

كلمة العدد

التنقل في المشهد الأخلاقي لاستخدام الذكاء الاصطناعي التوليدي

(GenAI) في النشر العلمي: نهج حذر لضمان النزاهة

محمد نجيب المرزوق

أبرار رضا الحسن

جامعة الكويت، الكويت

في السنوات الأخيرة، شهد النشر العلمي بروز الذكاء الاصطناعي التوليدي والتكنولوجيات المدعومة بالذكاء الاصطناعي. ولقد أحدثت هذه الأدوات تحولاً في مراحل مختلفة من عملية البحث والكتابة، مما مكّن الباحثين من تحسين انسيابية سير العمل، وتحسين جودة المخطوطات، وزيادة التركيز على إنشاء المحتوى العلمي الموضوعي (Holmes et al., 2022). ويمثل هذا الأمر حقيقة واقعة يجدر بالباحثين والمراجعين والمحررين تقبّلها. وعلى الرغم من ذلك، لم يصل مجتمع البحث العلمي بعد إلى نقطة الاستيعاب الكامل بشأن الاستخدام الصحيح لهذه الأدوات (Nguyen et al., 2023). وإننا نأمل أن تساهم كلمة العدد هذه في هذه المناقشة المتواصلة، وأن نقرب أكثر من التوصل إلى توافق حول الاستخدام المقبول للذكاء الاصطناعي التوليدي في كتابة الأبحاث العلمية ونشرها.

يمكن للذكاء الاصطناعي التوليدي تحسين العديد من المهام التي يؤديها محررو الأبحاث العلمية أو مراجعوها أو مؤلفوها، فهو يمتلك القدرة على إنشاء التلخيصات والملخصات (Abstracts)، والتحقق من الأخطاء الموجودة في النص وتصحيحها، وتحليل البيانات، وإنتاج التصورات البصرية، بالإضافة إلى المساعدة في تحديد الأدبيات ذات الصلة وإدارة الاقتباسات (Flores-Vivar & García-Peñalvo, 2023). ورغم استخدامنا العديد من الأدوات في الماضي للمساعدة في تحسين هذه المهام، إلا أن ما يميز الذكاء الاصطناعي التوليدي هو طبيعته التوليدية.

ولكن من ناحية ثانية، تفتقر أدوات الذكاء الاصطناعي التوليدي إلى القدرة على المساهمة في "النظرية العميقة"، أي صياغة الحجج الأصلية وربط الأفكار بأساليب جديدة، فبرغم قدرة الذكاء الاصطناعي التوليدي على إعادة الصياغة وتنظيم المحتوى القائم بكفاءة، فإنه يكرّر ما هو معروف بالفعل بدلاً من تقديم رؤى مبتكرة (Dobrin, 2023; Mabus, 2024). وقد يفضي هذا الاعتماد على الذكاء الاصطناعي إلى الحد من عمق الأعمال العلمية وأصالتها، مما يقوّض الإسهامات المبتكرة الضرورية للتطور العلمي. وهنا يكمن التحدي الأخلاقي الجديد الذي تطرحه هذه الأدوات: هل يمكن لشخص ما استخدام الذكاء الاصطناعي التوليدي لإنتاج محتوى يدعي صحته؟ وما هي حدود استخدام الذكاء الاصطناعي التوليدي في هذا السياق؟ وما هي الاعتبارات الأخلاقية المترتبة على ذلك؟

أمّا التحدي الآخر فيتعلّق باستخدام الذكاء الاصطناعي التوليدي في البحث العلمي بجودة المخرجات المنتجة التي قد تتسبب في تقويض رصانة الأبحاث العلمية وقابلية تكرارها، فمن

المعروف أنّ الذكاء الاصطناعيّ التوليديّ قد ينتج عنه أحياناً معلومات غير صحيحة ومضلّة لكنّها متماسكة، وتبدو مقنعة (Ferrara, 2024). وعلاوةً على ذلك، فإنّ عمليات الاستبّاط التي يجريها الذكاء الاصطناعيّ التوليديّ ليست متسقة؛ إذ إنّها تعتمد على الاستدلال الإحصائيّ المستند إلى بيانات التدريب عوضاً عن التفكير المنطقيّ الحقيقيّ غير المصطنع (Mirzadeh et al., 2024).

