

Generative artificial intelligence and human language: The sociology of algorithmic suspicion of tools of plagiarism between bias and delusion of accuracy, and penetration of human language.

Abdelilah Farah^{*1}, Mohammed Bazi¹, Sabir Ben Daoud², Abdellatif Talbi¹, Hassan Habrane^{2,3}

¹ PhD researcher in sociology at Ibn Tofail University, Kenitra, Morocco.

² Doctorate in sociology, Ibn Tofail University, Kenitra, Morocco.

³ Professor of sociology at the Social Services Center – Al Akhawayn University, Ifrane.

*Corresponding Author: abdelilah.farah89@gmail.com

الذكاء الاصطناعي التوليدي واللغة البشرية: سوسيولوجيا الشك الخوارزمي لأدوات الانتحال بين التحيز ووهم الدقة، واختراق اللغة البشرية.

عبد الإله فرح^{1*}، محمد بزي¹، صابر بن داود²، عبد اللطيف طالبي¹، حسن حبران^{2,3}

¹ طالب باحث في سلك الدكتوراه، تخصص علم الاجتماع بجامعة ابن طفيل بالقنيطرة، المغرب.

² دكتوراه في علم الاجتماع، من جامعة ابن طفيل بالقنيطرة، المغرب.

³ أستاذ السوسيولوجيا بمركز الخدمات الاجتماعية - جامعة الأخوين - افران.

* المؤلف المراسل: abdelilah.farah89@gmail.com

Received: 27 -9- 2025

Accented: 14- 11- 2025

تاريخ الاستلام: 27- 9- 2025 تاريخ القبول: 14- 11- 2025

DOI: <https://doi.org/10.48185/sjhss.v1i4.1824>

ISSN (online): 3080-1648

المخلص:

الدراسة تنتمي إلى تخصص سوسيولوجيا الخوارزميات، وهي دراسة نوعية تبحث في تأثير برامج تمييز النصوص المولدة بالذكاء الاصطناعي، وترى في هذه البرامج بأنها لا تهدف إلى حماية البحث العلمي وإقامة العدل بين الأكاديميين بقدر ما أنها تعيد إنتاج لغة وثقافة مهيمنة، إذ إنها متأثرة بأشكال من التحيزات اللغوية والعرقية التي تؤدي إلى إقصاء فئة واسعة من الناطقين غير الأصليين باللغة الإنجليزية، ولاسيما من الأصول العرقية كالسود واللاتينيين. وأن الثقة في خوارزميات برامج الكشف عن محتوى الذكاء الاصطناعي إنما هي وهم تتوارى خلفه مصالح اقتصادية ورمزية، بالمقابل تؤكد على أن اللغة البشرية تتأثر بالهايتوس الخوارزمي الناتج عن التنشئة الخوارزمية للذكاء الاصطناعي التوليدي وأن هذا الهايتوس بات مؤثرا في طرائق الكتابة والتفكير والتعبير في الخطاب البشري. كل هذا يستدعي إلى انتهاز سبل نقدية ووضع سياسات تعليمية عربية بصيرة.

الكلمات المفتاحية: الخوارزميات، برامج الانتحال، الذكاء الاصطناعي، الهايتوس، السوسيولوجيا، اللغة.

Abstract

This research belongs to the sociology of algorithms. It is a qualitative study that investigates the impact of programmes designed to detect AI-generated texts. The study argues that these programmes do not truly aim to safeguard academic research or ensure fairness among scholars; rather, they reproduce a dominant language and culture. They are shaped by linguistic and racial biases that push aside many non-native English speakers, especially those of Black or Latin American heritage. Trust in the algorithms that claim to expose AI content is, the study maintains, an illusion that conceals economic and symbolic interests. At the same time, it affirms that human language itself is being reshaped by an algorithmic habitus—a disposition formed through the algorithmic upbringing of generative artificial intelligence. This habitus now influences the ways of writing, thinking, and expressing ideas within human discourse. All of this calls for critical approaches and for the development of far-sighted Arab educational policies.

Keywords: algorithms, plagiarism-detection programmes, artificial intelligence, habitus, sociology, language.

للاقتباس: فرح، عبد الإله. بزي، محمد. بن داود، صابر. طالبي، عبد اللطيف. حبران، حسن. (2025). الذكاء الاصطناعي التوليدي واللغة البشرية: سوسيولوجيا الشك الخوارزمي لأدوات الانتحال بين التحيز ووهم الدقة، واختراق اللغة البشرية، مجلة سبأ للعلوم الإنسانية والاجتماعية، مج1، ع(4): 154 - 187

Cite this article as: Farah, Abdelilah. Bazi, Mohammed. Ben Daoud, Sabir. Talbi, Abdellatif. Habrane, Hassan.. (2025). Generative artificial intelligence and human language: The sociology of algorithmic suspicion of tools of plagiarism between bias and delusion of accuracy, and penetration of human language. Saba Journal of Humanities and Social Sciences Volume 1, Issue (4), Pages: 154- 187

المقدمة:

شهد العالم في العقد الأخير عصرا جديدا تتدفق فيه مختلف المعلومات والبيانات بانسياب كبير بفضل تطور تكنولوجيا الاتصال والذكاء الاصطناعي الذي يعتمد على الخوارزميات، الأمر الذي جعل من الأخيرة تلعب دورا كبيرا في هذا العصر الذي تجاوزت فيه البعد التقني حتى وصل الأمر بها إلى الهيمنة على باقي الحقول الاجتماعية بما فيها الحقل الأكاديمي. فقد غدت خوارزميات الذكاء الاصطناعي شأننا يوميا يألفه الباحثون والطلاب من مختلف الثقافات والجنسيات، فالبحوث والمقالات والرسائل أصبحت تصاغ اليوم وتنقح بمعونة الخوارزميات في برامج الذكاء الاصطناعي، مثل: نماذج اللغة الكبيرة، وبرامج التدقيق اللغوي إلخ، وهو ما أثار جدلا مستفيضا حول حدود "الأصالة" في البحث العلمي. وإثر ذلك ظهرت أدوات للكشف عن محتوى الذكاء الاصطناعي (AI detectors)، مثل: تيرنيتين (Turnitin)¹، وزيروجيتي (ZeroGPT)²، وكوبيليك (Copyleaks)³، وأوريجيناليتي (Originality)⁴، وغيرها من البرامج التي تزعم بأنها عادلة في الفصل بين كتابة البشر وكتابة برامج الذكاء الاصطناعي.

غير أن هذه البرامج ما لبثت أن أظهرت عيوبها؛ بسبب افتقارها إلى الكثير من الدقة؛ وهو ما تسبب في إثارة الخوف والريب في صفوف الأكاديميين والباحثين والطلبة، حتى أنها لم تستطع إلى اليوم أن تعزز الثقة عند الجامعات الغربية التي اختارت أن تقوم بتعطيل أداة الكشف عن محتوى الذكاء الاصطناعي، وذلك بعد اختبارات عديدة، فقد وجدت أن هذه الأدوات متحيزة، ويمكن أن توسم نصوصا بشرية أصيلة على أنها ذكاء اصطناعي، أو نصوص مولدة بالذكاء الاصطناعي على أنها نصوص بشرية. إذ وجدت الدراسات على أن خوارزمياتها تمارس الحيف ضد الأساليب الكتابية البسيطة أو المركبة، أو التي تغيب فيها تنوع الأساليب اللغوية، أو التي تصدر عن السود أو الذين يعيشون الهشاشة. وهنا يتجاوز الأمر جدلا تقنيا محضاً ليغدو مسألة اجتماعية قائمة بذاتها، تتصل بإعادة إنتاج سلطان المعرفة، وترسيخ ضروب من "العنف الرمزي" في الحقول الأكاديمية.

لقد أُمست الكتابة الأكاديمية في ظل هذه الحال أسيرة لتقديرات حسابية قد تخطئ في احتمالاتها، وصار الطالب أو الباحث الأكاديمي موضع اتهام بين هذه الأدوات. كما بات الحقل العلمي مهددا بالتحول إلى مجال للرصد والإذعان لسلطة الخوارزميات؛ بعد أن كان يرتجى من أدوات الذكاء الاصطناعي أن تكون عوناً على الابتكار والإبداع والسمو بأسلوب الكتابة الأكاديمية. وهو ما يطرح بإلحاح سؤال الثقة: هل هذه أدوات تهدف إلى تعزيز النزاهة الأكاديمية أم أنها سلع تجارية، تبتغي فرض هيمنة رمزية واقتصادية على الجامعة والأكاديمي؟

بالنسبة إلى أهداف هذه الدراسة، فهي تتمثل في محاولة الإجابة عن الأسئلة الآتية:

¹. منصة Turnitin هي شركة عالمية تهتم بتطوير تكنولوجيا التعليم، وهي متخصصة في كشف الانتحال، وفحص أصالة الأبحاث والكتابات الأكاديمية عبر تقنيات الذكاء الاصطناعي.

². أداة متخصصة في كشف النصوص المكتوبة بواسطة الذكاء الاصطناعي، تعتمد على خوارزميات تحليل الأسلوب (Perplexity & Burstiness) لتقدير احتمالية أن يكون النص مولداً آلياً أو مكتوباً بشرياً.

³. منصة متقدمة لفحص الانتحال (Plagiarism Detection)، مهمتها هو فحص الانتحال، والتعرف على النصوص المولدة بالذكاء الاصطناعي في الكتابة.

⁴. أداة تستخدم لفحص أصالة المحتوى للانتحال والذكاء الاصطناعي.

- ✓ ما مدى دقة برامج الكشف عن محتوى الذكاء الاصطناعي في السياق الأكاديمي؟
- ✓ هل تنطوي هذه الأدوات على تحيزات لغوية وعرقية تهدد مبدأ العدالة الأكاديمية؟
- ✓ إلى أي حد تتوافق نتائج أدوات الكشف مع الواقع التجريبي؟
- ✓ كيف تعيد أدوات الذكاء الاصطناعي تشكيل اللغة الأكاديمية والعلاقة بين الإنسان والآلة؟

منهجية الدراسة

قد لا نجد أفضل من الاعتماد على المنهج الوصفي التحليلي للموضوع الذي نحن بصدد فحصه وتحيصه؛ فهو منهج يعتمد على فحص البيانات والوثائق بهدف وصفها وتحليلها، الأمر الذي يجعله ملائماً لطبيعة الدراسة القائمة على مقارنة سوسيولوجية للخوارزميات. فالدراسة تبحث في المخاطر الكامنة في خوارزميات الذكاء الاصطناعي، وخصوصاً برامج الكشف عن المحتوى المولد بالذكاء الاصطناعي التوليدي (ChatGPT). ولتحقيق الأمر؛ تطلب منا إنجاز مراجعة للأدبيات العلمية والتقارير التي تولي اهتماماً بموضوع الخوارزميات في علاقتها بالمجتمعات، بالإضافة إلى رصد ما تم نشره من بيانات رسمية من قبل الجامعات والمعاهد بشأن استخدام برامج الكشف عن AI. وتوثيق حالات واقعية من منصة ريديت (Reddit)¹ التي عبر فيها المستخدمون عن تجاربهم مع برامج الكشف للمحتوى الآلي أيضاً. وهذا يعني بأن نمارس الملاحظة من دون المشاركة، فالتخفي عن المستخدمين سيفيدنا في تحديد المضامين والبيانات الأكثر قيمة بهدف؛ "التعرف على هذا العالم كما يفعل المشاركون في البحث من الداخل" (Charmaz, 2006, p.14). وقد قمنا باختبار أدوات الكشف عن محتوى AI مثل: برنامج (iThenticate)² التابع لشركة تيرنيتين بالاعتماد على نصوص متنوعة: نصوص بشرية لأساتذة سبق لهم الشكوى من هذه البرامج، ونصوص مولدة بالذكاء الاصطناعي، وأخيراً نصوص مؤنسة. فقمنا بتوثيق نتائج الاختبارات على شكل صور بنية استكشاف أبعاد السلطة والثقة والمعرفة التي تولدها تلك البرامج. ففي عملية التحليل اعتمدنا على العديد من الأطر النظرية في السوسيولوجيا؛ بغرض إظهار عمل الخوارزميات والذكاء الاصطناعي في البرامج. إذ نبتغي إبراز المفارقات حول ما تدعيه الأدوات من دقة، وتتبع أثرها في الحقل الأكاديمي.

