

## مسئولية المراجع الخارجي في ظل مخاطر التواطؤ داخل تقنيات الذكاء الاصطناعي

د/ ياسر سعيد محمود الورداني / د/ محمود محمد عبد المنعم منسي / د/ خالد جمعة السيد حسين  
مدرس المحاسبة والمراجعة / مدرس المحاسبة والمراجعة / مدرس المحاسبة والمراجعة  
كلية التجارة- جامعة الأزهر / كلية التجارة- جامعة الأزهر / كلية التجارة- جامعة الأزهر

### ملخص الدراسة

تهدف هذه الدراسة إلى تحليل أثر مخاطر التواطؤ داخل تقنيات الذكاء الاصطناعي على مستوى تحمل المراجع الخارجي لمسؤوليته المهنية والقانونية، في ضوء التحديات التقنية المتزايدة المرتبطة باستخدام الخوارزميات الذكية في بيئة المراجعة.

أعتمدت الدراسة على المنهج الوصفي التحليلي، وتم جمع البيانات ميدانياً باستخدام أداة استبيان موجهة إلى عينة من المراجعين في مصر، وشمل التحليل خمسة أبعاد أساسية تمثل أنواع المخاطر التقنية (مثل التحيز الخوارزمي، والتواطؤ التقني، وضعف الشفافية)، إضافة إلى متغيرات ترتبط بالمسؤولية المهنية والقانونية للمراجع. توصلت النتائج إلى أن بعض أدوات الذكاء الاصطناعي تُظهر تحيزات مدمجة تؤثر سلباً على قدرة المراجع على اكتشاف الأخطاء الجوهرية، وتوصلت أيضاً إلى أن زيادة مخاطر التواطؤ في البرمجة تؤدي إلى زيادة مسؤولية المراجع عن الإفصاح وتفسير القرارات، كما توصلت أيضاً إلى أن عدم فهم الخوارزميات المعقدة (Black-box AI) يقيد قدرة المراجع على ممارسة الحكم المهني، بالإضافة إلى أن نقص إختبارات الأداء الفني للأدوات الذكية يزيد من فرص المساءلة القانونية. أوصت الدراسة بضرورة مشاركة المراجع في تطوير أدوات الذكاء الاصطناعي، وفرض معايير مهنية ملزمة بشأن الإفصاح والحكم المهني في بيئة الذكاء الاصطناعي، وتوفير تدريب تقني مستمر للمراجعين حول أدوات الذكاء الاصطناعي وآليات عملها.

**الكلمات المفتاحية:** الذكاء الاصطناعي، التواطؤ التقني، المسؤولية القانونية، الإفصاح المهني، التحيز الخوارزمي.  
الكلمات المفتاحية: الذكاء الاصطناعي- التواطؤ التقني- المسؤولية القانونية- الإفصاح المهني- التحيز الخوارزمي

## **Abstract**

This study aims to analyze the impact of collusion risks within artificial intelligence (AI) technologies on the external auditor's level of professional and legal responsibility, in light of the growing technological challenges associated with the use of intelligent algorithms in the audit environment.

The study adopted a descriptive-analytical methodology and collected data through a field survey directed at a sample of auditors In Egypt. The analysis covered five core dimensions representing types of technical risks (such as algorithmic bias, technical collusion, and lack of transparency), In addition to variables related to the auditor's professional and legal responsibility.

The results revealed that some AI tools exhibit embedded biases that negatively affect the auditor's ability to detect material misstatements. The study also found that increased collusion risks in programming lead to greater responsibility for the auditor to disclose and explain decisions. Furthermore, the complexity and opacity of certain AI algorithms (Black-box AI) were found to limit the auditor's ability to exercise professional

judgment. Additionally, the lack of performance testing for intelligent tools increases the likelihood of legal accountability.

The study recommends the active involvement of auditors in the development of AI tools, the enforcement of mandatory professional standards regarding disclosure and judgment in AI-supported environments, and the provision of continuous technical training for auditors on how AI tools function.

**Keywords:** Artificial Intelligence, Technical Collusion, Legal Responsibility, Professional Disclosure, Algorithmic Bias.

## الفصل الاول

### أولاً: المقدمة

شهدت العقود الأخيرة تطوراً مذهلاً في تقنيات الذكاء الاصطناعي (AI)، مما أدى إلى إعادة تشكيل ملامح العديد من المهن، وعلى رأسها مهنة المراجعة الخارجية، حيث أصبحت أدوات الذكاء الاصطناعي تدخل بشكل متزايد في عمليات جمع الأدلة، وتقييم المخاطر، والتحقق من البيانات الضخمة، مما يعزز من كفاءة وجودة الأداء المهني (Issa, Sun, & Vasarhelyi, 2016). ومع التوسع في استخدام هذه التقنيات، ظهرت تحديات جديدة تتعلق بمصداقية الأنظمة الذكية، خصوصاً تلك التي تعمل كـ "صناديق سوداء" لا تسمح بفهم آلية اتخاذ القرار داخلها (Coglianese & Lehr, 2019). وفي هذا السياق، برزت تساؤلات حادة حول المسؤولية المهنية للمراجع الخارجي عندما يعتمد في حكمه المهني على أنظمة قد تكون معرضة لخلل أو حتى تواطؤ برمجي مقصود. لقد حذرت العديد من الدراسات من احتمالية استخدام الذكاء الاصطناعي بشكل متحيز أو مضلل داخل المؤسسات، إما نتيجة برمجة غير شفافة، أو تدخل بشري غير نزيه في تصميم الخوارزميات (Binns, 2018; Mittelstadt et al., 2016) وهو

ما يطرح إشكالية مركزية أمام مهنة المراجعة: هل يتحمل المراجع الخارجي مسؤولية قانونية أو مهنية إذا اعتمد على نتائج صادرة من نظام ذكاء اصطناعي متواطئ؟ وبالنظر إلى حساسية الدور الذي يؤديه المراجع الخارجي في تعزيز الثقة في المعلومات المالية، يصبح من الضروري دراسة هذه الإشكالية بعمق، خاصة مع غياب أدلة مهنية ومعايير واضحة تحكم طبيعة العلاقة بين الذكاء الاصطناعي والمسئولية المهنية في حالات التواطؤ التقني أو الإهمال الرقمي. ومع تزايد اعتماد الشركات على أنظمة الذكاء الاصطناعي في معالجة البيانات المالية، أصبحت احتمالية تواطؤ هذه الأنظمة مع أطراف داخلية أو خارجية أمراً يحتاج إلى دراسة متأنية، خاصة فيما يتعلق بمسئولية المراجع عن اكتشاف هذا النوع من التلاعب، والذي يختلف عن الأساليب التقليدية للتحايل (Deng, A, 2020, ) Le, H. H. 2025

### ثانياً: مشكلة الدراسة

مع التوسع المتسارع في استخدام تقنيات الذكاء الاصطناعي في بيئات الأعمال، أصبحت نظم المراجعة تعتمد بشكل متزايد على الخوارزميات الذكية في تنفيذ مهام جوهرية مثل تقييم المخاطر وتحليل البيانات (Issa, Sun, & Vasarhelyi, 2016) غير أن هذه النظم – رغم كفاءتها الظاهرة – تعاني من نقص في الشفافية، مما يجعل من الصعب على المراجع الخارجي فهم أو تفسير كيفية التوصل إلى النتائج (Coglianese & Lehr, 2022).

وبالرغم من الفوائد العديدة لهذه التقنيات، فإن استخدامها يطرح تحديات تتعلق بالثقة والشفافية والمسئولية، ومن أبرز تلك التحديات ما يُعرف بـ"التواطؤ التقني"، حيث يمكن استغلال هذه الأنظمة سواء من خلال تصميم خوارزميات متحيزة، أو برمجيات تتعمد إخفاء الانحرافات المالية مما يضع مسؤولية متزايدة على عاتق المراجع الخارجي (Crawford & Paglen, 2021; Binns, 2018; Mittelstadt et al., 2016).

وتكمن المشكلة الرئيسية في غياب معايير مهنية واضحة تحدد مسؤولية المراجع تجاه أدوات الذكاء الاصطناعي، وهو ما يعرضه لمخاطر قانونية وأخلاقية في بيئة مراجعة معقدة تقنياً (IAASB, 2020).

ومن هنا تنبع مشكلة الدراسة في التساؤل عن مدى مسؤولية المراجع الخارجي إذا اعتمد على نتائج أنظمة ذكية ثبت لاحقاً أنها تعرضت لتواطؤ أو تلاعب داخلي. ويمكن صياغتها في محاولة الإجابة عن التساؤل الرئيسي التالي: ماهي مسؤوليات المراجع الخارجي المهنية والقانونية في ظل مخاطر التواطؤ داخل تقنيات الذكاء الاصطناعي؟

ويشتق من هذا السؤال الاسئلة الآتية:

- ١- ما أهم أنواع التواطؤ المحتمل داخل نظم الذكاء الاصطناعي المستخدمة في المحاسبة والمراجعة؟
- ٢- إلى أي مدى تؤثر هذه المخاطر على قدرة المراجع الخارجي في أداء مهامه المهنية؟
- ٣- ما حدود المسؤولية القانونية للمراجع الخارجي في ظل الاعتماد المتزايد على الذكاء الاصطناعي؟
- ٤- كيف يمكن تعزيز الضوابط المهنية والأخلاقية للتعامل مع هذه التحديات الجديدة؟

### ثالثاً: أهداف الدراسة:

- تحديد أنواع وأبعاد مخاطر التواطؤ داخل تقنيات الذكاء الاصطناعي.
- تحليل أثر هذه المخاطر على مسؤولية المراجع الخارجي.
- تقديم تصور مقترح لتعزيز ممارسات المراجعة في بيئة مدعومة بالذكاء الاصطناعي.
- بيان حدود الحكم المهني للمراجع في ظل النماذج الخوارزمية المعقدة.

### رابعاً: أهمية الدراسة:

تنبع أهمية هذه الدراسة من كونها تتناول موضوعاً حديثاً ومتشعباً يمس صميم مهنة المراجعة، في ظل التحول الرقمي المعاصر، كما أنها تساهم في:

- إثراء الجانب العلمي حول المخاطر التقنية غير التقليدية في بيئة المراجعة (Brynjolfsson & McAfee, 2023).
- مساعدة المراجعين على فهم حدود مسؤوليتهم في بيانات مدفوعة بالذكاء الاصطناعي.
- تحفيز الجهات المهنية والتنظيمية على وضع ضوابط جديدة خاصة بالذكاء الاصطناعي في العمل المحاسبي.

### خامساً: فروض الدراسة

- الفرض الرئيسي الصفري:** لا يوجد تأثير ذو دلالة إحصائية لمخاطر التواطؤ داخل تقنيات الذكاء الاصطناعي على مستوى تحمل المراجع الخارجي لمسئولية المهنية والقانونية.
- الفرض الرئيسي البديل:** يوجد تأثير ذو دلالة إحصائية لمخاطر التواطؤ داخل تقنيات الذكاء الاصطناعي على مستوى تحمل المراجع الخارجي لمسئولية المهنية والقانونية.
- ويشتق من الفرض الصفري والبديل الفروض الفرعية الآتية:
- ١- **الفرض الاول الصفري:** لا يوجد تأثير ذو دلالة إحصائية لمخاطر التحيز الخوارزمي داخل أنظمة الذكاء الاصطناعي على مسؤولية المراجع الخارجي عن اكتشاف الأخطاء الجوهرية.
  - الفرض الاول البديل:** يوجد تأثير ذو دلالة إحصائية لمخاطر التحيز الخوارزمي داخل أنظمة الذكاء الاصطناعي على مسؤولية المراجع الخارجي عن اكتشاف الأخطاء الجوهرية.
  - ٢- **الفرض الثاني الصفري:** لا يوجد تأثير ذو دلالة إحصائية لمخاطر التواطؤ التقني في برمجة خوارزميات الذكاء الاصطناعي على التزام المراجع بمسئولية الإفصاح المهني والقانوني.
  - الفرض الثاني البديل:** يوجد تأثير ذو دلالة إحصائية لمخاطر التواطؤ التقني في برمجة خوارزميات الذكاء الاصطناعي على التزام المراجع بمسئولية الإفصاح المهني والقانوني.

