Frontiers in Psychology*, 15. <https://doi.org/10.3389/fpsyg.2024.1268549>
- Todd, P. (2016). *Strawson, Moral Responsibility, and the "Order of Explanation": An Intervention**.
- UNESCO. (2023). Guidance for generative AI in education and research. In *Guidance for generative AI in education and research*. UNESCO. <https://doi.org/10.54675/ewzm9535>
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly: Management Information Systems*, 27(3), 425–478. <https://doi.org/10.2307/30036540>
- Venkatesh, V., Walton, S. M., Thong, J. Y. L., & Xu, X. (2012). CONSUMER ACCEPTANCE AND USE OF INFORMATION TECHNOLOGY: EXTENDING THE UNIFIED THEORY OF ACCEPTANCE AND USE OF TECHNOLOGY. In *MIS Quarterly* (Vol. 36, Issue 1). <http://ssrn.com/abstract=2002388>
- Venter, I. M., Blignaut, R. J., Cranfield, D. J., Tick, A., & Achi, S. El. (2025). AI versus tradition: shaping the future of higher education. *Journal of Applied Research in Higher Education*, 17(7), 151–167. <https://doi.org/10.1108/JARHE-12-2024-0702>
- Yeralan, S., & Lee, L. A. (2023). Generative AI: Challenges to higher education. In *Sustainable Engineering and Innovation* (Vol. 5, Issue 2, pp. 107–116). Research and Development Academy. <https://doi.org/10.37868/sei.v5i2.id196>
- Yilmaz, F. G. K., Yilmaz, R., & Ceylan, M. (2023). Generative Artificial Intelligence Acceptance Scale: A Validity and Reliability Study. *International Journal of Human-Computer Interaction*. <https://doi.org/10.1080/10447318.2023.2288730>

Yusuf, A., Pervin, N., & Román-González, M. (2024). Generative AI and the future of higher education: a threat to academic integrity or reformation? Evidence from multicultural perspectives. *International Journal of Educational Technology in Higher Education*, 21(1). <https://doi.org/10.1186/s41239-024-00453-6>

Zlateva, P., Steshina, L., Petukhov, I., & Veleev, D. (2024). A Conceptual Framework for Solving Ethical Issues in Generative Artificial Intelligence. *Frontiers in Artificial Intelligence and Applications*, 381, 110–119. <https://doi.org/10.3233/FAIA231182>