1) الجامعة وإشكالية الدقة والخطأ في برامج الكشف

1.1. تسليع النزاهة الأكاديمية

أضحت أدوات الكشف عن الانتحال والذكاء الاصطناعي ركناً مهماً في ميدان الممارسة الأكاديمية ومراكز البحوث؛ لكونها تهدف إلى حماية الباحثين والطلاب في عملية الكتابة الأكاديمية من خلال تعزيز سياسة النزاهة الأكاديمية بنية حماية حقوق النشر. بيد أن هذه التطبيقات والبرامج التي باتت في حرب مع أدوات إضفاء السمة الإنسانية على الذكاء الاصطناعي لم تبلغ درجة عالية من الدقة، مع ما يروج له عند الزبائن من الجامعات ومراكز البحوث والأكاديميين، خصوصاً، وأن هذه الأدوات تعتمد على خوارزميات لا تفكر في قصيدة النص أو فحواه؛ فما تستند إليه هو عبارة عن

¹ شبكة اجتماعية ضخمة تبنى على مجتمعات متخصصة تسمى (subreddits). مجتمع يدور حول موضوع محدد: من العلوم والفلسفة إلى الألعاب، ومن النصائح الشخصية إلى الأبحاث الأكاديمية.

² أداة أكاديمية شهيرة لفحص الانتحال العلمي والأدبي والذكاء الاصطناعي، تعد مرجعاً في الجامعات ودور النشر. وهي من منتجات شركة تيرنيتين.

مؤشرات وقياسات، مثل: تواتر الكلمات أو النمط أو شكل الجملة في الطول، وهي خصائص النصوص الاصطناعية. لكنها في البتة لا تفسر سبب التصنيف على أنه نص آلي، هل بسبب كلمة أو جملة أم إيقاعها أو بنيتها؟ لأنه حتى الاقتباس المدرج في النص يتم اعتباره على أنه ذكاء اصطناعي.

وعلى سبيل المثال، فمنصة تيرنيتين (Turnitin) تخبرنا بأنها تبحث عن الأنماط الشائعة في الكتابة بالذكاء الاصطناعي، لكنها لا تشرح أو تحدد ما هي تلك الأنماط، أو من أي مصدر جاءت؟ وما يثبت هذه المعضلة أن حتى شركة OpenAI التي أصدرت ChatGPT أغلقت موقعها AI Classifier الخاص بالكشف عن المحتوى AI في 20 يوليو 2023 بعد أن أطلقتته في يناير 2023؛ بسبب عدم دقته؛ لأن معدل دقته منخفض جدا. فالأداة كانت تتعرف فقط على 26% من النصوص التي كتبها الذكاء الاصطناعي "نسبة إيجابية صحيحة"، لكنها أيضا وعن خطأ كانت تعتبر 9% من النصوص البشرية على أنها ذكاء اصطناعي "إيجابيات كاذبة" (OpenAI, 2023, July). (20)

ومن المعلوم أن برامج الكشف عن محتوى AI برزت بسبب ظهور الذكاء الاصطناعي التوليدي ChatGPT، وما شابهه من الأدوات التي تولد النصوص الآلية. وهو ما دفع بمجموعة من الشركات بما فيها تيرنيتين إلى استغلال حالة التوجس والقلق الذي ساد الحقل الأكاديمي، فقامت بطرح منتجها الخاص بالكشف عن محتوى الذكاء الاصطناعي بعد ستة أشهر من إطلاق ChatGPT، إذ نتج عن ذلك زيادة في مدة العقود، وزيادة في قيمتها المالية، فهناك جامعات أبرمت معها عقود طويلة "عشر سنوات"، وبمبالغ ضخمة تقترب إلى مليون دولار بحجة حماية النزاهة الأكاديمية (García Mathewson, T. 2025, June 26). ومن هنا بات الذعر من الذكاء الاصطناعي التوليدي وسيلة ممتازة، تمكن الشركات من جني أرباح ضخمة بالملايين الدولارات. بالإضافة إلى أن شركة تيرنيتين دفعت بالمؤسسات الأكاديمية إلى التخلي جزئيا عن حقوق الملكية للأعمال التي ينجزها الطلاب والأساتذة بهدف استغلالها كبيانات لغرض تطوير أدواتها وتسويقها، فيما بقيت المشكلات التي تعاني منها خوارزمياتها على مستوى الدقة والتمييز والانحياز من دون حلول جذرية (بالتصرف: Ibid).

إن ما يثير القلق هو المبالغة في الادعاء بأن برامج الكشف عن محتوى الذكاء الاصطناعي تصل دقتها إلى ما يفوق 90%، مثل برنامج Originality 100% دقة من دون أخطاء إيجابية، في حين أن دقة Copyleaks تصل حتى 95% دون إشعار بالخطأ. أما برنامج Turnitin تصل حتى 94%، في حين تصل دقة ZeroGPT إلى 96%. غير أن هذا الادعاء يبقى تسويقيا بالدرجة الأولى، وخاضعا للمنافسة، والهدف منه كما نلاحظ هو السيطرة على حقل الكتابة الأكاديمية في المؤسسات الجامعية. ذلك أن أثر هذه البرامج مرتبط بتنامي دور القطاع الخاص في التعليم وفي تحويل الجامعات إلى سلعة، فأصبح الطالب ينظر إليه بشكل متزايد بوصفه زبونا، والأكاديمي بوصفه مقدم خدمة، وأصبحت المقالة الأكاديمية بما تحمله من أرصدة دراسية ساحة للتبادل الاقتصادي: كتابة أكاديمية مقابل رصيد، ورصيد مقابل شهادة، وشهادة مقابل وظيفة، وهكذا دواليك" (Matzner, 2024, p.30). وهو ما يدخل في باب ما يسميه بيير بورديو بالرأسمال الرمزي. وكما يقول بورديو في كل حقل اجتماعي هناك قوانين عمل ثابتة (بورديو، 2012، ص 181)، وفي كل منها يلاحظ فيها صراع وتنافس، وهو ما يؤكد بورديو بقوله: "نعرف أننا سوف نجد صراعا في كل حقل (يتعين علينا أن نبحث، في كل مرة، عن أشكاله الخاصة) بين الداخل الجديد الذي يحاول أن يكسر أقفال حق

الدخول، والمهيمن الذي يحاول أن يدافع عن احتكاره ويبعد عنه المنافسة" (بورديو، 2012، ص 182). فالشركات المتخصصة في تكنولوجيا التعليم والاتصال غايتها هو مراكمة رصيدها من خلال تبني خطاب مشبع بالثقة في التقنية واليقين العلمي، حتى لو كان في العمق يفتقر إلى موضوعية تامة. وهنا تتدخل خوارزميات كشف الانتحال والذكاء الاصطناعي، وخصوصاً تيرنيتين بوصفها "تكنولوجيا للحكم والضبط تهدف إلى ضمان "قيمة" المقال الأكاديمي... [إذ] تقوم هذه الخوارزمية بتقييم أصالة النص عبر مقارنته بقاعدة بيانات. وهكذا تتحول الأصالة التي كانت في السابق معياراً مثيراً للجدل يخضع لتقدير القراء البشريين، إلى مهمة آلية لمطابقة الأنماط" (Matzner, 2024, p.30)، إلى جانب حساب الاحتمالية الإحصائية التي تشير إلى كون النص مولداً آلياً. وتلك النسب المئوية ما هي إلا آليات لإنتاج الهيمنة الرمزية في الحقل الأكاديمي، إذ يعاد إنتاج الثقة في هذه الأدوات بوصفها حكماً شرعياً في ميدان أكاديمي مشحون بالصراعات حول الشرعية والمعرفة. وذلك على الرغم من أن الذكاء الاصطناعي الذي يتم استخدامه في "برامج الكشف قاصر على الإحاطة الكاملة لفهم السياق، فهو إن أفلح في تمييز مشكلات القواعد أو العبارات المتكررة فإنه يعجز عن فهم الفروق الدقيقة في البراهين المستخدمة والإبداعات اللغوية التي تتسم بها الكتابة الأكاديمية" (Giray, L. 2024, p.6).

إن تعاضل نفوذ الشركات في مجال تكنولوجيا التعليم والبحث يتجاوز ما هو تقني؛ فمبتغاها الأساس هو الهيمنة على الأسواق الأكاديمية بما تقدمه من سلع ومنتجات، تهدف إلى ضبط المجتمع الأكاديمي والسيطرة عليه. ويتجلى ذلك من خلال طغيان المنطق الاقتصادي الذي ترسخه في المؤسسات الأكاديمية، والذي يدفع إلى تسليع العمل الذهني، وخصخصة خدمات التعليم والبحث العلمي؛ لأن الشركات التي تحركها الرأسمالية الأكاديمية لا تطبق المعرفة غير المسلعة، ولا تأبى للقيم والمثل العليا التي تدافع عن فكرة أن دور الجامعة هو خدمة المجتمع. وهو ما يلفت إليه بوب جيسوب بأن الرأسمالية الأكاديمية لا تدار فقط بمنطق السوق، وإنما عبر صفقات سياسية غير شفافة، إذ إن "أكبر المستفيدين من هذه المعاهدات التجارية هم مزودو التعليم الأميركيين الهادفين إلى الربح" (Jessop, B., 2018, p. 108)، ومن ثم فإن أدوات الكشف قد تسوق كتطبيقات محايدة، لكنها تشتغل وفق أولوية ربحية على الرغم من أنها منتجات تتركب أخطاء فادحة في تصنيف النصوص، "فهذه الأدوات أشبه بأنظمة النماذج اللغوية الكبيرة، إذ تعمل كصناديق سوداء التي تفتقر إلى الشفافية لتصنيف النصوص على أنها بشرية أو مولدة آلياً، فهي تستند في حكمها فقط إلى التعرف على الأنماط. وفوق هذا حتى لو أظهرت هذه الأدوات دقة بالغة في كشف النصوص المولدة آلياً، فإن أياً منها لا يزعم خلوها من التحيز ضد فئات معينة من المؤلفين أو أنماط الكتابة" (Pratama, A. R., 2025, p.3).

وبهذا المعنى يتحول خطاب الدقة والإحصاء عند هذه الأدوات إلى شكل من أشكال "العنف الرمزي" الذي يخفي رهانات السوق والتنافس الرأسمالي خلف ستار العلمية والحياد. وهو الأمر الذي يعبر عنه الباحث سهيل فيزي، مدير مختبر الذكاء الاصطناعي في جامعة ماريلاند (Maryland)، إذ يقول: "هناك الكثير من الشركات التي تجمع الكثير من التمويل، وترغم أن لديها أجهزة كشف يمكن استخدامها بشكل موثوق، لكن المشكلة هي أن أياً منها لا يشرح ما هو التقييم، وكيف يتم ذلك، إنها مجرد لقطات عابرة" (Quintana, C., 2024, February 9).

1. 2. انعدام الثقة في الحقل الأكاديمي

كما سيتبين في المحور الثالث من هذه الدراسة فإن جل هذه البرامج تصدر نتائج إيجابية كاذبة أمام النصوص التي يتم كتابتها بأسلوب منمق وصارم، أو عندما يتم نقلها بشريا من لغة إلى لغة أخرى، أو في حالة ترجمتها حرفيا؛ إذ تؤكد دراسة وير-وولف وآخرين بأن النصوص البشرية التي يتم ترجمتها عبر تطبيقات مساعدة، مثل: (Google Translate) يتم احتسابها ذكاء اصطناعياً، وهو ما يعبر عن انخفاض في جل البرامج بنسبة 20% (Weber- Wulff et al., 2023, p.19). وهي غير دقيقة عند استخدام البرامج والمواقع الخاصة بالقواعد اللغوية، إذ تكون برامج الكشف عن محتوى الذكاء الاصطناعي أكثر حساسية أمام صرامة القواعد اللغوية. وتكشف دراسة بأن برامج، مثل: Turnitin، و ZeroGPT تفشل عندما "يعاد توليد نص مولد بالذكاء الاصطناعي بواسطة برنامج لإعادة الصياغة، فالأدوات تعتبره نصا بشريا. حيث بلغت الدقة في هذا الاختبار 26% فقط؛ مما يعني بأن معظم النصوص الآلية تظل خافية بعد إعادة صياغتها" (Ibid, p.20).