- ٣- **الفرض الثالث الصفري:** لا يوجد تأثير ذو دلالة إحصائية لمخاطر الإهمال الرقمي عند إعداد أدوات الذكاء الاصطناعي على التزام المراجع بمسئوليته القانونية.
- الفرض الثالث البديل:** يوجد تأثير ذو دلالة إحصائية لمخاطر الإهمال الرقمي عند إعداد أدوات الذكاء الاصطناعي على التزام المراجع بمسئوليته القانونية.
- ٤- **الفرض الرابع الصفري:** لا يوجد تأثير ذو دلالة إحصائية لمخاطر ضعف الشفافية في الخوارزميات الذكية على ممارسة المراجع للحكم المهني.
- الفرض الرابع البديل:** يوجد تأثير ذو دلالة إحصائية لمخاطر ضعف الشفافية في الخوارزميات الذكية على ممارسة المراجع للحكم المهني.
- ٥- **الفرض الخامس الصفري:** لا تختلف تصورات المراجعين بشأن مخاطر الذكاء الاصطناعي ومسئوليتهم المهنية تبعًا لاختلاف خصائصهم الشخصية والمهنية.
- الفرض الخامس البديل:** تختلف تصورات المراجعين بشأن مخاطر الذكاء الاصطناعي ومسئوليتهم المهنية تبعًا لاختلاف خصائصهم الشخصية والمهنية.

### سادساً: حدود الدراسة:

- ١- اقتصرت الدراسة على قياس إدراك المراجعين للمخاطر والمسئولية المهنية، وهو ما يعكس تصوراتهم الشخصية، وليس الأداء الفعلي أو السلوكي داخل بيئة العمل الواقعية.
- ٢- ركزت الدراسة على أربعة أبعاد أساسية لمخاطر الذكاء الاصطناعي (مثل التحيز، التواطؤ، ضعف الشفافية، الإهمال الرقمي) وعلى محورين للمسئولية (المهنية والقانونية)، دون التوسع في جوانب أخرى محتملة كالمسئولية الأخلاقية أو المساءلة التكنولوجية.
- ٣- لم تتضمن الدراسة اختباراً ميدانياً أو عملياً لأداء أدوات الذكاء الاصطناعي، بل اعتمدت فقط على آراء المراجعين حول مخاطرها وفعاليتها المفترضة.
- ٤- لم يتم التطرق إلى سياسات الشركات أو مكاتب المراجعة تجاه استخدام الذكاء الاصطناعي، والتي قد يكون لها دور في تفسير مستوى التحيز أو التواطؤ أو الإفصاح.

### سابعاً: منهجية الدراسة وأدواتها

تعتمد هذه الدراسة على المنهج الوصفي التحليلي، بهدف تحليل الأدبيات المرتبطة بالذكاء الاصطناعي ومسئولية المراجع، ثم إجراء دراسة ميدانية باستخدام أداة الاستبيان الموجهة إلى المراجعين القانونيين في كبرى شركات المحاسبة في مصر، وسيتم تحليل البيانات باستخدام الأساليب الإحصائية المناسبة مثل تحليل التباين والانحدار لاختبار صحة الفروض.

### ثامناً: مصطلحات الدراسة

- الذكاء الاصطناعي (Artificial Intelligence) يقصد به الأنظمة والبرمجيات التي تحاكي قدرات الإنسان الذهنية مثل التعلم واتخاذ القرار، وتستخدم في مهام المراجعة لتحليل البيانات والتنبؤ بالمخاطر (Russell & Norvig, 2022).

- التواطؤ التقني (Algorithmic Collusion) يشير إلى التلاعب المتعمد داخل النظم أو بين أطراف متعددة لتحقيق غايات غير قانونية أو مضللة، وقد يتم من خلال

- برمجة الذكاء الاصطناعي بطريقة تؤدي إلى إخفاء أدلة الغش أو الانحراف المالي (Ezrachi & Stucke, 2016, OECD, 2017).
- المسؤولية المهنية للمراجع: التزامه الأخلاقي والتقدير في أداء المراجعة باحترافية وفقاً للمعايير المهنية المعتمدة (IFAC, 2010).
- المسؤولية القانونية للمراجع: التزامه القانوني تجاه الأطراف المتعاملة في حال ثبوت الإهمال أو التقصير أو عدم الإفصاح عن المخاطر الجوهرية (IAASB, 2020).

### تاسعاً: خطة الدراسة

تنقسم الدراسة إلى خمسة فصول رئيسية على النحو التالي:

الفصل الأول: الإطار العام للدراسة

الفصل الثاني: الدراسات السابقة وتحليلها النقدي.

الفصل الثالث: الإطار النظري – الذكاء الاصطناعي، التواطؤ، المسؤولية المهنية والقانونية.

الفصل الرابع: منهجية الدراسة وتحليل البيانات.

الفصل الخامس: النتائج والتوصيات.

### الفصل الثاني: الدراسات السابقة وتحليلها النقدي

#### أولاً: عرض الدراسات السابقة:

تناولت العديد من الدراسات الحديثة قضايا الذكاء الاصطناعي والتواطؤ والمسؤولية المهنية للمراجع من زوايا متعددة، فقد ناقشت دراسة (Burrell 2016) مشكلة "الصندوق الأسود" في الأنظمة الذكية، وأكدت صعوبة تفسير مخرجات الذكاء الاصطناعي مما يؤثر على قدرة المراجع في تتبع منطق القرار المهني، كما ركزت دراسة (Binns 2018) على مفاهيم العدالة والمساءلة في الأنظمة الذكية، وحذرت من إمكانية التلاعب الخفي في تصميم الخوارزميات، كما حذرت دراسة (Ezrachi & Stucke 2016) من التواطؤ الخوارزمي بين الأنظمة الذكية، موضحة أن الخوارزميات قد تتعاون ضمناً دون وجود تواصل بشري مباشر، مما يعقد مهمة

اكتشاف الغش أو التلاعب .

أما دراستي (Diakopoulos, 2016, Johnson, 2021) فقد ركزت على ضرورة وجود تشريعات تلزم مطوري الخوارزميات بالشفافية، وذلك لحماية مهنة المحاسبة والمراجعة، وتكاملت معها دراسة (Raisch & Krakowski, 2021) التي تناولت إعادة توزيع المسؤولية في ظل هيمنة الذكاء الاصطناعي، مشيرة إلى أن مسؤولية المراجع لم تعد حصرية بل أصبحت موزعة بين الإنسان والنظام. وفي ذات الإطار، ناقشت دراسة (Falah, et al, 2025) تحديات المراجع في التعامل مع الأدلة الناتجة عن أنظمة الذكاء الاصطناعي، وتطرق إلى قضية المسؤولية عند الاعتماد عليها دون فهم آلية عملها، وهدفت دراسة Munoko et al, (2020) إلى دراسة الآثار الأخلاقية لاستخدام أنظمة الذكاء الاصطناعي من خلال الجمع بين إطارين أخلاقيين مُستقبليين، وأكدت الدراسة على أنه يمكن أن تتعرض مبادئ المراجعة الأساسية مثل الشك المهني والكفاءة والرعاية والحكم للخطر من خلال أنظمة الذكاء الاصطناعي غير القابلة للتفسير والتي تعيق الرقابة البشرية والفهم، بالإضافة إلى ذلك، فإن الاعتماد على الذكاء الاصطناعي في المراقبة والتقييم المستمر يزيد من مخاطر الخصوصية والأمن السيبراني، ويمكن أن يؤدي عدم الثقة في الذكاء الاصطناعي الغامض إلى خلق فجوات في المساءلة، مما يؤثر على تصورات المسؤولية العامة، ويؤكد مونوكو وآخرون على أهمية الشفافية والاستقلال والتعاون بين الإنسان والذكاء الاصطناعي في وظائف المراجعة، حتى مع الأتمتة، كما أوضحت دراسة Murikah et al. (2024) أن دمج الذكاء الاصطناعي في المراجعة يطرح تحديات أخلاقية كبيرة، بما في ذلك التحيزات الخوارزمية والشفافية والمساءلة والإنصاف، كما استكشفت الدراسة مصادر التحيز والمخاطر التي تشكلها أنظمة الذكاء الاصطناعي المطبقة في المراجعة والتفاعلات المعقدة اللاحقة لها وتأثيراتها.

تشير هذه الدراسات في مجملها إلى أن الذكاء الاصطناعي يحدث تحولاً في بيئة المراجعة، يفرض على المراجعين إعادة فهم مسؤولياتهم المهنية والقانونية، ويستدعي إطاراً قانونياً وأخلاقياً جديداً يواكب هذا التحول.

### ثانياً: التحليل المقارن للدراسات السابقة

تكشف الدراسات السابقة عن تباين في طرق تناول موضوع الذكاء الاصطناعي وتأثيره على مهنة المراجعة، فمنها ما ركز على الجانب القانوني لمسئولية المراجع في ظل الأنظمة الذكية مثل دراسة (Raisch & Krakowski, 2021) ومنها ما اهتم بالإشكاليات الأخلاقية والتقنية مثل الشفافية والتواطؤ الخفي (Ezrachi & Burrell, 2016, Stucke, 2016) كذلك ظهرت دراسات ذات طابع تجريبي وميداني (Munoko et al., 2022؛ Murikah et al., 2024) تسلط الضوء على أثر الأنظمة الذكية على تقدير المراجع وسلوكه المهني.

### ثالثاً: أوجه الاتفاق والاختلاف

اتفقت معظم الدراسات على أن الذكاء الاصطناعي يمثل نقلة نوعية في المراجعة، لكنه في ذات الوقت يُنتج تحديات جديدة تتعلق بالشفافية والتحيز والتواطؤ الخفي لم تكن موجودة في المراجعة التقليدية، وأبرزت جميع الدراسات تقريباً الحاجة إلى تعزيز الشفافية ومهارات التقدير المهني، بينما اختلفت في معالجة مسؤولية المراجع؛ فبعضها رأى أن المسؤولية لا تزال قائمة ولكن بصيغة موسعة، في حين ركزت دراسات أخرى على أن الذكاء الاصطناعي يخفف من عبء المسؤولية الشخصية عن طريق تفويض الخوارزميات باتخاذ القرار) مثل Johnson, 2021 (Diakopoulos, 2016).

### رابعاً: الفجوة البحثية

يتبين من التحليل أن هناك فجوة واضحة في الدراسات العربية والميدانية التي تتناول العلاقة بين مخاطر التواطؤ الخوارزمي ومسئولية المراجع الخارجي، معظم الدراسات ركزت على الإطار النظري، أو على المحيط القانوني في الدول المتقدمة،

بينما لم تعالج كيفية إدراك المراجعين في البيئة العربية لهذه المخاطر، وتأثيرها على أدائهم، كما تفتقر البحوث الحالية إلى كيفية توضيح تعامل المراجع مع الأنظمة الذكية المتواطئة، ومدى مسؤوليته عن نتائجها، خاصة في ظل غياب معايير مهنية واضحة تُنظّم هذه العلاقة، وفي ضوء مراجعة الأدبيات، تتضح فجوة بحثية تتعلق بغياب نموذج ميداني تطبيقي يدمج بين الأبعاد الأخلاقية والقانونية لمخاطر الذكاء الاصطناعي، وتأثيرها على مسؤولية المراجع الخارجي، وتسعى هذه الدراسة إلى سد هذه الفجوة من خلال بناء إطار تحليلي يتيح قياس وعي المراجعين بتلك المخاطر وتقدير انعكاساتها على ممارساتهم المهنية.

### الفصل الثالث : الإطار النظري – الذكاء الاصطناعي، التواطؤ، المسؤولية المهنية والقانونية:

#### المبحث الأول: تقنيات الذكاء الاصطناعي ودورها في المراجعة الخارجية

##### ١- مفهوم تقنيات الذكاء الاصطناعي:

يشير الذكاء الاصطناعي إلى الأنظمة أو الأجهزة التي تُحاكي الذكاء البشري لأداء المهام وتحسين الأداء ذاتيًا من خلال التعلم من البيانات (Russell & Norvig, 2021) وفي مجال المراجعة الخارجية، يُستخدم الذكاء الاصطناعي في مهام مثل تحليل البيانات، واكتشاف الأنماط غير العادية، وتقييم المخاطر، والكشف عن التلاعب المحتمل مما يساعد في تحسين جودة وكفاءة المراجعة (Issa et al., 2016, Russell & Norvig, 2022).

##### ٢- تطبيقات الذكاء الاصطناعي في المراجعة:

من أبرز أدوات وتقنيات الذكاء الاصطناعي المستخدمة في المراجعة : تقنيات التعلم الآلي (Machine Learning) ، ومعالجة اللغة الطبيعية (NLP) ، وتحليل البيانات الضخمة (Big Data Analytics) (Appelbaum et al., 2017) ، وتُمكن هذه الأدوات المراجع من استيعاب كمٍّ هائل من المعلومات في وقت قصير، مما يحسن من كفاءة أداء الفحص والتحقق، حيث تُستخدم الخوارزميات لاكتشاف الأنشطة غير الطبيعية وتحليل المعاملات المالية بكفاءة عالية، وتساهم هذه الأنظمة

في تعزيز كفاءة المراجع وتحسين قدرته على اكتشاف الأخطاء أو التلاعب، كما تسمح له بتخصيص وقته وجهده نحو المناطق ذات المخاطر الأعلى (Yoon, Hoogduin, & Zhang, 2021).