وبالنظر إلى هذه الخصائص التي تتسم بها النتائج الصادرة عن الذكاء الاصطناعيّ التوليديّ، ومع افتراض الاستخدام الأخلاقيّ له، هل يمكن للمستخدم أن يصادق على صحة مخرجاته أو يتقّبها؟ هل تتمتع هذه المخرجات بالموثوقيّة الكافية بحيث يمكن للباحثين الآخرين استخدامها لإنتاج نتائج متسقة؟ وهل النتائج البالغة الأهمية التي تُرصد في إحدى الدراسات هي مجرد نتائج زائفة تولّدت إثر استخدام الذكاء الاصطناعيّ التوليديّ، أم أنّها تعكس بالفعل الظاهرة قيد الدراسة؟

إذن، كيف يمكن للمؤلف أو المحرر أو المراجع تحديد ما إذا كانت الطريقة التي يستخدمون بها الذكاء الاصطناعيّ تحظى بالقبول؟ يمكننا تصنيف التساؤلات التي نعتقد أنّ أيّ شخص يستخدم الذكاء الاصطناعيّ التوليديّ لإنتاج الأبحاث العلميّة أو تحريرها سيطرحها ضمن فئتين عريضتين؛ حيث تشمل الفئة الأولى الأسئلة التي تبحث فيما إذا كان استخدام الذكاء الاصطناعيّ لأداء مهمة معيّنة أمراً مقبولاً أو أخلاقياً، أمّا الفئة الثانية فتتعلق بالمخاطر التي قد يشكّلها استخدام الذكاء الاصطناعيّ على جودة العمليّة العلميّة وصحتها.

وبناءً على هاتين الفئتين، نقترح نهجاً استرشادياً من سؤاليّن يمكن للمحررين والمراجعين والباحثين الاستفادة منهما لتحديد ما إذا كانت المهمة التي يستعينون فيها بالذكاء الاصطناعيّ لأدائها مقبولة. وهذان السؤالان هما:

- 1 - هل سيكون من المقبول أخلاقياً الاستعانة بمصادر خارجيّة لأداء المهمّة؟
- 2 - هل يمكنك التحقق من جودة هذه المهمّة؟

وإذ إنّنا لا ندّعي أنّ هذه الأسئلة ستقدم الحلول لجميع المشكلات المتعلّقة باستخدام الذكاء الاصطناعيّ، لكننا نعتقد أنّها تمثّل نهجاً استرشادياً ممتازاً لتحديد ما إذا كان الاستخدام أخلاقياً، وما إذا كان يمكنك التحكم في الجودة إلى المستوى الذي يُمكن المستخدم من التحقق من جودة عمل الذكاء الاصطناعيّ التوليديّ. وحين تكون الإجابة على أحد هذين السؤالين "لا" فإنه يمكننا الإشارة إلى استخدام الذكاء الاصطناعيّ التوليديّ على أنّه "استخدامٌ كسولٌ للذكاء الاصطناعيّ". وبافتراض أسوأ سيناريوهات الاستخدام الكسول للذكاء الاصطناعيّ التوليديّ، سيُعدّ المستخدمون أنّهم قد ارتكبوا سرقةً أدبيّةً أو أخطاءً تؤدّي إلى سحب النشر. وعلى أفضل تقدير، فإنّ هذا سيُفضي على الأرجح إلى تخفيض مهارات المستخدم.

وفي حين أنّ السؤال المطروح حول الاستعانة بمصادر خارجيّة لأداء المهمّة يمثّل استجابةً ثنائيّةً واضحة لا لبس فيها في معظم الحالات، فإن العبء الذي يقع على عاتق المستخدم تجاه

الذكاء الاصطناعيّ التوليديّ يتمثّل في إثبات جودة النتائج المنتجة، وهو ما قد يكون أكثر سهولةً في بعض حالات الاستخدام مقارنةً بأخرى.

ولننظر في بعض السيناريوهات المختلفة التي يمكن أن توضح لنا بالأمثلة هذه النقطة، ففيما يتعلق بالمؤلفين، قد تكون السيناريوهات التي يمكنهم فيها تحقيق الاستفادة من الذكاء الاصطناعيّ التوليديّ على النحو الآتي:

- السيناريو 1: الطلب من الذكاء الاصطناعيّ التوليديّ كتابة مقال أو قسم أو فقرة.
- السيناريو 2: الطلب من الذكاء الاصطناعيّ التوليديّ تحسين مقال أو قسم أو فقرة.
- السيناريو 3: الطلب من الذكاء الاصطناعيّ التوليديّ تحليل البيانات أو تصوّرها بصرياً.
- السيناريو 4: الطلب من الذكاء الاصطناعيّ التوليديّ كتابة برنامج لتحليل البيانات وتصورها بصرياً. وفيما يتعلق بالمراجعين، فيمكن الأخذ في الاعتبار السيناريوهات الآتية:
- السيناريو 5: الطلب من الذكاء الاصطناعيّ التوليديّ إجراء مراجعة لإحدى الدراسات.
- السيناريو 6: الطلب من الذكاء الاصطناعيّ التوليديّ تنظيم ملاحظات المراجعة في تقرير متماسك.
- السيناريو 7: الطلب من الذكاء الاصطناعيّ التوليديّ تلخيص الدراسة أو العثور على تناقضات فيها.
- السيناريو 8: الطلب من الذكاء الاصطناعيّ التوليديّ تحسين لغة تقرير المراجعة لزيادة وضوحها أو رفع مستوى عنصر التعاطف بها.
- وفيما يتعلق بالمحرّرين، يمكن الأخذ في الاعتبار السيناريوهات الآتية:
- السيناريو 9: الطلب من الذكاء الاصطناعيّ التوليديّ فحص الطلبات المقدّمة.
- السيناريو 10: الطلب من الذكاء الاصطناعيّ التوليديّ اتخاذ قرار تحريريّ بناءً على تقارير المراجعين.
- السيناريو 11: الطلب من الذكاء الاصطناعيّ التوليديّ تلخيص تقارير المراجعين أو تحديد القواسم المشتركة أو التناقضات عبر التقارير.
- السيناريو 12: الطلب من الذكاء الاصطناعيّ التوليديّ كتابة خطاب ردّ تحريريّ بناءً على ملاحظات المحرّر وقراره.
- السيناريو 13: الطلب من الذكاء الاصطناعيّ التوليديّ كتابة مقال افتتاحيّ.

الملاحظات	هل يتم التحقق من الجودة؟	هل يمكن الاستعانة بمصادر خارجية؟	المؤلف	السيناريو
	صعب	لا	المؤلف	1
يعمل الذكاء الاصطناعيّ التوليديّ كمحرر نسخ.	يختلف حسب الموقف	نعم	المؤلف	2
كلما كبر حجم النص المُتولد، ازدادت صعوبة التحقق من صحته. يفضل إنتاج الجمل والفقرات بدلاً من الأقسام.	صعب	نعم	المؤلف	3
- يعمل الذكاء الاصطناعيّ التوليديّ كأداة للتحليل الإحصائي. - ستتطلب عملية التحقق من الصحة الوسائل الكفيلة لتقييم الثقة في النتائج مثل الموثوقية بين المقيمين (IRR).	سهل	نعم	المؤلف	4
- الاستعانة بالذكاء الاصطناعيّ التوليديّ لكتابة التعليمات البرمجية (الكود). - تمثّل عملية إثبات صحة التعليمات البرمجية المصدريّة (الكود المصدري) أمراً أكثر حتميّة، ويمكن ضمان وإرساء الصحة عن طريق إجراء الاختبارات.	صعب	لا	المراجع	5
- مماثل للسيناريو 1 بالنسبة للمؤلفين. - يتطلب المقارنة بالمراجعين البشريين، فلماذا نفعل ذلك؟ - حتى بعد التدريب على مهمة المراجعة، من المرجح أن يتم استبعاد الأفكار المبتكرة أو تجاهل سوء ممارسات البحث أو المنطق. - يؤدي ذلك إلى تخفيض مهارات المراجعين.	صعب	نعم	المراجع	6
- على الرغم من أنّ هذا الأمر مقبول، إلا أنه من المرجح أن يقلل الثقة في المراجع، حيث يمكن للمحررين استشعار التقارير التي كتبها الذكاء الاصطناعيّ التوليدي. - أداة فحص تساعد المراجع على إيجاد المشكلات التي يمكن له متابعتها لتأكيدّها أو نفيها. - وقوع عبء التحقق من الصحة على المراجع. - يمكن تحسين عمق المراجعات إذا تم استخدامه على النحو الصحيح.	سهل	نعم	المراجع	7