علينا أن ندرك جيدا بأن برامج الكشف عن محتوى AI تعمل على التخمين أو الشك في المحتوى، فهي لا تحدد بدقة الموضوع الخاص بالذكاء الاصطناعي وإنما تكتفي بوضع إطار على النص بلون مغاير كإشارة على أنه أسلوب يشابه أسلوب الذكاء الاصطناعي. فبرنامج تيرنيتين يعتمد على درجة انتظام النصوص، فالناطق الأصلي يميل إلى الانفجار من خلال تنويع الجمل والأساليب. أما الناطق غير الأصلي فيكتب بشكل بسيط ومتوازن. وأمام هذه المشكلة، فإن جامعة فاندربيلت (Vanderbilt University) بولاية فلوريدا الأمريكية قامت بتعطيل الكشف بالذكاء الاصطناعي الخاص بشركة Turnitin بعد أشهر من استخدامها (Coley, M., 2023, August 16)، فيدعي البرنامج بأن نسبة الإنذارات الكاذبة لا تتجاوز 1% (Turnitin, 2023, April 5). أي أن من بين 500 باحث أو طالب، يمكن أن تظهر 5 إحالات إنذار كاذبة. وعلى الرغم من ادعاء الشركة بعدم التحيز ضد الناطقين غير الأصليين، فإن ما تراكم من بحوث علمية وشهادات على منصة ريديت (Reddit) (الصورة رقم 1)، ومن بينها شهادات لناطقين أصليين، يقدم صورة مغايرة تماما.

الصورة رقم 1: توضيح خطأ من Turnitin يضع علامة على عمل بشري باعتباره مكتوبا بالذكاء الاصطناعي



المصدر: فريق البحث، بتاريخ 12 شتنبر 2025

لم يتوقف إلغاء منصة الكشف عن محتوى AI في تيرنيتين على جامعة فاندربيلت، بل تبعها جامعات أخرى، مثل: جامعة نيويورك (NYU) التي عطلت الميزة بتاريخ 8 شتنبر 2023، فقد اكتشفت أن شركة تيرنيتين تدعي أن معدل الخطأ في التعرف على النص البشري كنص مولد بالذكاء الاصطناعي هو 1%، لكن الأبحاث الصيفية أشارت إلى أن هذا

المعدل أقرب إلى 4% (Shirky, C., & Maddox, B., 2023, September 7). كما قامت جامعة جورج تاون (Georgetown University) ابتداء من أكتوبر 2023، بإيقاف ميزة كشف الكتابة بالذكاء الاصطناعي في فحوصات التشابه التابعة لتيرنيتين؛ وذلك بسبب مخاوف اللجنة التنفيذية لأعضاء هيئة التدريس في الحرم الجامعي الرئيسي، ومجلس الشرف من دقة الأداة، واعتقادهم بأن الأضرار الناتجة عن النتائج الإيجابية الكاذبة أسوأ من أي فوائد محتملة للأداة (University Information Services, n.d). بالإضافة إلى جامعة كاليفورنيا في بيركلي التي اختارت عدم تفعيل خاصية الكشف عن محتوى AI بعد فترة تجريبية من خريف 2023 حتى ربيع 2025؛ وذلك بناء على نتائج ودراسات، فاعتبرت أن البرنامج لا يلي معاييرها (RTL – Research, Teaching, & Learning, 2023 ; 2025).

لقد انخرطت العديد من الجامعات الأمريكية والكندية والأسترالية إلى إلغاء خاصية الكشف عن محتوى الذكاء الاصطناعي في تيرنيتين، ومنعت الأساتذة من استخدام أي تطبيق أو برنامج يدعي الكشف عن محتوى الذكاء الاصطناعي؛ وذلك لأسباب عديدة، تتعلق بالخصوصية، إلى جانب أن مثل هذه الشركات المتخصصة في الكشف عن محتوى الذكاء الاصطناعي ترفض الكشف عن كيفية عمل خوارزمياتها، أي أنه لا توجد شفافية في كيفية حساب الأداة لدرجاتها، أو تحليل المنهجية التي استخدمتها الأنظمة، فهي بمثابة صندوق أسود. وهو ما دفع بالعديد من الجامعات اليوم بأن تكون غير قادرة على الوثوق بنتائجها.

وقد كان إلغاء خاصية الكشف عن AI في منصة تيرنيتين من الأسباب التي دفعت الشركة نفسها إلى تقديم إشعار في واجهة الأداة ينص على أن نتائجها "قد لا تكون دقيقة دائماً (إذ قد تولد نماذج الذكاء الاصطناعي نتائج إيجابية أو سلبية زائفة)، لذلك لا ينبغي الاعتماد عليها وحدها كأساس لاتخاذ إجراءات سلبية ضد الطالب" (تيرنيتين، إشعار نافذة الكشف). فلا نرى في إشعار الشركة إلا اتصالاً من المسؤولية إذا ما لحق الباحثين أو الطلبة حيف؛ بسبب النتائج الإيجابية الكاذبة للبرنامج؛ لأن الإشعار يلقي المسؤولية على عاتق الأستاذ المشرف أو الإدارة. بالمقابل ولمواجهة هذه المشكلة، قامت الشركة نفسها في 2025 بطرح تطبيق (Turnitin Clarity) الذي يستخدم الذكاء الاصطناعي من أجل مساعدة الطلبة في الجامعات والمعاهد على كتابة النصوص من حيث أنه يقوم باقتراح تحسينات، ويعطي ملاحظات كما يساعد على تحسين الأسلوب. ومع ذلك، فإنه رغم هذه الجهود، تم حظر خاصية الكشف عن AI في العديد من الجامعات، ولا تزال العديد من الجامعات الغربية حتى اليوم لا تعتمد على خاصية الكشف عن محتوى AI في برنامج (iThenticate/Turnitin) التابع لتيرنيتين وغيرها من البرامج التابعة لها. وحتى كتابة هذه الدراسة، فقد أعلنت جامعة واترلو (University of Waterloo) الكندية في بداية شهر شتنبر عن إيقاف استخدام خاصية الكشف عن الذكاء الاصطناعي في برنامج تيرنيتين. فقد "وجدت الاختبارات الداخلية التي أجرتها مجموعة التقنيات التعليمية والخدمات الإعلامية (ITMS) بالجامعة أن مزايا أداة الكشف لم تكن حاسمة. على سبيل المثال، في أكثر من حالة، قام المنتج بوضع علامة على النص المكتوب بواسطة الإنسان على أنه تم إنشاؤه بنسبة 100% بواسطة الذكاء الاصطناعي، وهو أمر مثير للقلق" (University of Waterloo, 2025, September). كما قررت جامعة كيب تاون في جنوب إفريقيا عن توقفها في استخدام أداة Turnitin's AI Score في الكشف عن محتوى الذكاء الاصطناعي اعتباراً من 1 أكتوبر 2025 (University of Cape Town, 2025, July 24). علاوة على ذلك، نجد جامعة كورنيل الأسترالية التي قررت بدورها توقيف برنامج الكشف عن محتوى AI في تيرنيتين ابتداء من 1 يناير 2026،

مع إبقاء على فحص التطابق التقليدي (Originality) نشطا (Curtin University, 2025, September 04). كما أنه بدلا من محاربة محتوى الذكاء الاصطناعي، بدأت مجموعة من الجامعات والمعاهد في التكيف معه، من خلال تطوير تقنيات تعتمد على الذكاء الاصطناعي في إنجاز البحوث والمقالات، مثل: منصة ستورم (<https://storm.genie.stanford.edu>) التي أطلقتها جامعة ستانفورد لمساعدة الأساتذة والطلبة على كتابة مقالاتهم وبحوثهم وغيرها.

لقد أفضى ما قرره الجامعات بوقف خدمة تيرنيتين إلى وضع مسألة الثقة في صميم الجدل الأكاديمي؛ فقد كشفت التجارب أن مثل هذه البرامج المبتكرة التي تهدف إلى تعزيز الأمانة العلمية قد جعلت من التقنية وسيلة لإثارة الخوف المنظم حتى طالت ربيتها نصوصا تم صياغتها بيد البشر بشكل خالص. وما يهم هنا سوسيولوجيا "ليس ما يحدث في دماغ الآلة الاصطناعي، وإنما ما تخبر به الآلة مستخدميها، وما يترتب عنها من نتائج. فبفضل القدرات الثقافية والتواصلية الجديدة التي طورتها من خلال استغلال البيانات البشرية المولدة عبر الإنترنت، تحولت الخوارزميات في جوهرها إلى وكلاء اجتماعيين" (Airoldi, M. 2022, p.6). لذلك فالشغف ببلوغ مراتب عليا من الدقة حدا بالشركات إلى اختراع برامج قائمة على احتمالات بوسعها أن تصف الكتابات البسيطة والمحايدة بأنها كتابة آلية. وهو ما يمكن أن يؤثر في أي باحث في المستقبل، ويدفعه إلى الدفاع عن أصالة ما خطه بيده أمام أي مؤسسة. وهذا الأمر هو نوع من التضيق؛ لأنه يحاكي ما اعتبره بورديو بـ «العنف الرمزي»، إذ إن هذه الأدوات صارت في زمن الذكاء الاصطناعي تمتلك سلطة مرجعية قاطعة، علما أنها تقوم على احتمالات احصائية واهنة، وهو ما يتقاطع مع ما ذهب إليه ونكون هونغ (Soonkwan Hong) حين كتب: "لقد أعيد في هذا العصر تصور الثقافة وتشكيلها استنادا إلى ثقة مرتجلة في الخوارزميات، حيث تعمل الأجيال الحاضرة بمقتضاها وكأنهم عملاء أحرار للنظام. فالثقة هي الفضيلة التي تعلي هذا النظام الثقافي بنية صون كيانه واستمراره؛ ومن ثم، وجب أن يعاد صياغة ثقافة الاستهلاك حين تغدو الثقة العابرة شرطا لازما على الكافة؛ للحفاظ على مقامهم كفاعلين في منظومة السوق، ولعل الذاتية قد غدت قيمة بالية حل محلها بديل إحصائي ما هو إلا مقارنة خوارزمية للذات" (Hong, 2020, p.45).

إن هذه الموثوقية في الأساس مبنية على الإيمان المفرط بالتقنية الذي يتأثر بالعمليات الاجتماعية المثقلة برهانات الأسواق والرأس المال الرمزي. فخطاب شركات التكنولوجيا بمختلف أنواعها الذي يتخذ من الثقة شعارا، ويدعي العلم هو في جوهره يضمن تسابقا تجاريا ضروسا، وهو ما يبينه كين فوشير في دراسته حول رأس المال الاجتماعي على الإنترنت، إذ يشير إلى أن "دافع الربح لم يتوارى في خضم التحول نحو مجتمع المعلومات؛ فقد خابت المزايم المثالية التي بشرت بأن الاعتماد على تكنولوجيا المعلومات والاتصال سيصرف الأنظار عن الربح إلى نفع الناس. وظل الربح هو الغاية الأولى، وإن تم تزيينه بمظهر المنفعة العامة. وقد استقر هذا الدافع في صميم وسائل التواصل حتى أنه سوق للأفراد على أنه ميزة تتيح لهم أن يصيروا أرباب أعمال يستغلون تلك الشبكات لمنافعهم الخاصة. وما هذا في بعض جوانبه إلا ضرب من التسويق الهرمي" (Faucher, 2018, p.34). وهذا الأمر لا يختلف عن أدوات الكشف عن محتوى الذكاء الاصطناعي، فما يطرح في هذه الأدوات من أخطاء ما هو إلا دليل على أن الغاية تكمن في الربح بالدرجة الأولى. وهكذا أصبحت الشركات الخاصة بالانتحال والكشف عن محتوى الذكاء الاصطناعي من أكثر الخصوم داخل الحقل الأكاديمي؛ لأنها تسعى إلى التحكم في ممارسة الأكاديميين عن طريق إجبارهم على احترام خوارزمياتها.

2) التمييز اللغوي والعرقي في عمل الخوارزميات

تظهر الدراسات أن الخوارزميات الخاصة بالذكاء الاصطناعي أياً كان مبتكرها ومطورها، وأياً كان جوهرها الرياضي تستقي من بيانات مشبعة بالتحيز. وتتميز بأنها وسيط اجتماعي يعمل على السيطرة على اللغة بشكل كبير، فتفرض مقاييس جديدة للفكر ولأسلوب الكتابة، كما أنها تفرض واقعا لغوياً مستجداً (Anderson, B., Galpin, R., & Juzek, T. S, 2025). فما يميزها أنها تعيد إنتاج علاقات القوة واللامساواة. بحيث يمكن أن تعرض الناطقين غير الأصليين أو المنتمين إلى تخصصات معينة إلى النبذ والتهميش، واتهامهم بالانكفاء على أدوات الذكاء الاصطناعي، في حين أن سعيهم هو صقل أسلوبهم وقدرتهم على التعبير والانخراط في ميدان العلم والتنافس مع الناطقين الأصليين. وهو ما يشير إليه الباحث لوي جيراي بالقول: "تتعدد المشكلة أكثر بسبب التحيزات المتأصلة داخل خوارزميات الذكاء الاصطناعي. فهذه الأنظمة، التي دربت على بيانات سابقة الوجود، قد تعيد إنتاج التحيزات الموجودة أصلاً في العالم الأكاديمي من غير قصد. على سبيل المثال، قد يجد الباحث القادم من خلفية غير غربية أن عمله يخضع لتدقيق أشد من نظرائه الغربيين، فقط لأن الكاشف الآلي غير معتاد على أسلوب كتابته أو سياقه الثقافي" (Giray, L., 2024, p.2).