وتتميز أدوات التعلم الآلي بقدرتها على التعرف على العلاقات غير الظاهرة بين المتغيرات المالية، مما يمكن النظام من تقديم مؤشرات مبكرة عن وجود خلل محتمل أو احتيال مالي، حتى دون وجود دليل مباشر، أما تقنيات NLP، فتستخدم حاليًا في مراجعة العقود، والكشف عن البنود غير العادية، ومقارنة محتوى النصوص مع السياسات المحاسبية المعتمدة (Rozario & Vasarhelyi, 2018). وتشير الممارسات الحديثة في كبرى شركات المحاسبة إلى استخدام الذكاء الاصطناعي في المراجعة المستمرة (Continuous Auditing) وهي عملية قائمة على تحليل آني ودوري للبيانات المالية، مما يقلل الاعتماد على العمل الميداني ويعزز فعالية الرقابة المستمرة (Raji et al, 2020).

ورغم هذه المزايا، تشير الدراسات إلى أن الاستخدام المتزايد لتقنيات الذكاء الاصطناعي يفرض تحديات مهنية، تتعلق بقدرة المراجع على فهم منطق هذه الأدوات، وما إذا كانت النتائج التي تنتجها قابلة للتفسير بما يتوافق مع متطلبات معايير المراجعة (Zhang et al., 2022; Du & Xiong, 2025; Lehner et al, 2022).

### ٣- دور المراجع في بناء وتطوير تقنيات الذكاء الاصطناعي

رغم التوسع المتسارع في تطبيقات الذكاء الاصطناعي في مجال المراجعة، إلا أن المراجع الخارجي غالبًا ما يُستبعد من مراحل تطوير وتصميم هذه التقنيات. فأنظمة الذكاء الاصطناعي تُبنى عادةً من قبل فرق تقنية تتكون من علماء بيانات ومهندسين ومبرمجين، دون إشراك فعلي للمراجعين في تحديد منطق عمل الخوارزميات أو آليات اتخاذ القرار داخل النظام، مما قد يؤدي إلى فجوات بين مخرجات الذكاء الاصطناعي ومتطلبات مهنة المراجعة القائمة على الحكم المهني والشفافية (Richins et al., 2017; Appelbaum et al., 2017)، ويُضعف من قدرة

المراجع على فهم مخرجاته أو التدخل في تقييم دقتها، مما يؤدي إلى حالة من عدم الثقة والحذر تجاه استخدام هذه الأدوات (Moll & Yigitbasioglu, 2019)، كما يشعر بعض المراجعين أن دورهم قد أصبح هامشيًا في ظل أنظمة لا يستطيعون تفسير قراراتها، وهو ما يثير مخاوف مهنية ومسئولية قانونية محتملة. وقد أشار Appelbaum et al. (2020) إلى أن افتقار المراجع للفهم التقني المتعمق في تصميم الخوارزميات يضعف من قدرته على فحص افتراضات النمذجة أو اكتشاف التحيزات المدخلة عن غير قصد في بيانات التدريب، كما أظهرت دراسات حديثة أن غياب المشاركة الفعالة للمراجع في مراحل بناء النماذج يجعل من الصعب عليه التحقق من ملائمة النظام لطبيعة بيئة العمل، أو منطق اتخاذ القرار الذي تُبنى عليه استنتاجاته (Moll & Yigitbasioglu, 2019).

هذا الفصل بين الدور المهني والدور التقني يُسهم في خلق فجوة ثقة، ويؤثر على استعداد المراجع للاعتماد على نتائج هذه الأنظمة دون تحفظ، ما لم يتمكن من فهم كيف ولماذا اتخذ النظام قرارًا معينًا، ولتجاوز هذه الفجوة، تشير الأدبيات الحديثة إلى أهمية تبني ما يُعرف بأنظمة الذكاء الاصطناعي القابلة للتفسير (Explainable AI)، والتي تسمح للمستخدم بفهم منطق القرار وتقييمه بشكل نقدي، بما يعزز من كفاءة المراجع ويعيد الثقة في استخدام تلك الأدوات (Tambe et al., 2019; Parker, et al, 2022; Zhang et al., 2022).

#### ٤- طبيعة دور المراجع في ظل أدوات الذكاء الاصطناعي

في ظل تزايد استخدام أدوات الذكاء الاصطناعي في المراجعة، يثور تساؤل جوهري حول طبيعة الدور الذي يجب أن يقوم به المراجع الخارجي: هل يقتصر دوره على تنفيذ مخرجات النظام الذكي دون مراجعتها؟ أم أنه مطالب بتحليل النتائج وفهم منطقها؟ أم يتعين عليه أن يُفسّر هذه النتائج أمام الجهات المعنية؟ تشير الدراسات الحديثة إلى أن الأدوات الذكية قد تضع المراجع في موقع "المنفذ الفني" وليس "المحكم المهني"، مما قد يهدد استقلالته ويضعف سلطته التقديرية (Libby, Witz, 2024)، وفي المقابل تؤكد (IAASB (2020 أن الحكم المهني لا

يجب أن يُستبدل بأي نظام ذكي، وأن على المراجع أن يحتفظ بدوره التفسيري في تقييم كفاية وملاءمة الأدلة، بغض النظر عن التقنية المستخدمة في استخراجها. وقد اقترح بعض الباحثين مفهوم "**المراجع الرقمي (Digital Auditor)**"، وهو مراجع يمتلك مزيجاً من الفهم المهني والخبرة التقنية في مجال الذكاء الاصطناعي وتحليل البيانات، هذا المفهوم يعزز من دور المراجع ليصبح فاعلاً ونقدياً في التعامل مع أدوات الذكاء الاصطناعي، بدلاً من أن يكون دوراً سلبيّاً أو تقنيّاً بحثاً. فالمراجع الرقمي لا يكتفي بفحص النتائج فقط، بل يتفاعل مع منطق عمل الأنظمة، ويقيم جودة البيانات، ويفهم حدود الخوارزميات، مما يعزز الثقة ويقلل من المخاطر المهنية والقانونية المرتبطة باستخدام الذكاء الاصطناعي في المراجعة (Leocádio et al., 2024; Mokander, 2024, Raisch & Krakowski, 2021).

### المبحث الثاني: مخاطر تطبيق تقنيات الذكاء الاصطناعي في المراجعة الخارجية

رغم الفوائد الكبيرة لاستخدام الذكاء الاصطناعي في المراجعة، إلا أن الأدبيات العلمية تشير إلى عدد من المخاطر التي قد تهدد جودة وكفاءة واستقلالية عملية المراجعة، وفيما يلي أبرز هذه المخاطر:

#### ١- مخاطر التواطؤ التقني: (Technical Collusion)

يشير مصطلح "التواطؤ التقني" أو "الخوارزمي" إلى تلك الحالات التي يتم فيها استغلال أنظمة الذكاء الاصطناعي أو توجيهها بشكل متعمد لإخفاء معلومات أو تحريف نتائج مراجعة أو محاسبية، أو إخفاء بيانات معينة عن التحليل، وقد يحدث هذا التواطؤ إما عن طريق المبرمجين، أو بفعل استخدام بيانات مُضللة، أو حتى من خلال تحيزات مدمجة ضمن الخوارزميات ذاتها حيث تُصمم النماذج لتفضيل نتائج معينة أو تجاهل مؤشرات خطر (O'Neil, 2016; Barocas, Hardt, & Narayanan, 2019; Mittelstadt et al., 2016). وتبرز أهمية فهم هذه المخاطر في سياق المراجعة لضمان نزاهة وجودة النتائج، بالإضافة إلى الحاجة إلى تطبيق الذكاء الاصطناعي القابل لتفسير والرقابة البشرية. (Tjoa & Guan, 2020).

أظهرت دراسة قامت بها (Eubanks, 2018) حول استخدام الخوارزميات في المؤسسات الحكومية الأمريكية أن بعض الأنظمة كانت تتخذ قرارات متحيزة ضد الفئات الفقيرة عند تقييم الأهلية للدعم الاجتماعي، بسبب اعتمادها على بيانات تاريخية غير ممثلة. وفي مجال المراجعة، يشير (Murikah, et al, 2024) إلى أن أدوات تحليل البيانات الضخمة قد تعاني من تحيز ضمني إذا تم تدريبها على بيانات محاسبية صادرة من شركات ذات ممارسات مشكوك فيها، مما يؤدي إلى نتائج مغلوطة تُضلل المراجع.

## ٢- مخاطر الإهمال الرقمي: (Digital Negligence)

يقصد بالإهمال غير المتعمد في تصميم أو استخدام أدوات الذكاء الاصطناعي، مثل عدم اختبار صحة الخوارزميات بشكل كافٍ، أو الاعتماد على بيانات تدريب غير ممثلة أو متحيزة، أو غياب مراجعة دورية لأداء الأنظمة الذكية، بالإضافة إلى الاعتماد المفرط على هذه الأدوات دون فحص نقدي دقيق، مما يؤدي إلى تفويت أخطاء مادية قد تؤثر على جودة نتائج المراجعة (Amershi et al., 2019; Ribeiro et al., 2016, Binns, 2018).

والفرق بين التواطؤ التقني (Technical Collusion) والإهمال الرقمي (digital neglect) يكمن في نية الفعل وتأثيره، فالتواطؤ التقني هو خطر متعمد ينشأ عندما يتم برمجة النظام الذكي أو توجيهه بشكل يهدف إلى إخفاء حقائق أو تحريف نتائج، سواء من قبل المبرمجين أو الأطراف المعنية بالمعلومات المالية (O'Neil, 2016) في المقابل، الإهمال الرقمي هو خلل غير مقصود ينتج عن ضعف اختبار النظام، أو الاعتماد على بيانات تدريب ناقصة، أو غياب رقابة دورية على أداء الأنظمة الذكية، مما يؤدي إلى أخطاء غير مقصودة في النتائج (Amershi et al., 2019).

## ٣- مخاطر فقدان أو ضعف الشفافية: (Opacity Risk)

يمثل غياب الشفافية في الخوارزميات، والمعروف بمخاطر "الغموض الخوارزمي" (Opacity Risk)، تحديًا حقيقيًا في مهنة المراجعة الخارجية، خاصة مع تزايد اعتمادها على أدوات الذكاء الاصطناعي لتحليل كميات ضخمة من البيانات المالية، ويُضعف هذا الغموض من قدرة المراجع على تقييم جودة الأدلة الناتجة أو

تتبع منطق القرار، وهو ما قد يهدد مصداقية الرأي المهني. وبحسب Burrell (2016)، فإن الغموض في الخوارزميات يمكن أن يأخذ ثلاثة أشكال: غموضاً مقصوداً لحماية الملكية الفكرية، أو غموضاً ناتجاً عن عدم دراية تقنية لدى المستخدمين، أو غموضاً جوهرياً بسبب تعقيد الخوارزميات نفسها. كل هذه الأنواع تشكل عوائق حقيقية أمام المراجع الخارجي، الذي يعتمد على الفهم الكامل لآلية جمع وتحليل الأدلة قبل إصدار الحكم المهني. لذلك، فإن التعامل مع هذه المخاطر يتطلب تعزيز مفاهيم الذكاء الاصطناعي القابل للتفسير (XAI)، وفرض حد أدنى من الشفافية لمواءمة استخدام الخوارزميات مع المعايير المهنية في المراجعة.

### ٣- مخاطر تآكل التقدير المهني :

من أبرز التحديات التي تطرحها أنظمة الذكاء الاصطناعي في مهنة المراجعة الخارجية، خطر تآكل التقدير المهني (Erosion of Professional Judgment)، فعندما يعتمد المراجع بشكل مفرط على نتائج الأنظمة الذكية دون تطبيق التقييم النقدي أو الحكم المهني اللازم، قد يتحول دوره من فاحص خبير إلى منفذ سلبي لقرارات الخوارزميات (Sutton, Holt, & Arnold, 2016)، ويزداد هذا الخطر في ظل ما يُعرف بـ"التحيز نحو الأتمتة" (automation bias)، أي ميل المهنيين لقبول مخرجات الأنظمة المؤتمتة دون تشكيك، حتى عندما تكون النتائج خاطئة أو غير منطقية (Peters, C. P. H. 2022). هذا الانحراف يُضعف من جودة المراجعة ويقوّض مبادئ الاستقلالية والشك المهني التي تقوم عليها مهنة المراجعة K لذلك فإن الحفاظ على التوازن بين الاستفادة من قدرات الذكاء الاصطناعي وممارسة الحكم المهني يظل عنصراً أساسياً لضمان جودة العمل الرقابي والامتثال للمعايير الدولية (Scientific African, 2024).