السيناريو	المستلزم	هل يمكن الاستعانة بمصادر خارجية؟	هل يتم التحقق من الجودة؟	الملاحظات
8	المراجع	نعم	سهل	• أفضل حالة استخدام للذكاء الاصطناعي التوليدي بالنسبة للمراجعين. • تحسين اللغة وعنصر التعاطف في التقارير. • انظر الملاحظات الخاصة بالسيناريو 2.
9	المحرر	نعم	يختلف حسب الموقف	• يمكن للذكاء الاصطناعي التوليدي استبعاد الأبحاث المبتكرة أو قبول الأبحاث الرديئة على سبيل الخطأ. • يفيد في تعزيز قرارات الفحص التي اتخذها المحرر، وليس استبدال المحرر.
10	المحرر	لا	صعب	• نظراً لاختلاف درجات جودة تقارير المراجعين وتضارب التقارير عموماً، فمن المرجح أن يتخذ الذكاء الاصطناعي التوليدي قراراً عشوائياً بدلاً من اتخاذ قرار مدروس بناءً على التقارير (Mirzadeh et al., 2024). • بصفتهم أمناء على مجال عملهم ولتحقيق المصالح الفضلى لأصحاب المصلحة الذين يقدمون لهم الخدمات، يتولى المحررون مسؤولية المشاركة في القرار.
11	المحرر	نعم	سهل	• إحدى أفضل حالات الاستخدام للذكاء الاصطناعي التوليدي فيما يتعلق بتعزيز قرارات المحرر.
12	المحرر	نعم	سهل	• حالة استخدام جيدة أخرى، لأن الرسائل قصيرة نسبياً وتستند إلى ملاحظات المحرر، وإثبات صحتها سهل نسبياً.
13	المحرر	لا	صعب	• مماثل للسيناريو 1.

تُبين السيناريوهات المقدمة للمؤلفين والمراجعين والمحررين أنه في حين أن بعض استخدامات الذكاء الاصطناعي التوليدي أكثر قبولاً من الناحية الأخلاقية (مثل تحسين اللغة أو تلخيص النتائج)، إلا أن هناك استخدامات أخرى محفوفة بمخاطر كبيرة.

فمثال ذلك، قد يكون الطلب من الذكاء الاصطناعي التوليدي تحسين قسم من مقال (السيناريو 2) أو تنظيم ملاحظات المراجعة (السيناريو 6) أمراً مقبولاً إذا كُشف عن ذلك على نحو صحيح، ولكن استخدامه لأداء المراجعة بالكامل (السيناريو 5) أو اتخاذ قرارات تحريرية (السيناريو 10) قد يقوض نزاهة العملية.

وبالنسبة للمؤلفين، فإنه يصعب تبرير عملية إنشاء محتوى جديد باستخدام الذكاء الاصطناعي التوليدي (السيناريو 1) دون إشراف قوي وإثبات حقيقي لصحة المحتوى، حيث تشكل عملية إرساء أصالة العمل وجودته أمراً بالغ الصعوبة. ومع ذلك، قد يكون تحسين المحتوى (السيناريو 2) أو استخدام الذكاء الاصطناعي التوليدي في مهام كتابة التعليمات البرمجية (السيناريو 4) أكثر قبولاً، شريطة أن يثبت المؤلف البشري صحة العمل ويتمكن من ضمان جودته والشهادة بها.

أما بالنسبة للمراجعين، فإن مطالبة الذكاء الاصطناعي التوليدي بإجراء مراجعة (السيناريو 5) يمثل مسألة جدالية من الناحية الأخلاقية؛ لأنه قد يتجاهل عناصر بحثية جوهرية أو أفكاراً مبتكرة. ورغم ذلك، يتيح استخدام الذكاء الاصطناعي التوليدي لتحسين المراجعة (السيناريو 7) عن طريق تحديد التناقضات أو تحسين اللغة (السيناريو 8) قابلية أكبر لدعمه وتبريره؛ لأن المراجع لا يزال منخرطاً على نحو جوهري في العملية.

وبالنسبة للمحررين، فبينما يحظى الذكاء الاصطناعي التوليدي بإمكانية المساعدة في تلخيص المراجعات (السيناريو 11) أو صياغة خطابات الرد (السيناريو 12)، فإن استخدامه لاتخاذ قرارات تحريرية نهائية (السيناريو 10) قد يعرض نزاهة العملية التحريرية لأخطار لا يستهان بها، فعلى سبيل المثال، هل سيكون الذكاء الاصطناعي التوليدي قادراً على إصدار أحكام جيدة على جودة المراجعات التي يستندون إليها في قراراتهم؟ هل يمكن للذكاء الاصطناعي التوليدي اتباع الحجج المنطقية التي قدمها المراجعون في حالة وجود تقارير مراجعة متضاربة للوصول إلى قرار غير متحيز ومنطقي؟ وبالنظر إلى دور المحررين كأمناء على مجال عملهم، فإنه يقع على عاتقهم واجب المشاركة في العملية التحريرية. وبينما يمكن للذكاء الاصطناعي التوليدي أن يمثل أداة ذات جدوى لفحص الطلبات المقدمة وتقارير المراجعة، إلا أنه لا ينبغي له أبداً أن يحل محل المحرر.