وقد أظهرت البحوث أن أداة مثل GPTZero على ما لها من دقة عالية (97%) قد وصمت قرابة 44% من نصوص البشر على أنه AI رغم نشرها قبل 2022. واعتبرت بأن النصوص "المحسنّة" التي يكتبها الناطقون بغير الإنجليزية بأنها نتاج AI بنسبة 44.6% مقابل 30.7% فحسب عند الناطقين الأصليين؛ كما أبانت أدوات DetectGPT وZeroGPT عن ضعف في الدقة (64% و 54%) مع جنوح ضد أصحاب تخصصات تقنية، مثل: العلوم التكنولوجية والهندسة (Pratama, A. R., 2025). وهذه النسب دليل على أن تفحص الخوارزميات لا يقتصر شأنها على تحليل النصوص، وإنما تعيد إنتاج التفاوت اللساني والمعرفي، وترسخ دعائمه في رحاب النشر الأكاديمي.

العديد من الدراسات العلمية تنتقد برامج الكشف عن محتوى الذكاء الاصطناعي. ففي إحدى الدراسات التي قام بها باحثون في جامعة ستانفورد، تم الاعتماد على 394 نصاً بشرياً أصلياً (91 مقالة TOEFL) للناطقين بغير الإنجليزية، 88 مقالا لطلاب الصف الثامن الأمريكيين، 70 مقالة جامعية، و145 ملخصاً علمياً، إضافة إلى 176 نصاً مولداً بالذكاء الاصطناعي (31 مقالة جامعية و145 ملخصاً علمياً تم إنشاؤها باستخدام ChatGPT-3.5). وقد أظهرت النتائج أن كواشف الذكاء الاصطناعي تعاملت بدقة شبه كاملة مع مقالات الطلاب الأمريكيين، في حين صنفت أكثر من 61% من مقالات (TOEFL) المكتوبة من قبل الناطقين بغير الإنجليزية على أنها نصوص آلية. كما أن 97.8% من هذه المقالات وُصفت من قبل كاشف واحد على الأقل على أنها مولدة بالذكاء الاصطناعي (Liáng, et al., 2023, p. 2)، وهو ما يعكس مستوى تحيز مرتفع ضد الكتاب ذوي الكفاءة اللغوية المحدودة. بالمقابل انخفض معدل الخطأ بشكل كبير عند تحسين التنوع اللغوي لتلك النصوص، الأمر الذي يؤكد أن التحيز مرتبط أساساً بمحدودية التنوع التعبيري لدى الناطقين بغير الإنجليزية أكثر من كونه مؤشراً حقيقياً على توليد آلي.

وقد لوحظ بأن هذه الأدوات قائمة على التمييز الاجتماعي، ومتحيزة ضد الأصول العرقية، فخطؤها لا يتوقف ضد الناطقين غير الأصليين، وإنما حتى ضد الأصول الإثنية حيث تمارس التمييز أيضا. فقد أشار تقرير بعنوان (The Dawn of the AI Era) نشر في سبتمبر 2024، بأن "أدوات الكشف عن المحتوى المولد بالذكاء الاصطناعي لا تؤثر بالتساوي في الجميع. فالمرهقون السود هم الأكثر عرضة للاتهام الخاطئ (20%)، يليهم اللاتينيون (10%)، ثم البيض (7%)". (Common Sense Media, 2024, p.10). وهذا يؤكد بأن أدوات الكشف عن محتوى AI تعيد إنتاج التفاوتات العرقية القائمة بدلا من تصحيحها؛ إذ من المعلوم أن عمل الخوارزميات في أي برنامج يقوم على مبادئ مشتركة التي تعتمد على التعلم الآلي، على الرغم من الاختلاف في طريقة تنفيذها. ومع ذلك فقد "صيغت خوارزميات التعلم الآلي في أصلها؛ لتبلغ بدقة التنبؤ منتهاها استنادا إلى ما لقنت من بيانات. وعليه، فإذا طغى تمثيل فئة سكانية ما في تلك البيانات؛ فإن تلك الخوارزمية تنجح إلى تجويد مقاييس أدائها لصالح تلك الفئة، فيترسخ بذلك ما كان سائداً من انحيازات" (Gorska, A.M & Jemielniak., 2025, p.215).

وبالتالي، فتأثير الخوارزميات أوسع وأكبر مما نتخيل على مستوى الأفراد والمجتمعات على حد سواء. وتشير كيلي-هولمز إلى أن ولوج الذكاء الاصطناعي إلى حقل علم اللغة الاجتماعي جعله فاعلا يعيد صياغة العالم، إذ يفرض منطقاً خوارزمياً جديداً يوجه نظرنا إلى اللغة وصلتها بالمجتمع، وهو ما يحدث تبديلاً في ماهية العمل نفسه، إذ تكتب: "ليس من المستغرب أن الذكاء الاصطناعي يتوقع أن يجلب مزايا وتحديات لعملائنا. المزايا تتمثل في جمع البيانات وتحليلها، والقدرة على التعامل مع كميات هائلة من البيانات، وربطها على نطاق واسع، لكن مع الجانب السلبي المتمثل في التحيزات المتأصلة في النماذج الخوارزمية" (Kelly-Holmes, 2024, p.4).

تفصح البحوث عن جنوح خوارزميات كواشف الذكاء الاصطناعي، وجورها على الأساليب اللغوية البسيطة والمعقدة. وإن كثرة وقوع الاتهامات الباطلة على الفئات المشهورة له أثر بالغ يلقي الرعب في ميدان العلم؛ مما يشيع الريبة والشك بين الأساتذة وطلبتهم، ويلغي الثقة فيما بينهم، وحتى أنه يمكن أن يثني عن المشاركة والخوض في غمار العلم، ويزعزع الثقة حول تحقيق العدالة في الامتحانات أو بحوث التخرج. إذ تزداد هذه الاتهامات الكاذبة ضد الطلبة الذين يعيشون المشاشة، أو المصابين بأمراض عصبية، أو أولئك الذين ينحدرون من أصول عرقية، مثل: السود واللاتينيين، ومن أصحاب الخلفيات اللغوية المتباعدة أيضاً، وهو ما تؤكد بحوث في سوسيولوجيا الخوارزميات التي ترى أن أصل التحيز في هذه النظم الخوارزمية، يكمن أساساً في بنيتها المعمارية والثقافية نفسها، والتي تعلي من ثقافة محددة على حساب ثقافة أخرى، فتسيطر عليها النخبة البيضاء، ولا سيما النخبة الغربية. وهو ما يسميه ماسيمو آيرولدي بـ "الإله في الآلة" (deus in machina)، ويقول: "هذه الاستعدادات معمارية وثقافية في الأساس؛ فهي تنشأ من الخيارات التصميمية، ومن ثم من الإله في الآلة. وينطبق ذلك أيضاً على أنظمة الذكاء الاصطناعي المتقدمة؛ فعلى سبيل المثال، يتم تحديد المعاملات الفائقة (hyperparameters) في الشبكات العصبية الاصطناعية مسبقاً من قبل منشئي النظام" (Airoldi, M., 2022, p.63). وعليه فإن هذه الخوارزميات مهما كانت طبيعتها ستبقى متحيزة؛ لأنها ناتجة عن تعلمها من النصوص العديدة التي تم إنتاجها من قبل النخبة البيضاء. وهذا التحيز عظيم الخطر؛ لأنه سيؤدي إلى كوارث عظيمة على الباحثين والطلبة من أصول غير إنجليزية على حد سواء، الأمر الذي قد يحرمهم من الاعتراف الأكاديمي. وهناك نقطة مهمة أشار لوبي جيراي إليها تتمثل في أن "إحدى أكبر المشكلات في أدوات الكشف بالذكاء الاصطناعي هي وجود تحيزات

خوارزمية. تحيزات مدججة غالبا في الخوارزميات؛ مما يجعلها تستهدف بشكل غير عادل أنماط كتابة معينة أو حتى الاختلافات الثقافية في استخدام اللغة. وعلى سبيل المثال، قد يكون أحد برامج الكشف أكثر ميلا لوسم أعمال كتبها الناطقون الأصليون بغير الإنجليزية، لمجرد أن تراكيب جملهم تختلف عن المعايير الشائعة في الإنجليزية. هذا التحيز يعكس ما سماه ستيل وأرونسون بـ "تهديد القوالب النمطية"، إذ يحكم على أشخاص من مجموعات معينة ظلما بناء على أنظمة معينة. وبدلا من تحقيق تكافؤ الفرص، يمكن لأدوات الذكاء الاصطناعي أن تفاقم هذه الفوارق من غير قصد" (Giray, 2024, p.5)، وربما عن قصد أيضا. صحيح بأن الانتقال من النظم البشرية إلى أنظمة التعلم الآلي يمكن أن تمنحنا قدرة على تجاوز الانحيازات البشرية. إلا أن "التعلم الآلي فيه من الخطر الذي قد يرسخ الانحيازات أكثر فأكثر إذا اعتمد النظام الخوارزمي دون تمحيص على بيانات تتضمن تحيزات بشرية. وما لم يصمم الخوارزم عمدا بطريقة متحيزة، فإن البيانات نفسها ستكون هي مصدر التحيز وليس الخوارزمية" (Coglianese, 2021, p.43).

يظهر هذا التحيز مع الناطقين غير الأصليين، ولا سيما مع النصوص التي يتم كتابتها بأسلوب بسيط أو بالاعتماد على مدقق لغوي، أو على مواقع وبرامج التدقيق اللغوي مثل Grammar المجانية و Google Docs التي توفرها شركة غوغل، أو الذين يترجمون من لغة معينة إلى اللغة الإنجليزية. إذ يجب أن نفهم ونذكر جيدا أنه في التعلم الآلي لبرامج الكشف عن محتوى AI أو في نماذج الذكاء الاصطناعي التوليدي فإنه "حتى إذا استبعدت بيانات صريحة عن العرق أو الجنس من مجموعة البيانات، فإن المخرجات الخوارزمية يمكن أن تتأثر بالعرق أو الجنس أو (اللغة) طالما أن المتغيرات الأخرى مرتبطة أو متأثرة بالانحيازات بشرية قائمة على العرق أو الجنس أو (اللغة) أيضا" (Ibid, p.44). ومع ذلك، تظهر المعطيات الميدانية أنه حتى مع الناطقين الأصليين ترتكب هذه البرامج أخطاء في تصنيف النصوص على أنها ذكاء اصطناعي، وتشير ويتني جيج-هاريسون (whitney gegg-harrison) -وهي أستاذة مشاركة في برنامج الكتابة والتحدث والحجة في جامعة روتشستر- إلى أن بعض مقالاتها القديمة أشير إليها هي الأخرى بأنها تحتوي على 30% من الذكاء الاصطناعي من طرف برنامج تيرنيتين (Gegg-Harrison, W., 2024, August 22). بالإضافة إلى أن مثل هذه البرامج ترتكب نتائج سلبية كاذبة للمحتوى، وهو ما نلاحظه من خلال المحور التالي الخاص بالمعطيات التجريبية.

3) المعطيات التجريبية لأدوات الكشف عن محتوى الذكاء الاصطناعي، نموذج iThenticate التابع لشركة تيرنيتين

3.1. التجربة الأولى

تأتي هذه التجربة في سياق ما طرحه أستاذ التاريخ ستيفن مينز (Steven Minz) بجامعة تكساس في أوستن¹، إذ يحكي تجربته عندما قام بأخذ نص كتبه في مدونة Higher Ed Gamma قبل ظهور ChatGPT ومرره عبر أداة كشف؛ فكانت النتيجة أن الأداة أعلنت بأنه نص آلي بنسبة 100% وهو ما جعله يتساءل، هل أنا فقط مجموعة تراكيب أسلوبية، أو بعبارة أخرى: هل أنا روبوت؟ (Mintz, S. 2025, April 2). وللتأكد قمنا باختبار بعض

¹ ستيفن مينز (Steven Mintz) هو مؤرخ أمريكي مشهور، ولد عام 1953 في ديترويت، ميشيغان. يشغل حاليا منصب أستاذ التاريخ في جامعة تكساس في أوستن. طوال مسيرته الأكاديمية، برز كخبير في تاريخ الأسرة والطفولة، ومراحل الحياة، بالإضافة إلى كونه رائدا في دمج التكنولوجيا في التعليم التاريخي.