## **المبحث الثالث: المسؤولية المهنية والقانونية للمراجع الخارجي**

### **١- طبيعة المسؤولية المهنية في البيئة التقليدية**

تقوم المسؤولية المهنية للمراجع الخارجي في البيئة التقليدية على الالتزام بمعايير المحاسبة والمراجعة، واستخدام الحكم المهني للكشف عن الأخطاء الجوهرية والغش، وتشمل هذه المسؤولية واجب الإخلاص والنزاهة، والتحقق من مدى صدق وعدالة القوائم المالية (IAASB, 2021).

### **٢- التحديات الجديدة في ظل الذكاء الاصطناعي**

أدى إدماج أنظمة الذكاء الاصطناعي في بيئة المراجعة إلى ظهور تحديات مهنية جديدة، من أبرزها: صعوبة تتبع منطق الخوارزميات (مشكلة الشفافية)، وغموض مصادر المخرجات (غياب المساءلة)، إلى جانب احتمالية وجود تحيز غير مقصود في النتائج. كما أن الاعتماد المفرط على هذه الأنظمة قد يُضعف من قدرة المراجع على تطبيق الحكم المهني، خاصة عندما تكون النتائج غير قابلة للتفسير أو تتطوي على تعقيد تقني يتجاوز مستوى فهمه. وتشير الدراسات إلى أن بعض المراجعين قد يترددون في رفض نتائج أنظمة تبدو دقيقة تقنياً، حتى عند تعارضها مع مؤشرات مهنية واضحة، مما قد يؤدي إلى تآكل تدريجي للحكم المهني (Burrell, 2016; Peters, 2022; Sutton et al., 2016).

### **٣- طبيعة المسؤولية القانونية للمراجع:**

تبرز المسؤولية القانونية للمراجع الخارجي كأحدى الركائز الأساسية في مهنة المراجعة، خاصة في ظل اعتماد أدوات الذكاء الاصطناعي لدعم الفحص والتحليل. ورغم الاستعانة بهذه الأدوات، يبقى المراجع مسؤولاً قانونياً عن الرأي الذي يصدره، ولا يُعفيه استخدام التكنولوجيا من مساءلته في حال ثبت تقصيره في تقييم الأدلة أو اعتماده على أنظمة لم تُختبر كفاءتها بالشكل الكافي. وتتمثل هذه المسؤولية في شقين: المسؤولية المدنية تجاه الأطراف المتضررة من تقرير مراجعة غير مدعوم، والمسؤولية الجنائية في حال ثبت الإهمال الجسيم أو التواطؤ (Arens et al., 2017; Kokina & Davenport, 2017; IFAC, 2025).

ووفقاً لمعيار ISA 240، فإن على المراجع أن يُبدي قدرًا عاليًا من الشك المهني في الحالات التي تُستخدم فيها الأنظمة الذكية، خاصة في تقييم ما إذا كانت الأدلة المستخلصة منها تُعتبر مناسبة وموثوقة (IAASB, 2020).

#### ٤- الفرق بين المسؤولية في المراجعة التقليدية والمدعومة بالذكاء الاصطناعي

تختلف طبيعة المسؤولية المهنية والقانونية للمراجع الخارجي بين بيئة المراجعة التقليدية وتلك المدعومة بالذكاء الاصطناعي. ففي المراجعة التقليدية، يتحمل المراجع المسؤولية الكاملة عن تقييم الأدلة واتخاذ القرار وفقاً لمعايير مهنية واضحة وباستخدام أدوات تحليل تقليدية، أما في بيئة الذكاء الاصطناعي، فيواجه المراجع تحديات متعلقة بصعوبة تفسير مخرجات الخوارزميات، وتحديد مصادر الخطأ أو التحيز عند حدوث خلل، مما يؤدي إلى توزيع غير واضح للمسؤولية بين المراجع ومطوري الأنظمة، كما تنخفض درجة الشفافية عند استخدام نماذج معقدة مثل التعلم العميق، ما يجعل من الصعب تتبع منطق القرار مقارنة ببيئة المراجعة التقليدية (Burrell, 2016; Sutton et al., 2016; Kokina & Davenport, 2017).

ورغم الاعتماد المتزايد على أدوات الذكاء الاصطناعي في المراجعة، فإن المراجع الخارجي لا يُعفى من المسؤولية القانونية أو المهنية لمجرد استخدامه لتقنيات ذكية، فقد أكدت المعايير الدولية للمراجعة، مثل ISA 200، أن المراجع يظل مسؤولاً عن إصدار رأي مهني مدعوم بأدلة كافية وملائمة، بغض النظر عن الأدوات المستخدمة، ويشمل ذلك التزامه بالشك المهني، وتقييم مخاطر التحريف الجوهرية، والتحقق من مصداقية مصادر البيانات حتى لو كانت ناتجة عن أنظمة ذكية (IAASB, 2021).

وفيما يلي جدول يوضح الفرق بين المسؤولية المهنية والقانونية للمراجع في كل من البيئة التقليدية والبيئة المدعومة بالذكاء الاصطناعي:

| وجه المقارنة          | المراجعة المدعومة بالذكاء الاصطناعي                                                  | المراجعة اليدوية التقليدية                                          |
|-----------------------|--------------------------------------------------------------------------------------|---------------------------------------------------------------------|
| إمكانية تتبع القرارات | منخفضة - الخوارزميات قد تكون غير شفافة                                               | مرتفعة - تعتمد على أدلة ورقية واضحة                                 |
| الحكم المهني          | قد يتراجع بسبب صعوبة تفسير نتائج بعض أدوات الذكاء الاصطناعي أو الاعتماد الزائد عليها | يعتمد المراجع على تقديره المهني في فحص الأدلة واتخاذ القرار النهائي |
| المساءلة القانونية    | تتقاسمها الأطراف المطورة للتقنية                                                     | مرتبطة بسلوك المراجع الفردي                                         |
| اكتشاف التواطؤ        | يتطلب أدوات تحليل خوارزمي وتعلم آلي                                                  | من خلال المقابلات والتحقيقات اليدوية                                |

المصدر: من إعداد الباحث استناداً إلى (Burrell, 2016؛ IFAC, 2023؛ Raisch & Krakowski, 2021).

وقد أكد العديد من الباحثين أن تطور أدوات المراجعة، لا سيما استخدام الذكاء الاصطناعي، لا يُسقط المسؤولية المهنية عن المراجع، بل يُعززها، فالمراجع مطالب اليوم بأن يجمع بين الحكم المهني والفهم التقني للأدوات الذكية التي يستخدمها، لضمان اتخاذ قرارات مدعومة بأدلة قابلة للتفسير والاعتماد (Sutton et al., 2016; Yoon et al., 2015). وبالتالي، فإن تطور التكنولوجيا يفرض على المراجع دوراً مضاعفاً في الفهم، والمساءلة، والرقابة.

## ٥- أنواع المسؤوليات الواقعة على عاتق المراجع الخارجي عند استخدام أدوات الذكاء الاصطناعي

مع تزايد الاعتماد على تقنيات الذكاء الاصطناعي في مجال المراجعة، بات من الضروري تحديد المسؤوليات التي يتحملها المراجع الخارجي عند استخدام هذه الأدوات، خاصة في ظل ما تفرضه البيئة المهنية والتنظيمية من التزامات، ويمكن تصنيف هذه المسؤوليات إلى ثلاث فئات رئيسية: مهنية، قانونية، وتقديرية.

**٥/١ المسؤولية المهنية:** تُعد المسؤولية المهنية من أبرز ركائز عمل المراجع الخارجي، حيث ينبغي له الالتزام بالمعايير الأخلاقية والمهنية حتى في ظل استخدام أدوات الذكاء الاصطناعي، إذ لا يُعفى من ممارسة الحكم المهني المستقل، ولا يجوز له الاعتماد الكامل على مخرجات الأنظمة التقنية دون مراجعتها وتقييمها، بما يضمن موثوقية النتائج واتساقها مع أهداف المراجعة (Leocádio et al., 2024; Baskarada et al., 2022; IAASB, 2020).

**٥/٢ المسؤولية القانونية:** يتحمل المراجع الخارجي مسؤولية قانونية عند استخدام أدوات الذكاء الاصطناعي، لا سيما فيما يتعلق بخصوصية البيانات والامتثال للتشريعات التنظيمية كالقوانين الخاصة بحماية البيانات، ويجب على المراجع التأكد من أن الأدوات المستخدمة لا تنتهك أي قوانين تنظيمية، وأن البيانات التي تُعالج عبر النماذج الذكية قد خضعت لإجراءات أمنية دقيقة، وتم الحصول عليها بطريقة قانونية (Kamareldawla, 2025; Raji et al., 2020).

**٥/٣ المسؤولية التقديرية:** تشير المسؤولية التقديرية إلى التزام المراجع الخارجي بتقييم مخرجات أدوات الذكاء الاصطناعي بصورة تحليلية، مع مراعاة مدى شفافيتها ومصداقيتها، وتحديد ما إذا كانت النتائج صالحة للاستخدام كأدلة إثبات، ويتطلب ذلك فحص جودة البيانات المستخدمة في تدريب النماذج، والانتباه لاحتمال وجود تحيز خوارزمي، بالإضافة إلى توظيف أدوات الذكاء الاصطناعي القابل للتفسير (XAI) مثل SHAP أو LIME لفهم آلية التوصل إلى الاستنتاجات (Leocádio et al., 2024; Müller et al., 2022; Raji et al., 2020).

**٦- تصور مقترح لتعزيز ممارسات المراجعة في بيئة مدعومة بالذكاء الاصطناعي**  
نظرًا لتزايد الاعتماد على تقنيات الذكاء الاصطناعي في أنشطة المراجعة، وظهور مخاطر مهنية جديدة تهدد جودة الأداء واستقلال المراجع، فقد أصبح من الضروري وضع إطار عملي يهدف إلى تعزيز ممارسات المراجعة من ناحية، وحماية المراجع من المخاطر التقنية من ناحية أخرى.  
ويقترح الباحثون في هذه الدراسة، أن يتكون الأطار المقترح من الأتي:

**٦/١ التأهيل المهني والتقني للمراجع:** ضرورة تضمين مهارات تحليل البيانات والذكاء الاصطناعي ضمن برامج إعداد وتأهيل المراجعين، مع تعزيز فهمهم للخوارزميات والأنظمة التنبؤية المستخدمة في المراجعة الذكية (Kokina, Pachamanova, & Corbett, 2021; IFAC, 2020).

٦/٢ **تفعيل الحكم المهني في ظل الأنظمة الذكية:** يجب الحفاظ على دور المراجع البشري في اتخاذ القرارات المهنية وعدم الاعتماد الكلي على مخرجات الذكاء الاصطناعي دون تحليل (Appelbaum, Kogan, & Vasarhelyi, 2022).

٦/٣ **إدارة المخاطر التقنية المرتبطة بالذكاء الاصطناعي:** يشمل ذلك إشراك خبراء متخصصين في تقييم الخوارزميات، والتحقق من خلوها من التحيز، ومراجعة التصميم لضمان الحيادية، بالإضافة إلى مراقبة أداء الأنظمة خلال فترة الاستخدام (Issa, Sun, 2022; Eilifsen & Messier, 2016; Vasarhelyi, 2016).

٦/٤ **تعزيز الإفصاح والمساءلة المهنية:** يجب على المراجع طلب توضيحات فنية مفصلة حول طريقة عمل النظم المستخدمة في إعداد التقارير المالية، كما يجب تضمين نتائج التحليل باستخدام الذكاء الاصطناعي في تقرير المراجعة بشكل واضح، من أجل تحديد مسؤولية المراجع عن التأكد من صلاحية الأنظمة المستخدمة (IAASB, 2023; IFAC, 2020).

٦/٥ **الرقابة المؤسسية والتشريعية:** ينبغي أن تصدر الجهات الرقابية معايير خاصة باستخدام الذكاء الاصطناعي في المراجعة، مع فرض الإفصاح عن طريقة عمل الأدوات الرقمية المستخدمة (IAASB, 2023; Kokina et al., 2021).