ورغم أن هذا الإطار يتيح الأساس لتقييم حالات الاستخدام الأخلاقية، فإنه لا يشمل جميع السيناريوهات المحتملة. ولا تزال هناك جوانب غير واضحة تتطلب حكماً دقيقاً ومناقشة مفتوحة بين المؤلفين والمراجعين والمحررين؛ فمثلاً، يفترض النموذج أن نظام الذكاء الاصطناعي التوليدي المستخدم تم تدريبه على بيانات جمعت من المصادر على النحو الصحيح وحصل عليها حسب المعايير الأخلاقية، ولكن إذا لم يتحقق هذا الأمر، فقد يُشكك في العملية برمتها.

ويتجلى الدور الأساسي للشفافية في استخدام الذكاء الاصطناعي التوليدي خلال عملية النشر العلمي، حيث إنه من الضروري أن يكشف المؤلفون والمراجعون والمحررون صراحةً عن أي استخدام لأدوات الذكاء الاصطناعي التوليدي لتمكين أصحاب المصلحة من تقييم دوره في عملية النشر وضمان المساءلة. وربما يتسبب عدم الكشف عن هذا الاستخدام في انتهاكات أخلاقية، ويلقي بظلال من الشك على نزاهة المعايير العلمية والتحريرية.

يحدث "الاستخدام الكسول للذكاء الاصطناعي التوليدي" حين يفرض المستخدمون في الاعتماد على الذكاء الاصطناعي التوليدي دون التحقق من جودة المخرجات أو إدارتها، مما يعرضهم لخطر السرقة الأدبية وسحب النشر وتخفيض المهارات. ويشكل هذا الاعتماد المفرط تهديداً على

التفكير النقدي ومهارات المراجعة لدى الجيل القادم، حيث يزيد المحتوى الذي يتم إنشاؤه دون إشراف بشري من خطر السرقة الأدبية. إضافةً إلى أنه يقلص من دور المهارات الأساسية في الكتابة والتحليل وصنع القرار. ويعتمد النشر العلمي على الرصانة الفكرية والثقة اللتان تتعرضان للخطر إثر الاستخدام المفرط للذكاء الاصطناعي التوليدي دون عملية إثبات الصحة. ويجدر بالمؤلفين والمراجعين والمحررين أن يستمروا في المشاركة بنشاط بغيّة التمسك بجودة العمل المنشور ونزاهته، وتزويد الباحثين في المستقبل بالمهارات اللازمة لتطوير مجالات العمل الخاصة بهم.

وختاماً، فإن هذا الإطار، وإن لم يكن شاملاً، لكنه يقدم نقطة انطلاق تساعد على تقييم الاستخدام المقبول للذكاء الاصطناعي التوليدي في النشر العلمي. وهو يسلط الضوء على المجالات التي يمكن فيها إرساء الثقة في مخرجات الذكاء الاصطناعي التوليدي وتتطلب توشي العناية الفائقة. ورغم ما سبق، فإنه في المواقف التي لا يتناولها هذا الإطار على نحو مباشر ينبغي على مستخدمي الذكاء الاصطناعي التوليدي توشي الحذر والكشف عن الأساليب التي اتبعوها بشفافية والاستعداد لإجراء المزيد من المناقشة. وسيطلب الاستخدام الأخلاقي للذكاء الاصطناعي التوليدي التفحص المستمر لضمان أن البيانات والمخرجات المتولدة تتوافق مع معايير المجتمع العلمي.

المراجع

- Dobrin, S. I. (2023). *AI and writing*. Broadview Press.
- Ferrara, E. (2024). GenAI against humanity: Nefarious applications of generative artificial intelligence and large language models. *Journal of Computational Social Science*, 1-21.
- Flores-Vivar, J.-M., & García-Peñalvo, F.-J. (2023). Reflections on the ethics, potential, and challenges of artificial intelligence in the framework of quality education (SDG4). *Comunicar*, 31(74), 37-47.
- Holmes, W., Persson, J., Chounta, I.-A., Wasson, B., & Dimitrova, V. (2022). *Artificial intelligence and education: A critical view through the lens of human rights, democracy and the rule of law*. Council of Europe.
- Mabus, A. H. (2024). *Teaching community college composition students how to use AI technology ethically through the research process*.
- Mirzadeh, I., Alizadeh, K., Shahrokhi, H., Tuzel, O., Bengio, S., & Farajtabar, M. (2024). *GSM-Symbolic: Understanding the limitations of mathematical reasoning in large language models* (arXiv:2410.05229). arXiv. <https://arxiv.org/abs/2410.05229>
- Nguyen, A., Ngo, H. N., Hong, Y., Dang, B., & Nguyen, B.-P. T. (2023). Ethical principles for artificial intelligence in education. *Education and Information Technologies*, 28(4), 4221-4241.