المقالات التي نشرت له على منصة (Inside Higher Ed)¹ قبل إطلاق برنامج (ChatGPT) في 30 نوفمبر 2022 ثم بعد هذا التاريخ؛ وذلك من خلال برنامج IThenticate التابع لشركة تيرنيتين وكانت النتيجة كالآتي:

1. قبل الإعلان عن نموذج ChatGPT، أي قبل 30 نوفمبر 2022:

✓ مقالة صادرة بتاريخ 01 يوليو 2013، بعنوان: Does History Matter؟، تم الإشارة إليها على أنها ذكاء اصطناعي بنسبة 32%، بتاريخ 23 شتبر 2025.

✓ مقالة صادرة بتاريخ 25 نوفمبر 2017، بعنوان: Would Public Universities Benefit From a Central Innovation Unit؟، تم اعتبارها ذكاء اصطناعياً بنسبة 21%، بتاريخ 9 شتبر 2025.

✓ مقالة صادرة بتاريخ 10 فبراير 2020، بعنوان: 6 Ideas Whose Time Has Come، تم اعتبارها ذكاء اصطناعياً بنسبة 72%، بتاريخ 9 شتبر 2025.

✓ مقالة صادرة بتاريخ 17 غشت 2020، بعنوان: Crisis-Informed Pedagogy، تم اعتبارها ذكاء اصطناعياً بنسبة 56%، بتاريخ 9 شتبر 2025.

✓ مقالة صادرة بتاريخ 28 أكتوبر 2021، بعنوان: Perceptual Blindness، تم اعتبارها ذكاء اصطناعياً بنسبة 64%، بتاريخ 9 شتبر 2025.

✓ مقالة صادرة بتاريخ 1 نوفمبر 2022، بعنوان: Creating the Learner and Learning-Centered University، تم اعتبار نسبة الذكاء الاصطناعي فيها 28%، بتاريخ 9 شتبر 2025.

✓ مقالة صادرة بتاريخ 27 نوفمبر 2022، بعنوان: The Midlife Crisis's 'Evil Younger Brother': The Quarter-Life Crisis، تم اعتبار نسبة الذكاء الاصطناعي فيها 30%، بتاريخ 9 شتبر 2025.

2. مقالات بعد 30 نوفمبر 2022، وبعد إضافة الكشف عن محتوى الذكاء الاصطناعي في شركة تيرنيتين بتاريخ 4 أبريل 2023:

✓ مقالة صادرة بتاريخ 21 يوليو 2023، بعنوان: Rebuilding the Academy of the 1970s، حدد نسبة الذكاء الاصطناعي فيها 55%، بتاريخ 9 شتبر 2025.

✓ مقالة صادرة بتاريخ 27 نوفمبر 2023، بعنوان: Mastering the Art of Quantitative Manipulation، لم يستطع البرنامج أن يحدده فأشار إليه ب: %*، بتاريخ 9 شتبر 2025.

✓ مقالة صادرة بتاريخ 31 يناير 2024، بعنوان: Nostalgia Is a Hell of a Drug، تم اعتبارها ذكاء اصطناعياً بنسبة 35%، بتاريخ 9 شتبر 2025.

✓ مقالة صادرة بتاريخ 1 فبراير 2024، بعنوان: What Should a Public Research University Be؟، تم اعتبارها ذكاء اصطناعياً بنسبة 62%، بتاريخ 9 شتبر 2025.

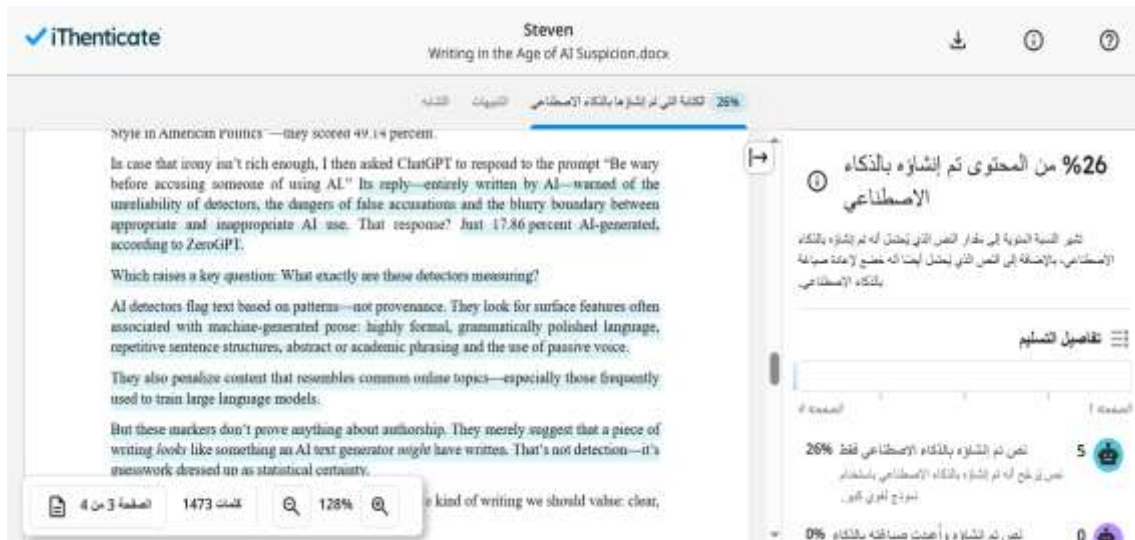
¹. رابط المنصة: <https://www.insidehighered.com>

✓ مقالة صادرة بتاريخ 2 أبريل 2025، بعنوان: Writing in the Age of AI Suspicion، تم اعتبارها ذكاء اصطناعياً بنسبة 26%، بتاريخ 6 سبتمبر 2025.

✓ مقالة صادرة بتاريخ 27 يونيو 2025، بعنوان: A Blueprint for Tomorrow، تم اعتبارها ذكاء اصطناعياً بنسبة 100%. بتاريخ 9 شتبر 2025.

كما قمنا بمقارنة مقالة الأستاذ ستيفن مينز التي تحمل عنوان: Writing in the Age of AI Suspicion، والتي ينتقد فيها برامج الكشف عن محتوى الذكاء الاصطناعي، ببعض الأدوات الخاصة بالكشف عن محتوى IA، وكانت النتيجة كالآتي:

الصورة رقم 2: توضح نسبة الذكاء الاصطناعي في نص ستيفن مينز من طرف برنامج تيرنيتين



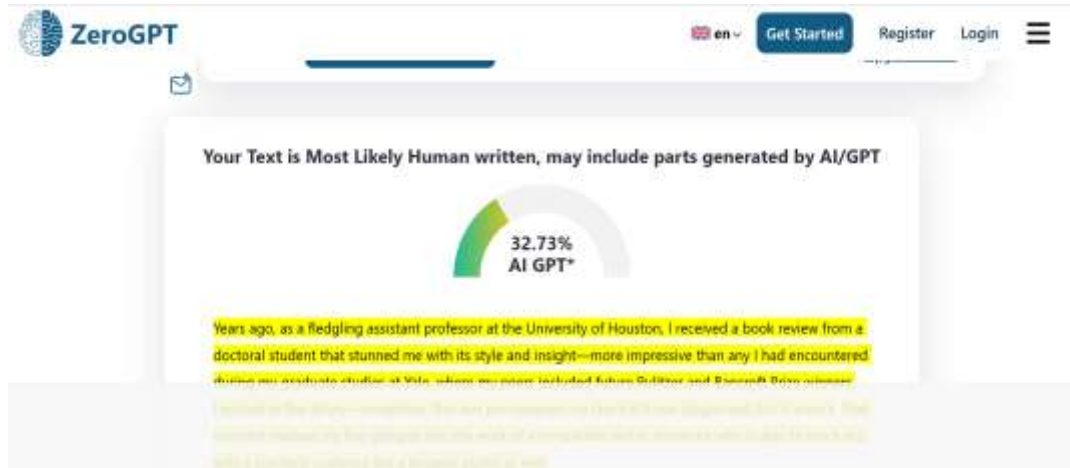
المصدر: فريق البحث، بتاريخ 06 شتبر 2025

الصورة رقم 3: توضح أن برنامج copleaks يعتبر أن نص ستيفن هو ذكاء اصطناعي بنسبة 99%



المصدر: فريق البحث، بتاريخ 06 شتبر 2025

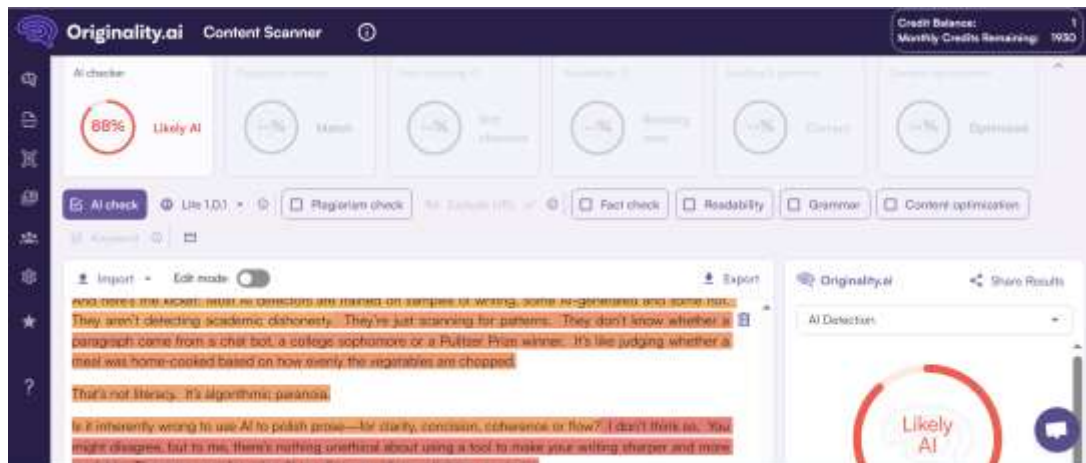
الصورة رقم 4: توضح أن برنامج ZeroGpt يرى في محتوى ستيفن بأنه مزيج بين الأسلوب البشري والذكاء الاصطناعي بنسبة 32.73%.



المصدر: فريق البحث، بتاريخ 06 شتنبر 2025

الصورة رقم 5: توضح بأن مقالة ستيفن تشير بأنها ذكاء اصطناعي بنسبة 88% من طرف برنامج

Originality

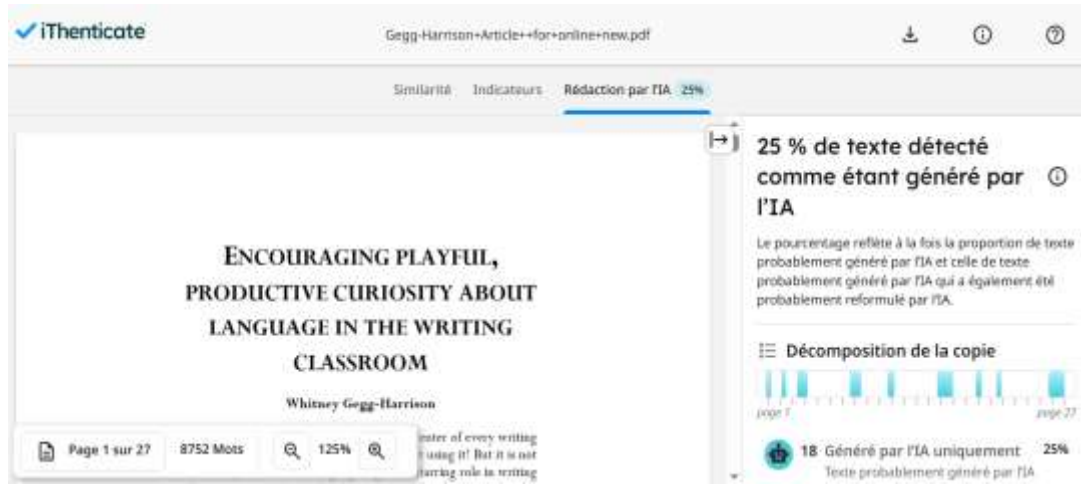


المصدر: فريق البحث، بتاريخ 06 شتنبر 2025

زيادة على ذلك، اخترنا بعض النصوص حسب الإمكانيات المتاحة؛ للوصول إلى ما أنتجته الباحثة ويتني جينغ-هاريسون، وكانت النتيجة أن بعض أعمالها التي نشرت قبل ظهور نماذج الذكاء الاصطناعي، لم يتمكن البرنامج من تمييزها 0%، في حين قام بوسم دراسة لها بنسبة من الذكاء الاصطناعي. على سبيل المثال، دراستها التي تحمل عنوان: Encouraging playful, productive curiosity about language in the writing

classroom، التي نشرت في سنة 2021، فتم الإشارة إلى أن محتواها يتضمن 25% من الذكاء الاصطناعي (الصورة رقم 6).