## الفصل الرابع: منهجية الدراسة

### تمهيد:

يُعد هذا الفصل بمثابة الأساس المنهجي الذي يستند إليه الباحثون في معالجة مشكلة الدراسة واختبار فروضها، ويتضمن هذا الفصل تحديد المنهج المستخدم، ووصف مجتمع وعينة الدراسة، بالإضافة إلى عرض أداة الدراسة المتمثلة في الاستبيان، وآلية تصميمه، وكيفية التحقق من صدقه وثباته، فضلاً عن الأساليب الإحصائية التي سيتم استخدامها لتحليل البيانات واختبار الفروض.

### أولاً: منهج الدراسة:

إعتمدت الدراسة على المنهج الوصفي التحليلي، نظراً لملاءمته لطبيعة أهداف البحث التي تركز على وصف الظاهرة موضوع الدراسة، وهي "مخاطر التواطؤ داخل تقنيات الذكاء الاصطناعي وتأثيرها على مسؤولية المراجع الخارجي"، وتحليل العلاقات بين المتغيرات المختلفة، واختبار الفروض المتعلقة بها.

### ثانياً: مجتمع الدراسة:

يتمثل مجتمع الدراسة في المراجعين القانونيين والمحاسبين العاملين بمكاتب المراجعة والمحاسبة الكبرى والمتوسطة بجمهورية مصر العربية، خاصة أولئك الذين لديهم دراية أو تعامل سابق (حتى وإن كان محدوداً) مع أدوات الذكاء الاصطناعي أو الأنظمة الرقمية الحديثة.

### ثالثاً: عينة الدراسة:

نظراً لصعوبة الوصول إلى كافة أفراد المجتمع، تم استخدام العينة القصدية، حيث تم اختيار مفردات العينة بناءً على معايير ملائمة مثل الخبرة العملية في مجال المراجعة، والمستوى العلمي، والانخراط في بيئة عمل تقنية أو رقمية، بلغ حجم العينة التي تم استهدافها من خلال توزيع أداة الاستبيان عدد (١٤٠) مفردة، وبلغ إجمالي الردود الفعلية الصالحة التي تم الإعتماد عليها في الدراسة ٩٧ رد.

### رابعاً: أداة الدراسة:

تمثلت أداة الدراسة في استبيان إلكتروني أعده الباحث باستخدام منصة Google Forms، وتكوّن من قسمين:

- القسم الأول: البيانات الديموغرافية (مثل: عدد سنوات الخبرة، المؤهل العلمي، نوع المكتب).
- القسم الثاني: محاور الاستبيان التي تعكس متغيرات الدراسة، وقد صُممت بنود هذا القسم استناداً إلى الأدبيات العلمية والدراسات السابقة، مع مراعاة وضوح العبارات

وسهولة فهمها، خاصة أن بعض المشاركين قد لا يمتلكون خلفية تقنية قوية عن الذكاء الاصطناعي.

تضمن الاستبيان ثمانية أبعاد رئيسية تمثل المتغيرات المستقلة والتابعة: وتشمل) مخاطر التحيز الخوارزمي، ومسئولية المراجع عن اكتشاف الأخطاء الجوهرية، ومخاطر التواطؤ التقني، ومسئولية المراجع عن الإفصاح المهني والقانوني، ومخاطر الإهمال الرقمي، والمسئولية القانونية للمراجع، ومخاطر ضعف الشفافية، وممارسة الحكم المهني.

وقد تم استخدام مقياس ليكرت الخماسي في قياس البنود (١: أعارض بشدة، ٢: أعارض، ٣: محايد، ٤: أوافق، ٥: أوافق بشدة).

#### خامساً: صدق وثبات أداة الدراسة:

• **الصدق الظاهري: (Face Validity)** تم عرض الاستبيان على مجموعة من المحكمين من الأساتذة المتخصصين في المحاسبة والمراجعة للتحقق من مدى ملائمة البنود ووضوحها.

• **الثبات (Reliability):** سيتم حسابه باستخدام معامل ألفا كرونباخ (Cronbach's Alpha) لكل بُعد من أبعاد الاستبيان، وذلك بعد جمع الاستجابات من العينة، للتأكد من اتساق الإجابات داخلياً.

سادساً: الأساليب الإحصائية:

سوف يتم تحليل البيانات باستخدام البرنامج الإحصائي (SPSS) أو ما يعادله، وتشمل الأساليب:

- الإحصاءات الوصفية: المتوسطات، الانحراف المعياري، النسب المئوية.
  - اختبار وتحليل التباين ( MNOVA ) لاختبار الفروق بين إجابات المراجعين بحسب الخصائص الديموغرافية.
  - اختبار الارتباط والانحدار لتحليل العلاقة بين المتغيرات المستقلة (مخاطر الذكاء الاصطناعي) والمتغيرات التابعة (مسئولية المراجع الخارجي).
- وفيما يلي عرض لخطوات التحليل الإحصائي على النحو التالي:

## ١ - تحليل سمات مفردات عينة الدراسة:

### ١/١ تبويب عينة الدراسة طبقاً للمؤهل الدراسي الذي حصل عليه المراجع:

#### جدول رقم (١)

توزيع افراد عينة الدراسة طبقاً لمتغير درجة المؤهل الدراسي

| م | المؤهل الدراسي الذي حصل عليه المراجع | العدد | %   | الترتيب |
|---|--------------------------------------|-------|-----|---------|
| ١ | بكالوريوس                            | ٥٧    | ٦٠  | ١       |
| ٣ | دبلوم                                | ٧     | ٧   | ٤       |
| ٤ | ماجستير                              | ٢٥    | ٢٥  | ٢       |
| ٥ | دكتوراه                              | ٨     | ٨   | ٣       |
| - | المجموع                              | ٩٧    | ١٠٠ | -       |

يتضح من الجدول رقم (١) أن توزيع مفردات عينة الدراسة وفقاً لمتغير " درجة المؤهل الدراسي الذي حصل عليه المراجع " يشير إلى أن اغلبيه عينة الدراسة من الحاصلين على البكالوريوس بنسبة (٦٠%)، بينما يأتي الحاصلين على مؤهل اعلى من البكالوريوس (ماجستير) في المرتبة الثانية بنسبة (٢٥%)، و(دكتوراه) في المرتبة الثالثة بنسبة (٨%)، و(دبلوم) في المرتبة الرابعة بنسبة (٧%) من إجمالي مفردات عينة الدراسة، مما يدل على تنوع العينة وثقلها في الاجابات التي يحتاجها الرد على إستبيان الدراسة.

### ١/٢ تبويب عينة الدراسة طبقاً للوظيفة التي يشغلونها في الوقت الحالي:

#### جدول رقم (٢)

توزيع افراد عينة الدراسة طبقاً لمتغير الوظيفة التي يشغلها المراجع

| م | الوظيفة التي يشغلها حالياً | العدد | %   | الترتيب |
|---|----------------------------|-------|-----|---------|
| ١ | شريك مكتب                  | ٥     | ٥   | ٤       |
| ٢ | مدير                       | ٢٦    | ٢٥  | ٢       |
| ٣ | رئيس مجموعة                | ٤٦    | ٤٨  | ١       |
| ٤ | مراجع                      | ١٥    | ١٧  | ٣       |
| ٥ | مراجع مساعد                | ٥     | ٥   | ٥       |
| - | المجموع                    | ٩٧    | ١٠٠ | -       |

يتضح من الجدول رقم (٢) أن توزيع مفردات عينة الدراسة وفقاً لمتغير " الوظيفة التي يشغلها المراجع " يشير إلى أن اغلبيه عينة الدراسة يشغلون وظيفة مدير ورئيس

مجموعة، حيث يمثلون نسبة (٧٣%)، يليها العاملين بوظيفة مراجع، ويمثلون نسبة (١٧%)، ثم العاملين بوظيفة شريك مكتب بنسبة ووظيفة مراجع مساعد، وكل منهم يمثل (٥%) من إجمالي مفردات عينة الدراسة.

### ١/٣ تبويب عينة الدراسة طبقاً لمدة الخبرة في مجال ممارسة المهنة:

#### جدول رقم (٣)

#### توزيع افراد عينة الدراسة طبقاً لمتغير مدة الخبرة

| م | مدة الخبرة بالسنوات | العدد | %   | الترتيب |
|---|---------------------|-------|-----|---------|
| ١ | من ١ - ٥            | ٩     | ٩   | ٤       |
| ٢ | من أكثر من ٥ - ١٠   | ٤٠    | ٤١  | ١       |
| ٣ | من أكثر من ١٠ - ١٥  | ٣     | ٣٨  | ٢       |
| ٤ | أكثر من ١٥          | ١٢    | ١٢  | ٣       |
|   | المجموع             | ٩٧    | ١٠٠ | -       |

يتضح من الجدول رقم (٣) أن توزيع مفردات عينة الدراسة وفقاً لمتغير " مدة خبرة المراجع " يشير إلى أن أغلبية عينة الدراسة مدة خبرة المراجع فيها أكثر من ٥ سنوات حتى ١٥ سنة حيث يمثلون نسبة (٧٩%)، يليها من بلغت مدة خبرته أكثر من ١٥ سنة بنسبة (١٢%)، ثم من بلغت مدة خبرته من سنة حتى ٥ سنوات، حيث يمثلون نسبة (٩%)، مما يدل على درجة الوعي المهني الذي يتمتع به أفراد العينة والذي أفاد كثيراً في درجة الوعي بموضوع الدراسة.

### ١/٤ تبويب عينة الدراسة طبقاً لجهة العمل:

#### جدول رقم (٤)

#### توزيع افراد عينة الدراسة طبقاً لمتغير جهة العمل

| م | جهة العمل       | العدد | %   | الترتيب |
|---|-----------------|-------|-----|---------|
| ١ | مكتب مراجعة خاص | ١٠    | ١٠  | ٣       |
| ٢ | شركة تضامن      | ٣٥    | ٣٧  | ٢       |
| ٣ | شركة مساهمة     | ٤٤    | ٤٥  | ١       |
| ٤ | جهة حكومية      | ٨     | ٨   | ٤       |
|   | المجموع         | ٩٧    | ١٠٠ | -       |

يتضح من الجدول رقم (٤) أن توزيع مفردات عينة الدراسة وفقاً لمتغير " جهة العمل " يشير إلى أن أغلبية عينة الدراسة تعمل في شركات مساهمة، حيث يمثلون

نسبة (٤٥%)، يليها من يعمل في شركات تضامن، بنسبة (٣٧%)، ثم من يعمل بمكتب مراجعة خاص، حيث يمثلون نسبة (١٠%)، وأخيراً من يعمل بجهة حكومية رقابية، حيث يمثلون نسبة (٨%) مما يدل على درجة تنوع جهة العمل ويساعد في اختبار فروض الدراسة.

## ٢- اختبار ألفا كرونباخ لقياس ثبات وصدق محتوى استبيان الدراسة.

تم استخدام معامل ألفا كرونباخ لاختبار تجانس متغيرات الدراسة وأسئلة الاستبيان المستخدمة في الاستقصاء، واتضح أن قيمة معامل الثبات والصدق تقترب من الواحد الصحيح (٧٧.٢%)، (٨٧.٧%) مما يعنى أنه يمكن الاعتماد على إجابات المستجيبين على العبارات التي تشملها الدراسة في اختبار الفروض من خلال إجراء الاختبارات الاحصائية عليها والاعتماد على نتائجها.

## ٣- مقياس ليكارت الخماسي لمعرفة اتجاهات آراء المستجيبين:

تم حساب كل من المتوسط الحسابي المرجح والانحراف المعياري وتم تقييم العبارات على أساس المتوسط المرجح وفقاً لمعايير الموافقة والموافقة التامة من عدمها في إطار مقياس ليكارت الخماسي الاتجاه وذلك كمايلي:

### جدول رقم (٥)

#### المحور الاول: مخاطر التحيز الخوارزمي

| المحور الاول: مخاطر التحيز الخوارزمي                             | المتوسط | الانحراف المعياري | الاتجاه العام |
|------------------------------------------------------------------|---------|-------------------|---------------|
| بعض أدوات الذكاء الاصطناعي قد تعطي نتائج غير دقيقة بسبب برمجتها  | ٤.٠٠    | ٠.٤٥٦             | أوافق         |
| أوافق على نتائج الذكاء الاصطناعي حتى لو لم أتأكد من دقتها        | ١.٧٤    | ٠.٥٠٦             | أعارض         |
| إذا كانت بيانات النظام غير ممثلة، فقد تؤدي لقرارات مراجعة خاطئة  | ٤.٠٣    | ٠.١٧٤             | أوافق         |
| لا بد من مراجعة مصدر بيانات النظام الذكي من أجل ضمان صدق النتائج | ٤.٣٣    | ٠.٥٣٥             | أوافق         |
| اجمالي المحور                                                    | ٤.١٥    | ٠.٢٤٠             | أوافق         |

ومن خلال تحليل المتوسطات الحسابية، والانحراف المعياري للاستجابات يتضح أن الاتجاه العام لاستجابات المستقصي منهم على معظم أسئلة الاستبيان الخاصة

بالمحور الأول (موافق) وذلك بمتوسط حسابي قدره (٤.١٥) وانحراف معياري قدره (٠.٢٤٠).

