

Editorial

From Plagiarism to AI-Generated Text, An Editor's Struggle to Ensure Originality

Abdalla Salem Elhwuegi¹ 

¹ Chief Editor, Libyan International Medical University Journal, Benghazi, Libya

Libyan Int Medical Univ J 2025;10:47–48.

Long before the discovery of artificial intelligence (AI), academia assured the originality of a research article by the use of plagiarism detectors. AI has been introduced in 2022 with different functions including creating a scientific text. This necessitates the development of detectors that can detect if a specific text is human or AI generated.

Unfortunately, most of today's available AI detectors have the disadvantages of producing a false positive result, a state that overwhelmed early plagiarism detectors. Just as early plagiarism software identified citations as plagiarism, AI detectors now mistake concise phrasing or technical terms as AI-written. Writers have to change their style to avoid

algorithmic scrutiny by AI detectors, sacrificing clarity for the sake of "human-like" metrics.

As a chief editor committed to ensuring the originality of submitted articles, I have evaluated several publicly available AI detection tools. These tools aim to identify whether a piece of text has been generated by AI models. However, I found that they produced a noticeable inconsistent result when applied to the same content. As a case study, I tested the abstract of a peer-reviewed article published in 2012¹—well before the introduction of generative AI—using four different detectors. The results were astonishing: the AI-generated probability scores ranged from 0 to 89% (see ►Fig. 1).

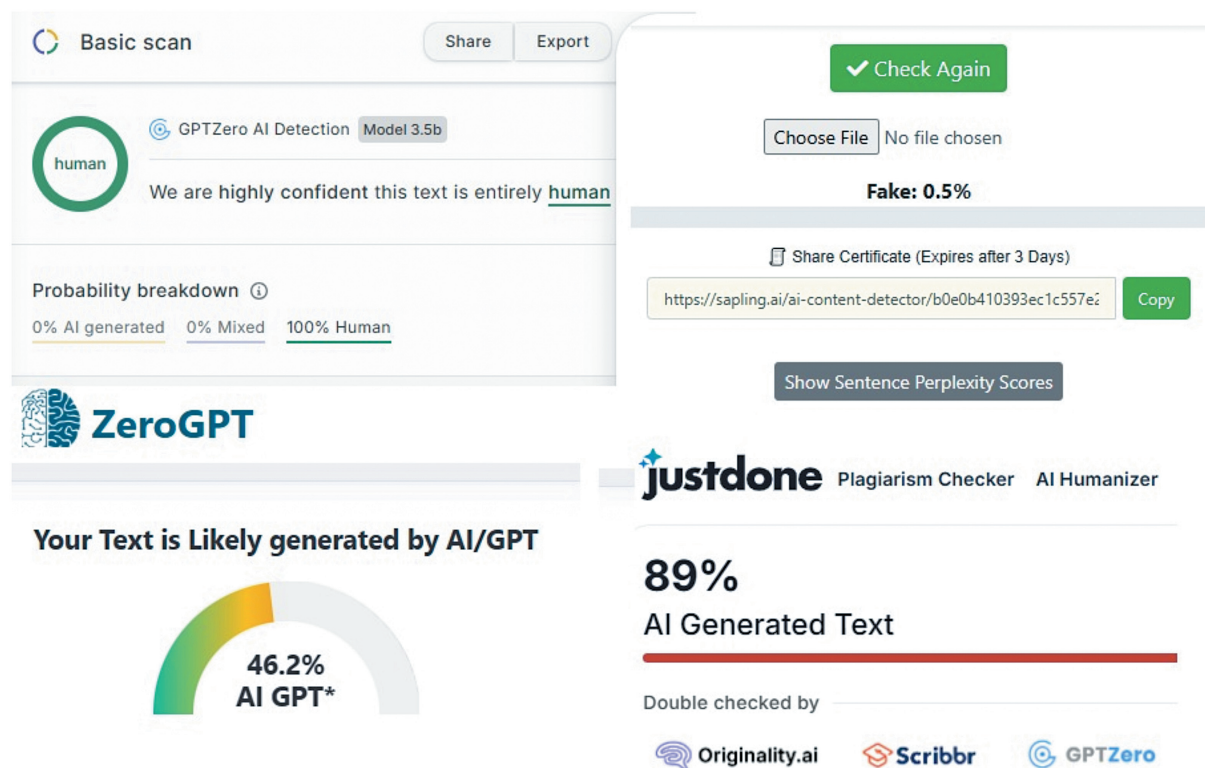


Fig. 1 The results of artificial intelligence (AI) detection scores of the same article using different AI detectors.

Address for correspondence
Abdalla Salem Elhwuegi, PhD,
Emeritus Professor of
Pharmacology, Faculty of
Pharmacy, University of Tripoli,
Tripoli, Libya
(e-mail: hwuegi@hotmail.com).

DOI <https://doi.org/10.1055/s-0045-1810418>.
ISSN 2519-139X.

© 2025. The Author(s).

This is an open access article published by Thieme under the terms of the Creative Commons Attribution License, permitting unrestricted use, distribution, and reproduction so long as the original work is properly cited. (<https://creativecommons.org/licenses/by/4.0/>)
Thieme Medical and Scientific Publishers Pvt. Ltd., A-12, 2nd Floor, Sector 2, Noida-201301 UP, India

We all believe that original research was never about the novelty of the text, it is the product of rigorous inquiry, appropriate insight, and intellectual synthesis, qualities that no present AI detectors can measure. Plagiarism software failed to capture this logic, and AI detectors repeat the same mistake. They cannot distinguish between AI-assisted drafting and AI-generated thought, just as old plagiarism tools couldn't separate boilerplate text from conceptual theft.