الصورة رقم 6: توضح نسبة الذكاء الاصطناعي في مقالة ويتني جيج-هاريسون



المصدر: فريق البحث، بتاريخ 18 شتنبر 2025

كما أن هناك عددا من المقالات البشرية التي كتبها أساتذة على منصة (Inside Higher Ed) قبل 2022، حيث تم وسمها على أنها ذكاء اصطناعي بنسب تتراوح ما بين 22% إلى 79% على البرنامج التابع لشركة تيرنيتين¹، مثل: كارولين ليفاندر (Caroline Levander)²، وبيتر ديشري (Peter Decherney)³، ومايكل باتريك راتر (Michael Patrick Rutter)⁴، وغاي ماكلين روجرز (Guy MacLean Rogers)⁵، وغيرهم. وهذه النماذج تؤكد بالملاموس بأن الأمر لا يقتصر على الناطقين الأصليين بغير اللغة الإنجليزية، بل حتى من يمتلكون تنوعا لغوياً ودقة في التعبير أو أسلوباً منمقا. وهي دليل أيضا على أن الأمر لا يختص بحالات فردية معزولة، لكنها في رأينا قد تشمل فئة واسعة من الأساتذة والباحثين حول العالم. ولهذا يبدو بأن الحقل الأكاديمي العالمي يدخل في مرحلة معقدة وخطيرة؛ لأن أدوات الكشف عن المحتوى الآلي باتت تفرض على الكاتب أو الباحث أن يكتب بطريقة ترضي خوارزمياتها، وذلك على حساب الإبداع والابتكار والصرامة اللغوية، إلخ. وهذا الخضوع لخوارزميات الكشف هو ناتج عن التوجه العالمي نحو الرأسمالية الأكاديمية التي تهدف إلى السيطرة على الحقل الأكاديمي، وتحقيق الربح في مجال التعليم والبحث العلمي، وهو توجه يدفع بجميع المؤسسات الأكاديمية إلى تقليد الاتجاهات، واتخاذ نفس الخطوات والمسارات

¹. تم فحصها بتاريخ 23 شتنبر 2025، على سبيل المثال، مقالة كارولين ليفاندر بعنوان: Education Is a Team Sport، نشرت بتاريخ 26 ماي 2020، تم وسمها بأنها ذكاء اصطناعي بنسبة 30%، ومقالة الأستاذ ومايكل باتريك راتر، بعنوان: Where Art Oh Analytics, Thou?، نشرت في 23 يناير 2017، تم اعتبارها ذكاء اصطناعياً بنسبة 62%.

². نائبة عميد المبادرات البيئية والتعليم الرقمي، وأستاذة كرلسون في العلوم الإنسانية، وأستاذة الأدب الإنجليزي.

³. أستاذ دراسات السينما والإعلام بجامعة بنسلفانيا وباحث في الثقافة الرقمية وصانع أفلام وثائقية.

⁴. مدير الاتصالات الاستراتيجية والعلاقات الإعلامية في كلية الهندسة بمعهد ماساتشوستس للتكنولوجيا (MIT School of Engineering).

⁵. أستاذ كيمبر للكلابسيكيات والتاريخ في كلية ويلسلي (Wellesley College) في الولايات المتحدة.

التكنولوجية من دون فحص أو تمحيص، وهي ظاهرة سبق أن وصفها بوب جيسوب حين قال: "أضحى هذا التطور أكثر عالمية بفضل احتدام المنافسة بين المؤسسات المعنية (وهو ما يتجلى، من بين أمور أخرى، في الاعتمادات الدولية للتدريس والتصنيفات الدولية للبحث)، واحتدام السباق بين الفضاءات الاقتصادية والسياسية الأوسع التي تندرج فيها هذه المؤسسات، ويضاف إلى ذلك ما يسمى بالسلوك التقليدي الأعمى الذي يدفع الفاعلين الاجتماعيين إلى مساهمة كل موجة جديدة" (Jessop, B., 2018, p.104).

إن هذه الأمثلة لبعض المقالات والدراسات التي نشرت قبل 30 نونبر 2022، تدل جميعها على أن العديد من الكتابات البشرية سوف تبقى في منطقة الشك، إذ يمكن أن تقرأ على أنها ذكاء اصطناعي، ولا سيما إن اعتمدت على الجمل المزوجة "ليس فقط، بل... (But....; Not Only ...)", أو التقطع السريع لجمل افتتاحية أو ختامية (That's not literacy. It's algorithmic paranoia)، أو جمل واضحة وبسيطة وغيرها. وتشير نتائج إحدى الدراسات إلى أن "تفاوت مستويات أداء هذه الكواشف الخاصة بالمحتوى النصي الناتج عن الذكاء الاصطناعي يكشف عن مدى التعقيد والصعوبات المرتبطة بالتمييز بين النصوص البشرية، وتلك المولدة آلياً" (Elkhatat et al., 2023, p.11). هذا التباين يثير مخاوف حول مدى موثوقية هذه الأدوات، خصوصاً في السياقات الحساسة، مثل: التحقيقات المتعلقة بالنزاهة الأكاديمية (Ibid, p.13). ولا سيما وأن الخوارزميات لا تفرق بين "نية الكاتب"، و"بنية النص"، فعملها يقوم على مؤشرات شكلية: التكرار، والإيقاع، والترابط الداخلي للجمل، وسلاستها الزائدة، أو افتقار النص إلى الارتجال والتشويش الذي هو سمة التعبير البشري غالباً، "فعلى الرغم من أن التقنيات الرقمية تقدم اليوم خدمات تتزايد في طابعها الشخصي إلا أنها تبقى في جوهرها غير مكترثة تماماً بالأفراد وخصائصهم الشخصية. فلا قيمة للأفعال والتفاعلات الفردية إلا بقدر ما يمكن تتبعها رقمياً وقياسها وتحليلها" (Romele, 2024, p.8). ومن هنا تكمن المفارقة: كلما أمعن الكاتب أكثر في وضوحه وتنظيم أسلوبه اقترب نصه من خاتمة الذكاء الاصطناعي.

3. 2. التجربة الثانية

طلبنا من ChatGPT المجاني (من دون حساب، أو حساب من دون اشتراك) أن يكتب ستة مقالات مكونة من 12 إلى 15 فقرة، من دون الاستعانة ببرامج التي تضيف الطابع البشري على الذكاء الاصطناعي. جاءت النتيجة على برنامج iThenticate أنه صنف معظمها 100% ذكاء اصطناعياً.

✓ المقالة الأولى بعنوان: Globalization and Its Relationship with Artificial Intelligence

تم تصنيفها على أنها ذكاء اصطناعي بنسبة 100%، بتاريخ 10 شتنبر 2025

✓ المقالة الثانية بعنوان: Understanding AI Detectors: Function, Challenges, and Implications

تم تصنيفها ذكاء اصطناعياً بنسبة 100%، بتاريخ 10 شتنبر 2025.

✓ المقالة الثالثة بعنوان: Literature and Artificial Intelligence

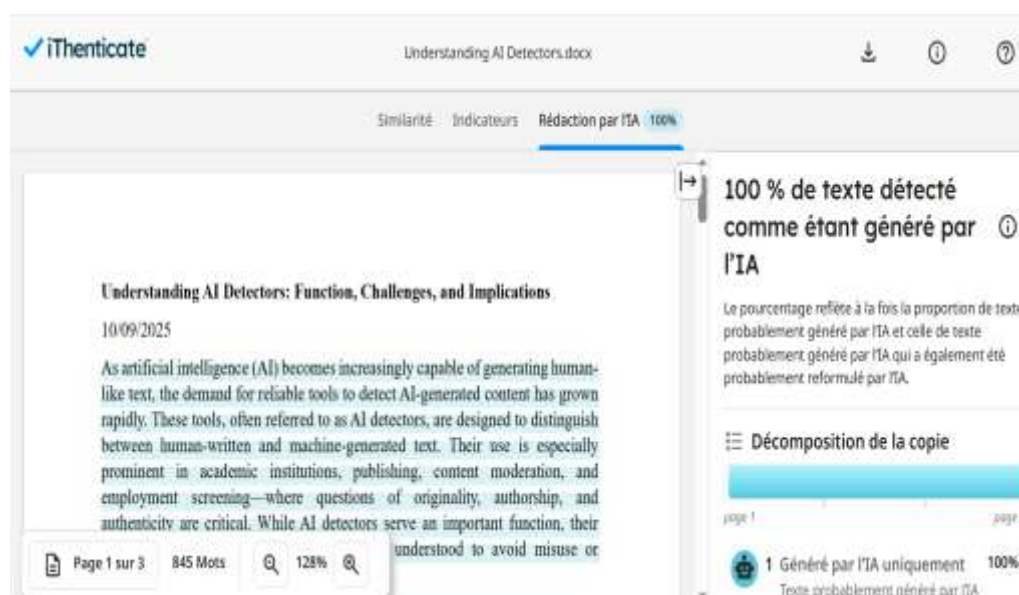
بنسبة 100%، بتاريخ 10 شتنبر 2025

✓ مقالة رابعة بعنوان: Artificial Intelligence and Sociology، تم تصنيفها ذكاء اصطناعياً بنسبة 99%، بتاريخ 10 شتنبر 2025.

✓ مقالة خامسة بعنوان: History and Sociology: Interwoven Threads of Human Understanding، تم تصنيفها ذكاء اصطناعياً بنسبة 100%، بتاريخ 11 شتنبر 2025.

✓ مقالة سادسة بعنوان: Artificial Intelligence and the Natural Sciences، نسبة الذكاء الاصطناعي 100%، بتاريخ 11 شتنبر 2025.

الصورة رقم 7: توضح كشف برنامج شركة تيرنيتي لحتوى الذكاء الاصطناعي المولد بنموذج ChatGPT



المصدر: فريق البحث، بتاريخ 10 شتنبر 2025

لكن بعد طلبنا من ChatGPT أن يعيد صياغة المقالات بأسلوب خصائص الكتابة البشرية، وجدنا أن النسبة انخفضت في النصوص التي قمنا بإعادة صياغتها، وهي كالتالي:

✓ المقالة الأولى بعنوان: Globalization and Its Relationship with Artificial Intelligence، تم اعتبارها ذكاء اصطناعياً بنسبة 100%، بتاريخ 10 شتنبر 2025.

✓ المقالة الثانية، بعنوان: W Understanding AI Detectors: Function, Challenges, and Implications، لم يتم التعرف عليها أو تحديدها %*، بتاريخ 10 شتنبر 2025.

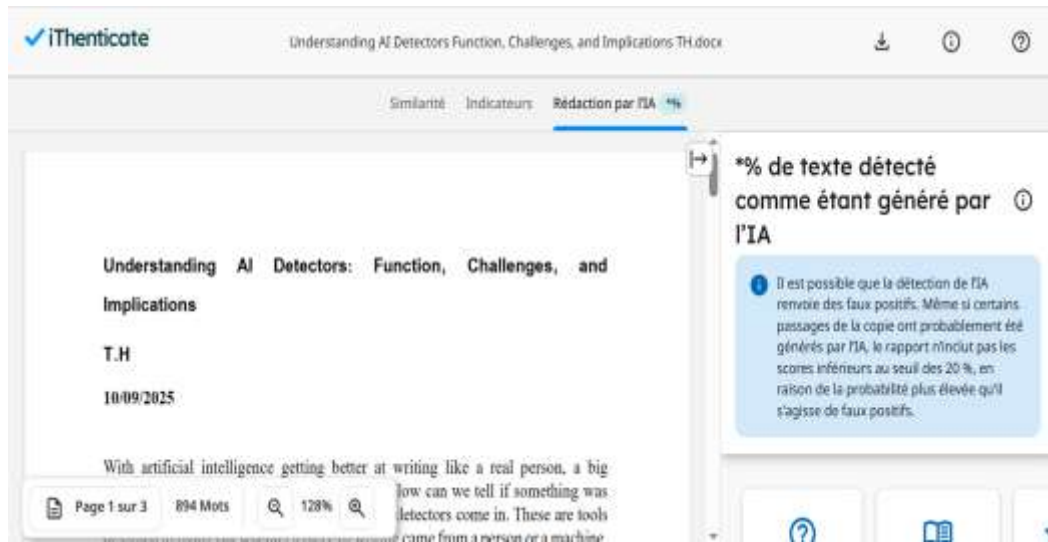
✓ المقالة الثالثة بعنوان: Literature and Artificial Intelligence، تم تصنيفها بنسبة 93% ذكاء اصطناعياً، بتاريخ 10 شتنبر 2025.