## جدول رقم (٦)

### المحور الثاني: قدرة المراجع على اكتشاف الأخطاء الجوهرية

| الاتجاه العام | الانحراف المعياري | المتوسط | المحور الثاني: قدرة المراجع على اكتشاف الأخطاء الجوهرية                                  |
|---------------|-------------------|---------|------------------------------------------------------------------------------------------|
| أوافق         | ٠.٤٥٦             | ٤.٠٠    | الاعتماد الكامل على الذكاء الاصطناعي قد يجعلني أغفل عن أخطاء مهمة                        |
| أعارض         | ٠.٣٤٨             | ٢.٠٦    | اثق في أدوات الذكاء الاصطناعي مهما كانت معقدة أو غير مفهومة                              |
| أوافق بشدة    | ٠.٤٩٩             | ٤.٥٦    | غموض طريقة عمل النظام الذكي يقلل من قدرتي على اكتشاف الانحرافات                          |
| أعارض         | ٠.٥٠٦             | ١.٧٤    | لا أحتاج لتحليل النتائج يدوياً طالما أن النظام قدمها لي                                  |
| أوافق         | ٠.٥٣٦             | ٤.٣٣    | كلما زاد اعتماد المراجع على الذكاء الاصطناعي قلت قدرته على ملاحظة الأخطاء الجوهرية بنفسه |
| أوافق         | ٠.٢٧٠             | ٤.١٨    | اجمالي المحور                                                                            |

ومن خلال تحليل المتوسطات الحسابية، والانحراف المعياري للاستجابات يتضح أن الاتجاه العام لاستجابات المستقصى منهم على معظم أسئلة الاستبيان الخاصة بالمحور الثاني (موافق) وذلك بمتوسط حسابي قدره (٤.١٨) وانحراف معياري قدره (٠.٢٧٠).

وهذا ما يثبت صحة الفرض الاول البديل وهو وجود علاقة معنوية ذات دلالة إحصائية بين مخاطر التحيز الخوارزمي داخل أنظمة الذكاء الاصطناعي ومسئولية المراجع الخارجي عن إكتشاف الأخطاء الجوهرية.

## جدول رقم (٧)

### المحور الثالث: مخاطر التواطؤ التقني

| الاتجاه العام | الانحراف المعياري | المتوسط | المحور الثالث: مخاطر التواطؤ التقني                                              |
|---------------|-------------------|---------|----------------------------------------------------------------------------------|
| أوافق         | ٠.٤٦٠             | ٤.٣٠    | قد يتم التلاعب بنتائج المراجعة إذا تدخل مبرمجون غير نزيهين في تصميم النظام       |
| أوافق بشدة    | ٠.٥١١             | ٤.٦١    | لا بد من مشاركة المراجع فعلياً وعملياً في برمجة النظام                           |
| أوافق بشدة    | ٠.٥١١             | ٤.٦١    | بعض أدوات الذكاء الاصطناعي قد تبرمج لتجنب اكتشاف الانحرافات أو الأخطاء المحاسبية |
| أعارض         | ٠.٢٦٦             | ٢.٠٣    | لا أعتقد أن التلاعب ممكن في أنظمة الذكاء الاصطناعي                               |
| أوافق بشدة    | ٠.٤٩٩             | ٤.٥٦    | من الخطر عدم إشراك المراجع في تطوير أدوات الذكاء الاصطناعي                       |
| أوافق         | ٠.٢٤٦             | ٣.٩٠    | اجمالي المحور                                                                    |

ومن خلال تحليل المتوسطات الحسابية، والانحراف المعياري للاستجابات يتضح أن الاتجاه العام لاستجابات المستقصى منهم على معظم أسئلة الاستبيان الخاصة بالمحور الثالث (موافق) وذلك بمتوسط حسابي قدره (٣.٩٠) وانحراف معياري قدره (٠.٢٤٦).

## جدول رقم (٨)

### المحور الرابع: الالتزام بمسئولية الإفصاح المهني والقانوني

| الاتجاه العام | الانحراف المعياري | المتوسط | المحور الرابع: الالتزام بمسئولية الإفصاح المهني والقانوني           |
|---------------|-------------------|---------|---------------------------------------------------------------------|
| وافق          | ٠.٥٢٠             | ٤.٤٦    | إذا شككت في دقة النظام الذكي، يجب أن أوضح ذلك في تقرير المراجعة     |
| وافق بشدة     | ٠.٥١١             | ٤.٦١    | أرى أن من مسؤوليتي أن أوضح إن كانت نتائج النظام واضحة أو غامضة      |
| وافق          | ٠.٤١٥             | ٣.٩٣    | أتحمل مهنيًا واجب الإفصاح عن أي خلل في أدوات المراجعة الذكية        |
| أعارض         | ٠.٢٢٦             | ٢.٠٣    | لا أرى ضرورة لشرح طريقة عمل النظام الذكي للجهات الرقابية أو العملاء |
| وافق          | ٠.٢٢٧             | ٤.٢٤    | إجمالي المحور                                                       |

ومن خلال تحليل المتوسطات الحسابية، والانحراف المعياري للاستجابات يتضح أن الاتجاه العام لاستجابات المستقصى منهم على معظم أسئلة الاستبيان الخاصة بالمحور الرابع (موافق) وذلك بمتوسط حسابي قدره (٤.٢٤) وانحراف معياري قدره (٠.٢٢٧).

وهذا ما يثبت صحة الفرض الثاني البديل وهو وجود علاقة معنوية ذات دلالة إحصائية بين مخاطر التواطؤ التقني في برمجة خوارزميات الذكاء الاصطناعي والالتزام المراجع بمسئولية الإفصاح المهني والقانوني.

## جدول رقم (٩)

### المحور الخامس: مخاطر الإهمال الرقمي

| الاتجاه العام | الانحراف المعياري | المتوسط | المحور الخامس: مخاطر الإهمال الرقمي                                          |
|---------------|-------------------|---------|------------------------------------------------------------------------------|
| وافق          | ٠.٤١٥             | ٣.٩٣    | من الخطر استخدام أدوات ذكية لم يتم اختبارها في بيئة العمل الفعلية            |
| وافق          | ٠.٣٦٠             | ٤.٠٧    | لا بد من اختبار أدوات الذكاء الاصطناعي حتى ولو كانت النتائج تبدو منطقية      |
| وافق          | ٠.١٠٢             | ٣.٩٩    | في بعض الحالات، لا يتم فحص أداء أدوات الذكاء الاصطناعي بشكل دوري أثناء العمل |
| أعارض         | ٠.٥٠٦             | ١.٧٤    | يمكن الاعتماد على النظام الذكي دون الحاجة لمراجعة خطواته يدويًا              |
| وافق          | ٠.٣٦٠             | ٤.٠٤    | يجب مراجعة أداء النظام الذكي بانتظام للتأكد من دقته                          |
| وافق          | ٠.١٤٠             | ٤.٠٥    | إجمالي المحور                                                                |

مسئولية المراجع الخارجي فهي ظل مخاطر التواطؤ داخل تقنياته الذكاء الاصطناعي

د/ محمود همد عبد المنعم منسي د/ محمود همد عبد المنعم منسي د/ محمود همد عبد المنعم منسي

ومن خلال تحليل المتوسطات الحسابية، والانحراف المعياري للاستجابات يتضح أن الاتجاه العام لاستجابات المستقصى منهم على معظم أسئلة الاستبيان الخاصة بالمحور الخامس (موافق) وذلك بمتوسط حسابي قدره (٤.٠٥) وانحراف معياري قدره (٠.١٤٠).

### جدول رقم (١٠) المحور السادس: المسؤولية القانونية للمراجع

| الاتجاه العام | الانحراف المعياري | المتوسط | المحور السادس: المسؤولية القانونية للمراجع                                                      |
|---------------|-------------------|---------|-------------------------------------------------------------------------------------------------|
| أعارض         | ٠.٥٠٦             | ١.٧٤    | تجاهل المراجع لعملية التحقق من أداء أدوات الذكاء الاصطناعي أثناء الفحص لا يحمله مسؤولية قانونية |
| أوافق         | ٠.٢٦٠             | ٤.٠٧    | يظل المراجع مسئولاً قانونياً عن النتائج حتى لو استخدم نظام ذكاء اصطناعي                         |
| أعارض         | ٠.٥٠٦             | ١.٧٤    | في حالة حدوث خطأ ناتج عن النظام الذكي، لا يتحمل المراجع اى مسؤولية القانونية                    |
| أوافق         | ٠.١٧٤             | ٤.٣     | استخدام أدوات غير مناسبة يعرضني للمساءلة القانونية كمراجع                                       |
| أوافق         | ٠.٢٦٠             | ٤.٠٧    | استخدام أدوات الذكاء الاصطناعي سوف لن يعفىني من مسؤولياتي القانونية كمراجع                      |
| أوافق         | ٠.٢٤٨             | ٤.١٣    | اجمالي المحور                                                                                   |

ومن خلال تحليل المتوسطات الحسابية، والانحراف المعياري للاستجابات يتضح أن الاتجاه العام لاستجابات المستقصى منهم على معظم أسئلة الاستبيان الخاصة بالمحور السادس (موافق) وذلك بمتوسط حسابي قدره (٤.١٣) وانحراف معياري قدره (٠.٢٤٨). وهذا ما يثبت صحة الفرض الثالث البديل وهو وجود علاقة معنوية ذات دلالة إحصائية بين مخاطر الإهمال الرقمي عند إعداد أدوات الذكاء الاصطناعي والتزام المراجع بمسئوليته القانونية.

### جدول رقم (١١) المحور السابع: ضعف الشفافية

| الاتجاه العام | الانحراف المعياري | المتوسط | المحور السابع: ضعف الشفافية                                          |
|---------------|-------------------|---------|----------------------------------------------------------------------|
| أوافق         | ٠.٢٩٢             | ٤.٠٩    | غموض الخوارزميات يجعل من الصعب تقييم صحة النتائج                     |
| أعارض         | ٠.٢٠٠             | ١.٩٦    | لا تهمني طريقة عمل النظام طالما أعطى نتائج نهائية                    |
| أوافق         | ٠.٢٩٢             | ٤.٠٩    | يصعب الوثوق بنتائج لا أعرف كيف تم التوصل إليها                       |
| أوافق         | ٠.٣٩١             | ٤.١٩    | إذا لم يفهم المراجع خوارزمية النظام، لا يمكنه استخدام نتائجه دون قلق |
| أوافق         | ٠.٢١٥             | ٤.١٠    | اجمالي المحور                                                        |

ومن خلال تحليل المتوسطات الحسابية، والانحراف المعياري للاستجابات يتضح أن الاتجاه العام لاستجابات المستقصى منهم على معظم أسئلة الاستبيان الخاصة بالمحور السادس (موافق) وذلك بمتوسط حسابي قدره (٤.١٠) وانحراف معياري قدره (٠.٢١٥).

### جدول رقم (١٢) المحور الثامن: ممارسة الحكم المهني

| الاتجاه العام | الانحراف المعياري | المتوسط | المحور الثامن: ممارسة الحكم المهني                                                      |
|---------------|-------------------|---------|-----------------------------------------------------------------------------------------|
| اوافق         | ٠.٢٩٢             | ٤.٠٩    | الخوارزميات لا يمكن الاعتماد عليها فقط بدون الحكم المهني البشري                         |
| اوافق         | ٠.٤٩٩             | ٤.١٩    | المراجع يجب أن يمارس تقديره المهني ولو في ظل وجود نتائج جاهزة من النظام الذكي           |
| اوافق         | ٠.٤٩٩             | ٤.١٩    | ادوات الذكاء لا تقني عن دور المراجع في الحكم والتقييم                                   |
| أعارض         | ٠.٢٠٠             | ١.٩٦    | ضعف الشفافية في نتائج أدوات الذكاء الاصطناعي لا يؤثر على اتخاذ القرار المهني            |
| أوافق         | ٠.٣٩١             | ٤.١٩    | كلما زادت درجة الغموض في نتائج أدوات الذكاء الاصطناعي قلت قدرتي على ممارسة الحكم المهني |
| أوافق         | ٠.٢٣٤             | ٤.١٢    | اجمالي المحور                                                                           |

ومن خلال تحليل المتوسطات الحسابية، والانحراف المعياري للاستجابات يتضح أن الاتجاه العام لاستجابات المستقصى منهم على معظم أسئلة الاستبيان الخاصة بالمحور السادس (موافق) وذلك بمتوسط حسابي قدره (٤.١٤) وانحراف معياري قدره (٠.٢٣٤).