I think that the solution is not by developing better detectors, but by reinforcing human judgment. Just as we once learned from plagiarism reports that it is not a verdict

but a starting point for evaluation, we must approach AI results with similar logic. The tools change by the advances of technology, but the challenge remains: to promote authentic scholarship without resorting completely to machine judgment. If we learned anything from the plagiarism era, it is that originality flourishes in discussion, not surveillance. Let's not repeat history's mistakes.

Reference

- 1 Elhwuegi AS, Darez AA, Langa AM, Bashaga NA. Cross-sectional pilot study about the health status of diabetic patients in city of Misurata, Libya. Afr Health Sci 2012;12(01):81–86

المقالة باللغة العربية

من الانتحال إلى النصوص المولدة بالذكاء الاصطناعي، يكافح المحرر لضمان الأصالة

المؤلف: عبدالله سالم الهويجي، رئيس تحرير مجلة الجامعة الليبية الدولية للعلوم الطبية، بنغازي، ليبيا، بريد إلكتروني: hwuegi@hotmail.com

قبل اكتشاف الذكاء الاصطناعي بوقت طويل، كانت الأوساط الأكاديمية تتأكد من أصالة المقالات البحثية باستخدام كاشفات الانتحال. وقد اكتشف الذكاء الاصطناعي عام ٢٠٢٢ بوظائف مختلفة، منها إنشاء نص علي. وهذا يستلزم تطوير كاشفات تمكن من اكتشاف ما إذا كان النص مُولداً من قبل البشر أم من قبل الذكاء الاصطناعي.

للأسف، تعاني معظم أدوات كشف الذكاء الاصطناعي المتاحة اليوم من مشكلة إنتاج نتائج إيجابية خاطئة، وهي إشكالية شبيهة بما واجهته أدوات كشف الانتحال في بداياتها. فكما كانت أنظمة كشف الانتحال تُصنّف الاستشهادات الأكاديمية على أنها انتحال، تُخطئ أدوات كشف الذكاء الاصطناعي اليوم في تفسير الأساليب المختصرة أو استخدام المصطلحات التقنية على أنها ناتجة عن الذكاء الاصطناعي. ونتيجة لذلك، يُضطر الكتاب إلى تعديل أساليبهم في الكتابة لتجنب الوقوع تحت التدقيق الخوارزمي، ولو على حساب الوضوح، سعياً لمواءمة معايير "الكتابة الشبيهة بالبشر".

بحكم مسؤوليتي كمحرر رئيسي حرص على ضمان أصالة المقالات المقدمة، قررت تجربة عدد من أدوات كشف الذكاء الاصطناعي المتاحة للجمهور، لأتحقق من مدى موثوقيتها. تهدف هذه الأدوات إلى تحديد ما إذا كانت النصوص قد أنشئت بواسطة نماذج ذكاء اصطناعي، لكن ما لاحظته كان مفاجئاً: النتائج جاءت متباينة بشكل واضح حتى عند تحليل نفس المحتوى. في تجربة عملية، استخدمتُ ملخصاً لمقال علي خضع للتحكيم نُشر عام 2012 (1)، أي قبل زمن طويل من ظهور الذكاء الاصطناعي التوليدي، وأجريت عليه اختبارات بواسطة أربعة أدوات كشف مختلفة. وكانت النتائج مذهلة فعلاً؛ حيث تراوحت نسب "احتمالية إنشاء بواسطة الذكاء الاصطناعي" من 0/0 إلى 89/0 (انظر الشكل 1)، هذا التفاوت يطرح تساؤلات جادة حول مدى دقة هذه الأدوات، خاصة عند استخدامها في السياقات الأكاديمية.

نحن نعلم أن جوهر البحث الأصلي لا يكمن في حداثة اللغة، بل في عمق الاستقصاء، ووضوح البصيرة، وتماسك الفكر. وهذه عناصر لا تستطيع أدوات كشف الذكاء الاصطناعي الحالية قياسها بأي شكل من الأشكال. لقد فشلت برامج كشف الانتحال سابقاً في إدراك هذا المفهوم، والآن تكرر أدوات الذكاء الاصطناعي نفس الخطأ، فهي لا تفرق بين الأسلوب المُصاغ بمساعدة الذكاء الاصطناعي وبين الأفكار المولدة منه، تماماً كما عجزت أدوات الانتحال في الماضي عن التمييز بين اللغة المتكررة والسرقة الفكرية.

وبرأيي، فإن الحل لا يتمثل في تطوير أدوات كشف أكثر دقة، بل في تعزيز الدور البشري في التقييم. كما تعلمنا من تجربة أدوات الانتحال، فإن نتائجها لم تكن أبداً حكماً نهائياً، بل نقطة انطلاق للمراجعة النقدية. ويجب أن نعامل نتائج كشف الذكاء الاصطناعي بذات المنطق: كمدخل للنقاش، لا كأداة للحسم.

قد تتغير الأدوات بتطور التقنية، لكن التحدي الحقيقي سيبقى: كيف نحني أصالة البحث العلمي من دون أن نخضع بالكامل للآلات؟ وإذا كان لنا أن نتعلم شيئاً من تجربة الانتحال، فهو أن الأصالة لا تُكتشف بالمراقبة، بل تُولد من الحوار. فلنتجنب تكرار أخطاء الماضي، ولنمنح الحكم البشري المساحة التي يستحقها.