✓ المقالة الرابعة بعنوان: Artificial Intelligence and Sociology، تم تصنيفها بنسبة 23% ذكاء اصطناعياً، بتاريخ 10 شتنبر 2025.

✓ المقالة الخامسة بعنوان: History and Sociology: Interwoven Threads of Human Understanding، تم تصنيفها ذكاء اصطناعياً بنسبة 47%، بتاريخ 11 شتنبر 2025.

✓ المقالة السادسة بعنوان: Artificial Intelligence and the Natural Sciences، تم تخفيض نسبة الذكاء الاصطناعي إلى 0%، بتاريخ 11 شتنبر 2025.

الصورة رقم 8: توضح عجز برنامج ترينتين عن تحديد محتوى الذكاء الاصطناعي بعد إعادة صياغته.

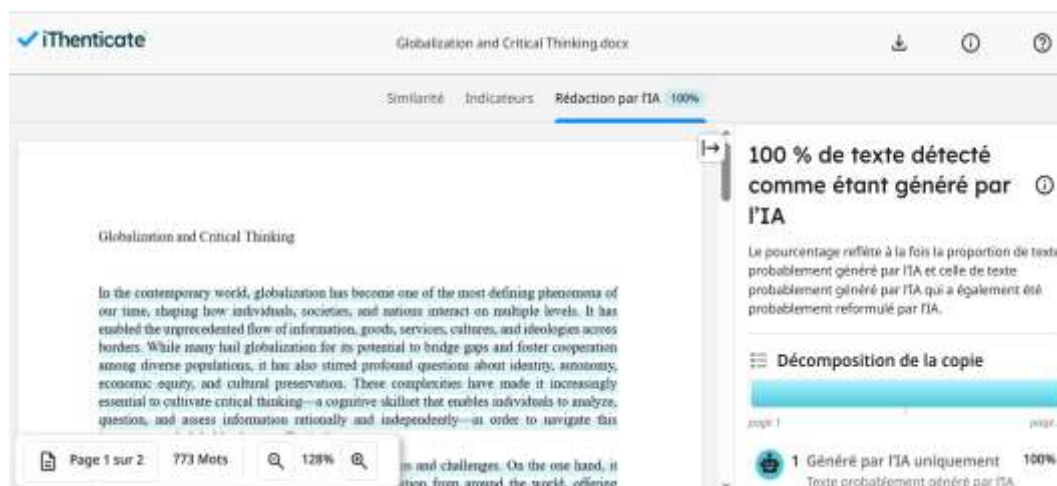


المصدر: فريق البحث، بتاريخ 10 شتنبر 2025

تفيد المعطيات أعلاه بشكل صريح بأن برنامج الكشف عن محتوى الذكاء الاصطناعي التابع لشركة ترينتين يمكن التلاعب به حتى من دون استخدام أي برنامج يعمل على أنسنة النصوص AI، على الرغم من أن الشركة قامت بتحديثه بتاريخ 27 غشت 2025؛ لكي يواجه النصوص المعادة صياغتها بالذكاء الاصطناعي. وهو ما يثبت عدم مصداقية البرنامج حتى مع النصوص الإنجليزية¹، زد على أنه يمكن للأفراد التحايل على مثل هذه الأدوات بمجرد إعادة الصياغة، أو إدراج المشاعر أو الحكايات، أو زيادة تنوع في الكلمات، أو إعادة إيقاع بنية النص. إذ يؤكد الباحث لوبي جيراري بأن "أدوات كشف الذكاء الاصطناعي يمكن أن تكون عرضة لهجمات مضادة، حيث تجري تعديلات بسيطة على النص المنتج بواسطة الذكاء الاصطناعي لخداع الكاشف. فعلى سبيل المثال، يمكن لإعادة الصياغة الطفيفة أو تغيير بنية الجمل أن تسمح بمرور المحتوى دون اكتشاف" (Girary, L., 2024, p.5). أما في حالة التلاعب باستخدام أحد برامج التي تؤنس محتوى الذكاء الاصطناعي، مثل: K(undetectable) أو أي برنامج آخر يهدف إلى أنسنة النصوص فإنه سوف يعطي نتيجة 0% أو 0%، وهذا ينطبق على جل البرامج التي تفشل في تحديد محتوى AI، ومن بينها iThenticate كما هو مبين في الصور الآتية:

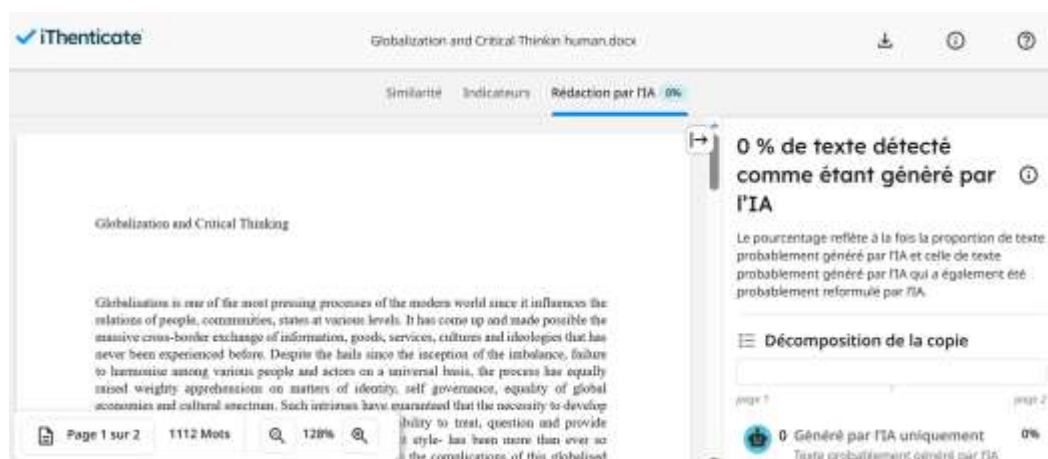
¹. تم إدراج اللغة الإسبانية في برنامج شركة ترينتين. وعلى غرار اللغة الإنجليزية، فإن البرنامج يقدم نتائج إيجابية وسلبية كاذبة بالنسبة إلى اللغة الإسبانية. وهو ما عبر عنه العديد من المستخدمين من أساتذة وطلبة على منصة ريديت. مما يظهر بأن المشكلة سوف تتفاقم مع اللغات الأخرى.

الصورة رقم 9: توضح مقالة بالذكاء الاصطناعي قبل إخضاعها لأحد برامج الأنسنة



المصدر: فريق البحث، بتاريخ 12 شتبر 2025

الصورة رقم 10: توضح نتيجة الأنسنة بأحد البرامج



المصدر: فريق البحث، بتاريخ 12 شتبر 2025

ومن جانب آخر عمدنا إلى مقالة صيغت بالذكاء الاصطناعي كلها عبر ChatGPT، دون إخضاعه لأي تغيير بشري أو برنامج للأنسنة، إذ يأتي برنامج (iThenticate) التابع لشركة تيرنيتين بتقديرات تجانب الحقيقة بالرغم من تحديثه الأخير في 27 غشت 2025 الذي يركز على اكتشاف النصوص الآلية المعاد صياغتها عبر برامج الأنسنة. ففي تقريره في أحد المقالات التي تبحث في علاقة السوسيولوجيا بالأدب الفرنسي وسم شطر من المقالة (فقرة ونصف الفقرة)

بأنها أُسست ببرنامج إعادة صياغة (36% مجموع الذكاء الاصطناعي: 24% مولد بالذكاء، و11% معاد صيغته بالذكاء)، والحق أن النص لم يخضع لأي إعادة صياغة قط. وقد تجلّى ذلك جلياً في الصورة المرفقة رقم (11). وهذا يبين عجز خوارزميات الكشف الراهنة وقصورها؛ لأنها تخلط أحياناً بين تشعب الأساليب، وتركيبية الجملة، وبين فعل الإنسان الحقيقي وقصديته؛ مما يوهن الثقة في نتائج تصنيفها، ويجعل أحكامها عرضة للنظر على الرغم مما أضيف للبرنامج من تحديثات أخيرة.

الصورة رقم 11: دليل بصري على خطأ خوارزميات الكشف عن إعادة الصياغة



المصدر: فريق البحث، بتاريخ 12 شتنبر 2025.

ولا مرية في أن أداء هذه البرامج قد يكون أقوى مع النماذج اللغوية الضعيفة أو عند مقارعتها لتلك البرامج الضعيفة التي تؤنسن، مثل: النماذج المجانية، وما شاكلها من النماذج القديمة. غير أنها قد تخطئ في الفصل بين النصوص البشرية، وما تم إنشاؤه من النصوص الآلية. إذ تبين البحوث التي عنيت بأدوات الكشف عن المحتوى AI أن "هناك تبايناً شديداً في قدرة الأدوات على تحديد النصوص، وتصنيفها بشكل صحيح، إما كنصوص مولدة بالذكاء الاصطناعي أو بشرية، إلا أن الغالب الأعم هو أن أدائها يكون أفضل أثناء تحديد المحتوى المولد بواسطة (GPT-3.5) مقارنة بالمحتوى المولد بواسطة (GPT-4)، أو ما ينتجه البشر من نصوص" (Elkhatat et al., 2023, p.14). وبالتالي فإن أداء هذه الأدوات غير مستقر، ويتفاوت حسب مدى تطور نموذج الذكاء الاصطناعي المستخدم في توليد النصوص. وهو ما تؤكد دراسة (Sadasivan et al., 2025) التي تقدم إثباتاً رياضياً، فكلما اقتربت النصوص المولدة آلياً من النصوص البشرية (كما يحدث مع تطور GPT-4 و GPT-5)، ينخفض معدل الكشف عن المحتوى بالذكاء الاصطناعي (Sadasivan et al., 2025, p.36)؛ إذ تصبح هذه البرامج سيئة جداً مع مرور الوقت. وبمعنى آخر، أنه مهما كان نوع البرنامج، ومهما تطور فإنه لن يتمكن من تقديم ضمان 100٪؛ مما سيجعله دائماً يقدم على ارتكاب الأخطاء في مستوى تصنيف النصوص، أو الفصل بينها؛ نتيجة تطور النماذج اللغوية الكبيرة، وبرامج الأنسنة التي تقترب في أدائها من اللغة البشرية.