وهذا ما يثبت صحة الفرض الرابع البديل وهو وجود علاقة معنوية ذات دلالة إحصائية بين مخاطر ضعف الشفافية في الخوارزميات الذكية وممارسة المراجع للحكم المهني.

#### ٤- إختبار فروض الدراسة الرئيسية، ومستوى الدلالة المعنوية:

٤/١ الفرض الاول الصفري: لا يوجد تأثير ذو دلالة إحصائية لمخاطر التحيز الخوارزمي داخل أنظمة الذكاء الاصطناعي على مسؤولية المراجع الخارجي عن اكتشاف الأخطاء الجوهرية.

أظهرت نتائج تحليل الانحدار وجود تأثير دال إحصائيًا لمخاطر التحيز الخوارزمي داخل أنظمة الذكاء الاصطناعي على مسؤولية المراجع الخارجي عن اكتشاف الأخطاء الجوهرية، حيث بلغ معامل الارتباط ( $R = 0.780$ )، ومعامل التحديد  $R^2 = 0.608$ ) وبلغ معامل التأثير ( $\beta = 0.876$ ) مع دلالة معنوية مرتفعة ( $p < 0.001$ ). يشير ذلك إلى أن زيادة احتمالات مخاطر التحيز في نتائج أدوات الذكاء الاصطناعي قد تؤدي إلى إضعاف قدرة المراجع على اكتشاف الانحرافات أو الأخطاء، وهو ما يعكس تأثيرًا سلبيًا للمخاطر التقنية على جودة الفحص المهني، وهذا يدعم الفرض الأول البديل، وينفي الفرض الصفري.

٢/٤ الفرض الثاني الصفري: لا يوجد تأثير ذو دلالة إحصائية لمخاطر التواطؤ التقني في برمجة خوارزميات الذكاء الاصطناعي على التزام المراجع بمسؤولية الإفصاح المهني والقانوني.

أظهرت نتائج تحليل الانحدار وجود تأثير دال إحصائيًا لمخاطر التواطؤ التقني في تصميم وبرمجة أدوات الذكاء الاصطناعي على التزام المراجع بالإفصاح المهني والقانوني، حيث بلغ معامل الارتباط ( $R = 0.785$ )، ومعامل التحديد  $R^2 = 0.617$ )، ومعامل التأثير ( $\beta = 0.757$ )، مع دلالة معنوية مرتفعة ( $p < 0.001$ ). وتشير هذه النتائج إلى أنه كلما زاد إدراك المراجع بوجود احتمالات بالتلاعب أو الانحياز أو التدخل غير المهني في برمجة أدوات الذكاء الاصطناعي، كلما زادت مسؤوليته تجاه الإفصاح، وأصبح مطالبًا بالتحقق من مدى نزاهة تصميم النظام، والتغيب عن مصادر النتائج وعدم التسليم بها، والإفصاح بوضوح في تقريره عن أي غموض أو خطر في أدوات الفحص المؤتمتة، وهذا يعزز شعوره بالمسؤولية المهنية، ويدفعه إلى التوسع في الإفصاح والشرح.

وبالتالي فإن الخطر التكنولوجي ينعكس في صورة عبء مهني إضافي يقع على عاتق المراجع لحماية جودة المراجعة والامتثال القانوني، وكل ذلك يدعم الفرض الثاني البديل، ويؤدي إلى رفض الفرض الصفري.

٤/٣ **الفرض الثالث الصفري:** لا يوجد تأثير ذو دلالة إحصائية لمخاطر الإهمال الرقمي عند إعداد أدوات الذكاء الاصطناعي على التزام المراجع بمسئوليته القانونية. أظهرت نتائج تحليل الانحدار وجود تأثير دال إحصائياً لمخاطر الإهمال الرقمي في إعداد أدوات الذكاء الاصطناعي على المسؤولية القانونية للمراجع، حيث بلغ معامل الارتباط ( $R = 0.790$ )، ومعامل التحديد ( $R^2 = 0.623$ )، ومعامل التأثير  $\beta$  ( $1.39 =$ ) عند مستوى دلالة ( $p < 0.001$ ).

وتشير هذه النتائج إلى أن الإهمال في إختبار أو تحديث أدوات الذكاء الاصطناعي يزيد من احتمالية ارتكاب أخطاء مهنية، مما يعرض المراجع لتحمل المسؤولية القانونية، حتى لو كانت الأداة هي المتسبب المباشر، وهذا يدعم الفرض الثالث البديل، وينفي الفرض الصفري.

٤/٤ **الفرض الرابع الصفري:** لا يوجد تأثير ذو دلالة إحصائية لمخاطر ضعف الشفافية في الخوارزميات الذكية على ممارسة المراجع للحكم المهني. أظهرت نتائج تحليل الانحدار وجود تأثير دال إحصائياً لمخاطر ضعف الشفافية في الخوارزميات الذكية على ممارسة المراجع للحكم المهني، حيث بلغ معامل الارتباط ( $R = 0.873$ )، ومعامل التحديد ( $R^2 = 0.763$ )، ومعامل التأثير ( $\beta = 0.948$ )، عند مستوى دلالة ( $p < 0.001$ ).

وتشير هذه النتائج إلى أن زيادة غموض الخوارزميات وآليات عمل النظام الذكي يضعف من قدرة المراجع على إتخاذ قرارات مهنية مستقلة، ويقلل من تطبيقه للحكم والتقدير المهني، وهو عنصر جوهري في جودة المراجعة، وهذا يدعم الفرض الرابع البديل، وينفي الفرض الصفري.

٤/٥ **الفرض الخامس الصفري:** لا تختلف تصورات المراجعين بشأن مخاطر الذكاء الاصطناعي ومسئوليتهم المهنية تبعاً لاختلاف خصائصهم الشخصية والمهنية. تم استخدام تحليل التباين المتعدد MANOVA بدلاً من ANOVA لإختبار الفرض الخامس، حيث يتميز MANOVA بقدرته على فحص التأثير المشترك للمتغيرات المستقلة (الخصائص الشخصية والمهنية للمراجع مثل سنوات الخبرة،

المؤهل، الدرجة الوظيفية، نوع جهة العمل) على متغيرين تابعين مرتبطين وهما: إدراك مخاطر الذكاء الاصطناعي، والمسئولية المهنية والقانونية، مما يوفر قوة إحصائية أكبر، وقد ظهرت النتائج كالآتي:

١/٥/٤: **بالنسبة لمتغير الوظيفة:** أظهرت نتائج تحليل MANOVA وجود فروق ذات دلالة إحصائية في إدراك المراجعين لكل من مخاطر الذكاء الاصطناعي ومسئولياتهم المهنية تبعًا لاختلاف درجتهم الوظيفية، حيث أشارت نتائج اختبار Wilks' Lambda إلى تأثير مشترك دال إحصائيًا للدرجة الوظيفية على إدراك المخاطر والمسئولية (Sig. = 0.026)، بحجم تأثير متوسط  $\text{Partial Eta}^2 = 0.074$ .

كما أظهرت إختبارات التأثير المنفصل فروقًا دالة في إدراك المخاطر (Sig. = 0.010,  $\text{Eta}^2 = 0.067$ )، والمسئولية (Sig. = 0.041,  $\text{Eta}^2 = 0.043$ ) وتوضح المتوسطات أن الإدراك يرتفع مع ارتفاع الدرجة الوظيفية: من مراجع مساعد (3.87) إلى مدير مكتب (4.10)، مما يشير إلى أن الدرجة الوظيفية تلعب دورًا في تشكيل تصورات المراجعين تجاه كل من المخاطر والمسئولية، مما يعكس أثر تراكم الخبرة في تنمية الوعي المهني والتقني لدى المراجعين.

٢/٥/٤: **بالنسبة لمتغير سنوات الخبرة:** أظهرت نتائج تحليل MANOVA وجود فروق ذات دلالة إحصائية في إدراك المراجعين لكل من مخاطر الذكاء الاصطناعي ومسئولياتهم المهنية تبعًا لاختلاف سنوات الخبرة، حيث أشارت نتائج اختبار Wilks' Lambda أن هناك تأثيرًا مشتركًا دالًا إحصائيًا لسنوات الخبرة على كل من إدراك المخاطر والمسئولية المهنية (Sig. = 0.025)، بحجم تأثير متوسط  $\text{Partial Eta}^2 = 0.075$ .

كما أظهرت اختبارات التأثير المنفصل فروقًا دالة في إدراك المخاطر (Sig. = 0.005,  $\text{Eta}^2 = 0.128$ ) والمسئولية المهنية (Sig. = 0.006,  $\text{Eta}^2 = 0.124$ ). وتشير المتوسطات إلى أن الإدراك يرتفع كلما زادت سنوات الخبرة، وأقل متوسط لدى من لديهم 1-5 سنوات (3.88) وأعلى متوسط لمن لديهم أكثر من 15 سنة

(4.06)، مما يشير إلى أن سنوات الخبرة تلعب دورًا في تشكيل تصورات المراجعين تجاه كل من المخاطر والمسئولية.

**٤/٥/٣ بالنسبة لمتغير نوع جهة العمل:** أظهرت نتائج تحليل MANOVA وجود فروق ذات دلالة إحصائية في إدراك المراجعين لكل من مخاطر الذكاء الاصطناعي ومسئولياتهم المهنية تبعًا لاختلاف نوع جهة العمل، حيث أشارت نتائج اختبار Wilks' Lambda أن هناك تأثيرًا مشتركًا دالًا إحصائيًا لنوع جهة العمل على كل من إدراك المخاطر والمسئولية المهنية ( $\text{Sig.} = 0.024$ ) بحجم تأثير متوسط ( $\text{Partial Eta Squared} = 0.076$ ).

كما أظهرت المتوسطات أن أعلى إدراك للمخاطر والمسئولية كان لدى العاملين في شركات المساهمة (4.07) يليهم العاملون في شركات التضامن (4.06)، في حين كان أقل إدراك لدى العاملين في مكاتب المراجعة الخاصة (3.93)، مما يعكس أن طبيعة جهة العمل تلعب دورًا في تشكيل وعي المراجعين بالمخاطر والمسئوليات المتعلقة باستخدام تقنيات الذكاء الاصطناعي في بيئة المراجعة.

**٤/٥/٤ بالنسبة لمتغير المؤهل الدراسي:** أظهرت نتائج اختبار Wilks' Lambda أن التأثير المشترك للمؤهل العلمي على إدراك مخاطر الذكاء الاصطناعي والمسئولية المهنية لم يكن ذا دلالة إحصائية، حيث كانت قيمة الدلالة ( $\text{Sig} = 0.273$ ) وهي أكبر من 0.05، مما يعني عدم وجود فروق ذات دلالة إحصائية عند مستوى 0.05 في التصورات بناءً على المؤهل العلمي.

والنتيجة الإجمالية للفرض الخامس هناك إختلافات جوهرية في إدراك المراجعين للمخاطر والمسئولية المهنية بناءً على الخبرة والوظيفة وجهة العمل، بينما لا يبدو أن المؤهل الأكاديمي يفرق إحصائيًا في هذه العينة، وهذا يدعم الفرض البديل وينفي الفرض الصفري بالنسبة للمتغيرات الشخصية والمهنية التي شملتها الدراسة، فيما عدا المؤهل الدراسي.

## الفصل الخامس: النتائج والتوصيات

### أولاً: أهم النتائج:

١. بعض أنظمة الذكاء الاصطناعي قد تتضمن تحيزات مدمجة ناتجة عن بيانات تدريب غير متوازنة أو برمجة موجهة، مما يحدّ من قدرة المراجع على إكتشاف الأخطاء الجوهرية.
٢. تشير الأدبيات إلى أن المراجعين قد يجدون صعوبة في الإفصاح عن الانحرافات الناتجة عن أدوات ذكية تم توجيهها عمداً أو برمجتها بصورة غير نزيهة، خاصة في حال عدم مشاركتهم في تطويرها.
٣. إستخدام أدوات "الصندوق الأسود" (Black-box AI) "يجعل من الصعب تفسير منطق القرارات، مما يقيّد قدرة المراجع على إصدار أحكام مهنية مستقلة.
٤. غياب المراجعة الفنية المنتظمة لأداء الأدوات الذكية، أو الاعتماد على أدوات غير مُختبرة، يزيد من احتمالية الوقوع في الأخطاء التي قد تعرّض المراجع للمساءلة القانونية.
٥. تختلف تصورات المراجعين عن مخاطر تقنيات الذكاء الاصطناعي ومسئولياتهم المهنية والقانونية باختلاف صفاتهم الشخصية والمهنية.

### ثانياً: التوصيات:

١. ممارسة درجة عالية من العناية المهنية عند استخدام الأدوات الذكية
٢. طلب توضيحات فنية مفصلة حول طريقة عمل النظم المستخدمة في إعداد التقارير المالية
٣. الحفاظ على التقدير المهني المستقل وعدم الاعتماد المطلق على نتائج الذكاء الاصطناعي.
٤. تعزيز وعي المراجعين بمخاطر الذكاء الاصطناعي من خلال التدريب المستمر على الجوانب التقنية لأنظمة الذكاء الاصطناعي المستخدمة في التقارير المالية.
٥. فرض ضوابط مهنية على إستخدام أدوات الذكاء الاصطناعي في المراجعة، خاصة تلك التي تتعلق بتوثيق منطق الخوارزميات ومصدر بيانات التدريب.

٦. التأكيد على ضرورة تدخل المراجع في مراحل تطوير أدوات الذكاء الاصطناعي لضمان ملاءمتها لمتطلبات المهنة وتقليل مخاطر التواطؤ أو الغموض.
٧. وضع معايير مهنية واضحة حول مسؤولية المراجع في بيئة الذكاء الاصطناعي، بما يشمل الحكم المهني، الإفصاح، وحدود المسؤولية القانونية.
٨. إلزام الجهات التنظيمية بتقارير تقييم دورية لأداء الأنظمة الذكية في المراجعة، ومراجعة الضوابط الأمنية والحوكمة التكنولوجية المرافقة لها.

### ثالثاً: مقترحات لدراسات مستقبلية

١. دراسة مدى وعي المراجعين القانونيين في مصر بمخاطر التحيز الخوارزمي.
٢. تحليل أثر التواطؤ التقني على جودة التقارير المالية المعتمدة على الذكاء الاصطناعي.
٣. قياس فعالية تطبيق أنظمة الذكاء الاصطناعي القابلة للتفسير في بيئة المراجعة.
٤. بحث العلاقة بين الخبرة المهنية للمراجع ومدى اعتماده على مخرجات الأدوات الذكية.
٥. تطوير إطار مهني مقترح لتقييم أدوات الذكاء الاصطناعي في ضوء معايير المراجعة الدولية.

## قائمة المراجع

- Amershi, S., Weld, D., Vorvoreanu, M., Fourney, A., Nushi, B., Collisson, P., ... & Horvitz, E. (2019). *Guidelines for Human-AI Interaction*. Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems, 1-13.
- Appelbaum, D. A., Kogan, A., & Vasarhelyi, M. A. (2017). *Big data and analytics in the modern audit engagement: Research needs*. Auditing: A Journal of Practice & Theory, 36(4), 1–27.
- Appelbaum, D., Kogan, A., & Vasarhelyi, M. A. (2020). Design and development of artificial intelligence – A research agenda for auditing. *Journal of Information Systems*, 34(4), 63–74.
- Appelbaum, D., Kogan, A., & Vasarhelyi, M. A. (2022). Artificial intelligence in audit: Current developments and future directions. *Journal of Emerging Technologies in Accounting*.
- Arens, A. A., Elder, R. J., & Beasley, M. S. (2017). *Auditing and assurance services: An integrated approach* (16th ed.). Pearson.
- Barocas, S., Hardt, M., & Narayanan, A. (2020). Fairness and machine learning. *Recommender systems handbook*, 1, 453-459.
- Baskarada, S., Shanks, G., & Draheim, D. (2022). Explainable Artificial Intelligence (XAI) in auditing. *International Journal of Accounting Information Systems*, 46, 100567.
- Binns, R. (2018). Fairness in machine learning: Lessons from political philosophy. *Proceedings of the 2018 Conference on Fairness, Accountability, and Transparency*, 149–159.
- Brynjolfsson, E., & McAfee, A. (2023). *The second machine age: Work, progress, and prosperity in a time of brilliant technologies*. W. W. Norton & Company.
- Burrell, J. (2016). How the machine ‘thinks’: Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), 1–12.

- 
- Coglianesi, C., & Lehr, D. (2019). Transparency and algorithmic governance. *Administrative law review*, 71(1), 1-56.
- Crawford, K., & Paglen, T. (2021). Excavating AI: The politics of images in machine learning training sets. *International Journal of Communication*, 15, 3702–3720.
  - Deng, A. (2020). Algorithmic collusion and algorithmic compliance: Risks and opportunities. *The Global Antitrust Institute Report on the Digital Economy*, 27.
  - Diakopoulos, N. (2016). Accountability in algorithmic decision making. *Communications of the ACM*, 59(2), 56-62.
  - Du, Y., & Xiong, Y. (2025). *Challenges and opportunities of artificial intelligence in auditing: A critical review*. *International Journal of Accounting Information Systems*, 48, 100580.
  - Eilifsen, A., & Messier, W. F. (2022). *Auditing and assurance services in a digital world*.
  - Eubanks, V. (2018). *Automating inequality: How high-tech tools profile, police, and punish the poor*. St. Martin's Press.
  - Ezrachi, A., & Stucke, M. E. (2016). *Virtual competition: The promise and perils of the algorithm-driven economy*. Harvard University Press.
  - Ezrachi, A., & Stucke, M. E. (2017). Artificial intelligence & collusion: When computers inhibit competition. *U. Ill. L. Rev.*, 1775.
  - Falah Alroud, S., Aljabr, M. A., & Al-Shorafa, A. J. (2025). The influence of artificial intelligence on electronic audit evidence: exploring the mediating role of digital transformation: evidence from Jordanian export firms. *EDPACS*, 70(2), 1-40.
  - IAASB. (2020). *Audit evidence*. International Auditing and Assurance Standards Board.
-

- 
- IAASB. (2020). The IAASB's Technology Working Group: Final Report. International Auditing and Assurance Standards Board.
  - IAASB. (2023). *Technology working group: The impact of technology on audit evidence*. International Auditing and Assurance Standards Board.
  - IFAC. (2020). *The role of professional accountants in the digital era*. International Federation of Accountants.
  - IFAC. (2023). *The role of the auditor in the age of AI*. International Federation of Accountants.
  - International Federation of Accountants (IFAC). (2025). *Artificial Intelligence & Accounting*. IFAC Knowledge Gateway. Retrieved from IFAC website
  - International Federation of Accountants, & International Ethics Standards Board for Accountants. (2010). *Code of ethics for professional accountants*. International Federation of Accountants.
  - Issa, H., Sun, T., & Vasarhelyi, M. A. (2016). Research ideas for artificial intelligence in auditing: The formalization of audit and workforce supplementation. *Journal of Emerging Technologies in Accounting*, 13(2), 1–20.
  - Johnson, D. G. (2021). Algorithmic accountability in the making. *Social Philosophy and Policy*, 38(2), 111-127.
  - Kamareldawla, N. M. (2025). External auditors' perceptions toward the use of artificial intelligence in the audit process and ethical challenges. *Journal of Governance & Regulation*, 14(1), 84–96.
  - Kokina, J., & Davenport, T. H. (2017). The emergence of artificial intelligence: How automation is changing auditing. *Journal of Emerging Technologies in Accounting*, 14(1), 115–122. <https://doi.org/10.2308/jeta-51730>

- Kokina, J., Pachamanova, D., & Corbett, A. (2021). The role of data and analytics in auditing: A research synthesis. *Accounting Horizons*.
- Le, H. H. (2025). *Who Should Take Responsibility for Artificial Intelligence Actions and Outcomes? Perception of Auditors As Users of AI Systems* (Doctoral dissertation, Louisiana Tech University).
- Lehner, O. M., Ittonen, K., Silvola, H., Wührleitner, A., & Ström, E. (2022). *Artificial intelligence-based decision-making in accounting and auditing: Ethical challenges and normative thinking*. *Accounting, Auditing & Accountability Journal*.
- Leocádio, D., Malheiro, L., & Reis, J. (2024). Artificial Intelligence in auditing: A conceptual framework for auditing practices. *Administrative Sciences*, 13(6), 136.
- Libby, R., & Witz, P. D. (2024). Can artificial intelligence reduce the effect of independence conflicts on audit firm liability?. *Contemporary Accounting Research*, 41(2), 1346-1375.
- Liu, J., Wang, H., & Zhang, M. (2022). Auditors' judgment under AI systems. *Accounting Research Journal*, 35(2), 88–105.
- Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S., & Floridi, L. (2016). The ethics of algorithms: Mapping the debate. *Big Data & Society*, 3(2), 1–21.
- Mokander, J. (2024). Auditing of AI: Legal, Ethical and Technical Approaches. *Digital Society*, 2(49).
- Moll, J., & Yigitbasioglu, O. (2019). The role of internet-related technologies in shaping the work of accountants: New directions for accounting research. *The British Accounting Review*, 51(6), 100833.
- Müller, V. C., Liao, Q. V., & Singh, A. (2022). Transparency and interpretability in AI systems for auditing. arXiv preprint.

- Murikah, W., Nthenge, J. K., & Musyoka, F. M. (2024). Bias and ethics of AI systems applied in auditing-A systematic review. *Scientific African*, 25, e02281.
- OECD, P. (2017). Algorithms and collusion: competition policy in the digital age. *Algorithms and Collusion: Competition Policy in the Digital Age*. IAASB. (2020). *Audit Evidence – Working Group Publication*. International Auditing and Assurance Standards Board.
- Parker, C. A., Cho, S., & Vasarhelyi, M. (2022). Explainable Artificial Intelligence (XAI) in Auditing. *International Journal of Accounting Information Systems*
- Peters, C. P. H. (2022). *Auditor automation usage and professional skepticism* (Working paper, Tilburg University).
- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation–augmentation paradox. *Academy of Management Review*, 46(1), 192–210.
- Raji, I. D., Binns, R., Shankar, S., & Hanna, A. (2020). Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency (FAT\*)*, 33–44.
- Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). “*Why Should I Trust You?*” *Explaining the Predictions of Any Classifier*. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135–1144.
- Richins, G., Stapleton, A., Stratopoulos, T. C., & Wong, C. (2017). *Big Data Analytics: Opportunity or Threat for the Accounting Profession?* *Journal of Information Systems*, 31(3), 63–79.
- Rozario, A. M., & Vasarhelyi, M. A. (2018). How do audit firms use AI? *The CPA Journal*, 88(6), 60–64.

- Russell, S. J., & Norvig, P. (2021). *Artificial intelligence: A modern approach* (4th ed.). Pearson.
- Scientific African. (2024). *Bias and ethics of AI systems applied in auditing: A systematic review*. Scientific African, 25, e02281.
- Sutton, S. G., Holt, M., & Arnold, V. (2016). The reports of my death are greatly exaggerated—Artificial intelligence research in accounting. *International Journal of Accounting Information Systems*, 22, 60–73.
- Tambe, P., Cappelli, P., & Yakubovich, V. (2019). *Artificial Intelligence in Human Resources Management: Challenges and a Path Forward*. California Management Review, 61(4), 15–42.
- Tjoa, E., & Guan, C. (2020). A survey on explainable artificial intelligence (XAI): Toward medical XAI. *IEEE Transactions on Neural Networks and Learning Systems*, 32(11), 4793–4813.
- Verma, S. (2019). Weapons of math destruction: how big data increases inequality and threatens democracy. *Vikalpa*, 44(2), 97–98.
- Yoon, K., Hoogduin, L., & Zhang, L. (2015). Big Data as complementary audit evidence. *Accounting Horizons*, 29(2), 431–438.
- Yoon, K., Hoogduin, L., & Zhang, L. (2021). The impact of AI on external audit. *Journal of Accounting Literature*, 46, 1–25.
- Zhang, Y., Dai, J., & Vasarhelyi, M. A. (2022). *Explainable Artificial Intelligence (XAI) in auditing: A research agenda*. International Journal of Accounting Information Systems, 44, 100565.
- Zhang, Y., Li, J., & Zhou, X. (2024). Deep learning and audit responsibility: A cognitive analysis. *Journal of Accounting Research*, 62(1), 30–55.