ومع أن جلّ أدوات الكشف AI تدعي أنها قادرة على فضح المحتوى الذي جرى تنقيحه ببرامج الأنسنة؛ ومن خلال ما نلاحظه فإنها لا تزال تقع في أخطاء فادحة عند تمييز النص البشري من النص الآلي، وحتى من دون الاعتماد على

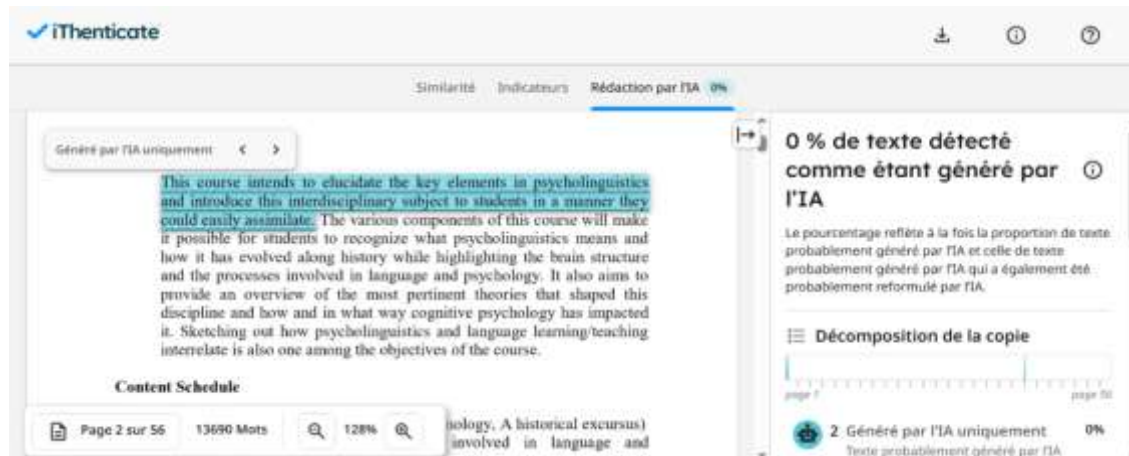
البرامج التي تؤنس. وتؤكد دراسة حديثة أن برامج كشف النصوص التي يصوغها الذكاء الاصطناعي يسهل خداعها بحيلة مبتكرة، تدعى إعادة الصياغة الخصامية (Adversarial Paraphrasing)، وهي صياغة موجهة تعتمد إلى انتقاء كلمات، وتراكيب تبدو في ظاهرها كأنها من نسج البشر. وبهذا الأسلوب يغدو النص الآلي مشابهاً لنص الإنسان لدرجة أنها تضلل مختلف الكواشف؛ حتى ما بلغ منها تقدماً وتطوراً أو ما يعتمد على العلامات المائية، إذ يتم كل ذلك من دون حاجة إلى تدريب زائد، أو من دون أن تنحط جودة النص الخطاطا يذكر (Cheng et al., 2025, p.1). كذلك يرجع ضعفها إلى تطور المولدات اللغوية إلى جانب أن الشركات التي تضيف على نصوص AI لمسة بشرية تطور بدورها من خوارزمياتها. وبالتالي فإن جوهر المشكلة يكمن في تلك الخوارزميات التي تتعلم من دون وعي، بالإضافة إلى أن مشكلتها تتباين من برنامج إلى آخر، وهو ما يجعلها غير موثوقة، فيختلط عليها التكرار اللفظي المعهود في كلام البشر، والأساليب الناعمة التي تنسجها الآلة؛ وذلك ما يوقعها في وسم كتابة البشر بأنها من نسج الذكاء الاصطناعي. فما يجعل من هذه البرامج والأدوات خطرة على البشر هو قدرتها على الانفصال عن تدخل البشر، إذ يشير بول هيمان إلى أنه "مع صعود التعلم الآلي (يسمى بشيء من التلطيف أو التخفيف الذكاء الاصطناعي)، أصبحت المعرفة الرقمية التي تنتج خوارزمياً أكثر انفصالاً عن التجربة الإنسانية. فالتصنيفات والإجراءات تولد الآن خوارزمياً وفق منطق تضعه الخوارزمية ذاتها، لا وفق ترميز مباشر من المبرمجين البشر. حيث يتطلب التعلم الآلي إعداد عدد هائل من المتغيرات الداخلية والروابط فيما بينها، ثم تزويد الخوارزمية ببيانات إدخال للتعلم... وتقوم الخوارزمية عبر عملية متكررة بتعديل هذه المتغيرات الداخلية والأوزان التي تربط بينها بهدف تحقيق أفضل تحويل ممكن من بيانات الإدخال إلى بيانات المخرجات. وبهذه العملية تنشئ الخوارزمية تصنيفاتها الخاصة، ومعرفتها الرقمية الخاصة بها؛ لتصل إلى أقوى الروابط الممكنة بين بيانات الإدخال والمخرجات" (Henman, P., 2021, p.25).

وبعبارة أخرى أكثر وضوحاً: تظهر التجارب ما يمكن أن تخلقه هذه الأدوات من توتر وقلق لدى الباحثين والطلبة. ولهذا لا يمكن في نظرنا الاعتماد عليها في الحكم ولو جزئياً؛ لأنها يمكن أن تكون ظالمة في حكمها. فمع تطور برامج الكشف عن محتوى AI وتعلمها من النصوص المولدة والنصوص المؤنسة، فإنها تبدأ في التعامل مع النصوص البشرية البسيطة أو المصقولة لغوياً على أنها نصوص آلية؛ وذلك لجرد شبهة الاصطناع، ولا سيما وأن النماذج اللغوية والبرامج المؤنسة باتت تقترب أكثر فأكثر من الأسلوب البشري. ويعزز هذا الأمر ما توصلت إليه دراسة هيربولد وزملائه (2023)، التي أجرت مقارنة واسعة النطاق بين مقالات دجها طلاب المرحلة الثانوية ومقالات أنجزها كل من ChatGPT-3، وChatGPT-4، فظهرت من نتائجها أنه في "حين أن نماذج ChatGPT تنتج جملاً أكثر تعقيداً، وتستخدم مزيداً من الأسماء المشتقة، يميل البشر بدلاً من ذلك إلى الإكثار من الأفعال المساعدة والتراكيب الإبتيمية، فالغنى المعجمي لدى البشر أعلى من تنوع ChatGPT-3، ولكنه أقل من ChatGPT-4. أما في توظيف أدوات الربط الخطائية فإنه لا وجود لفرق بين البشر وChatGPT-3. غير أن نموذج ChatGPT-4 يجري فيها أقل استعمالاً لأدوات ربط خطائية" (Herbold et al., 2023, p.8).

ومن هنا يمكننا الاستدلال بأنه من الصعب على خوارزميات الكشف أن تعتمد على مؤشرات أسلوبية تقليدية؛ لأن الفروق التي تعول عليها بين النصوص البشرية والآلية تنقلص بشكل لافت مع تطور النماذج. وإذا كان الباحث يريد أن يتجنب النتائج الإيجابية لأدوات الكشف التي تنتجها عن خطأ، فإن عليه أن يكتب بأسلوب أكثر تعقيداً بعيداً عن التبسيط اللغوي، والصرامة اللغوية، وهذا يتعارض مع الأسلوب العلمي الذي يتطلب تبسيطاً وصرامة لغوية ونحوية.

هناك العديد من الأعمال البشرية التي يتم تصنيفها على أنها ذكاء اصطناعي، كما أن هناك أعمالاً من إنجاز الذكاء الاصطناعي، يتم توصيفها على أنها نصوص بشرية (حسب الشهادات والتعليقات التي يتم نشرها على منصة ريديت)؛ وذلك بسبب اختلاف في التدقيق اللغوي، والأساليب التعبيرية وغيرها. وما يثير القلق هو أنه يمكن للخوارزميات أن تشير إلى جملة أو فقرة أو نص أو مقالة، أو حتى كتاب على أنه أسلوب ذكاء اصطناعي لمجرد تشكيك في أسلوب التعبير، وهو ما نلاحظه في الصورة الآتية:

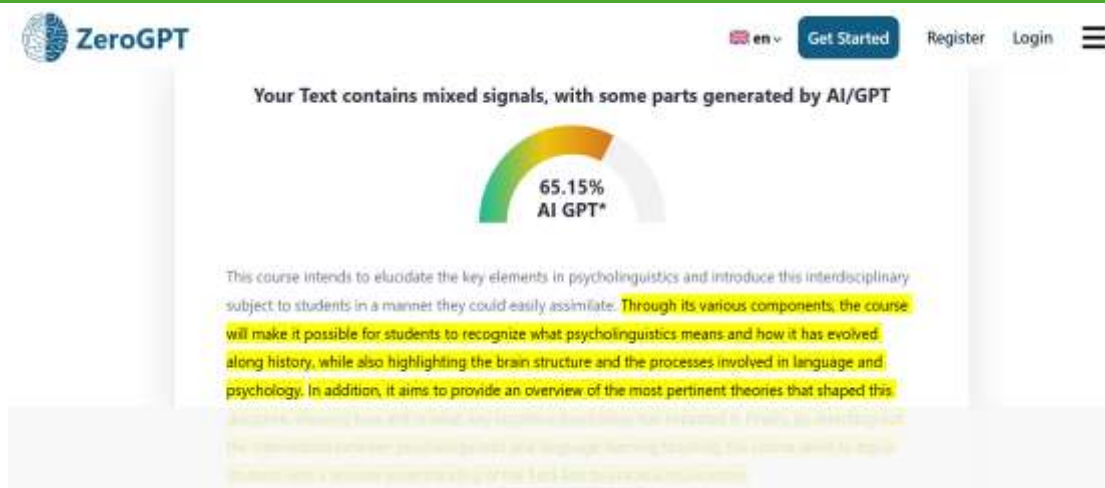
الصورة رقم 12: توضح أن برنامج Ithenticate رغم إقراره بأن النص بشري بنسبة 0% فإنه يشير خطأ إلى أن هناك AI



المصدر: فريق البحث، بتاريخ 6 شتنبر 2025.

فالصورة رقم (12)، تُوَمَّى إلى مقرر أستاذ جامعي في اللغة الإنجليزية بالمغرب، قام البرنامج الخاص بشركة تيرنيتين بإثارة الريب والتشكيك حول عبارة على أنها صياغة من الذكاء الاصطناعي مع أن البرنامج يظهر أن محتوى المقرر إنساني وخال من محتوى الذكاء الاصطناعي بنسبة 0%. كما أن المقرر منشور في سنة 2020 على موقع كلية الآداب والفنون في جامعة ابن طفيل بالقنيطرة، أي قبيل أن يتم الإعلان عن برنامج ChatGPT، أو النماذج الكبرى للذكاء الاصطناعي. وهذا يعني أن هامش الخطأ في هذه البرامج أكبر بكثير مما نتخيل؛ لأن نفس الفقرة تم اعتبارها ذكاء اصطناعياً على ZeroGPT بنسبة 65.15%، كما في الصورة رقم 13.

الصورة رقم 13: توضح أن نسبة الذكاء الاصطناعي في فقرة من المقرر الإنجليزي لأستاذ جامعي.



المصدر: فريق البحث، بتاريخ 6 شتنبر 2025.

قد يبدو من الصورة رقم (12) أن الخطر ضعيف بحجة أن نسبتها قليلة، مثل قول بعضهم: بأن "معدل القتل منخفض"، بينما هذه الأدوات تظهر بأن النصوص الأصلية في كثير من الأحيان منسوخة، أو أنها مولدة بالدكاء الاصطناعي. لكن لو كان أحد هؤلاء على وشك أن يقتل فإننا لا نعتقد على الإطلاق، بأن آخر ما سيخطر في ذهنه هو أن معدل جرائم القتل منخفض؛ لأن هذه البرامج تعتمد على التعلم الآلي الذي يحمل بدوره مخاطر تعزز من نتائج الإيجابيات الكاذبة، "فخوارزميات التعلم الآلي تعيد إنتاج وهم الموضوعية، بينما تخفي المقايضات بين الدقة والإنصاف داخل الخوارزميات. وقد طور باحثو العرق النقدي مفاهيم، مثل: "قانون جيم الجديد" (the New Jim Code)؛ للدلالة على أن خوارزميات البيانات الضخمة تعيد إنتاج أوجه اللامساواة القائمة، ومع ذلك تستمر في الظهور على أنها أكثر موضوعية أو تقدمية من الأنظمة التمييزية السابقة" (Ávila, et al., 2021, p.88).

إن مثل هذه الأمثلة دليل على الخطر الذي يمكن أن يواجهه الباحث أو الطالب على حد سواء، حيث يمكن للمقالات أو البحوث أو الأطاريح أو حتى مؤلفات أن ينظر إليها على أنها محتوية ذكاء اصطناعي؛ مجرد أنها لم تخضع للتعلم الآلي أو لعدم ثقتها في التمييز، أو لأنها خضعت للتدقيق اللغوي من طرف البشر، أو تم تدقيقه من طرف البرامج والمواقع اللغوية المساعدة (هناك أمثلة عديدة أكدها باحثون في المصادر التي اعتمدنا عليها في هذه الدراسة. على سبيل المثال، انظر إلى: Louie Giray, 2024). واليوم فإن العديد من الطلبة في العالم يستخدمون هذه البرامج في التدقيق اللغوي، أو تعديل النصوص، وهو أمر يعد مقبولا ومحموذا حتى اليوم لدى العديد من الجامعات حول العالم. إلا أن خوارزميات الكشف عن محتوى AI باتت تمتلك حساسية ضد أي أسلوب يتسم بلغة رسمية وموسومة ومصقولة أكثر من اللازم، أو من ذوي الأساليب الكتابية المميزة، أو لغة مجردة، إذ لا تفرق بين المحتوى المولد وبين التدقيق اللغوي. فهناك أدلة على أن تصنيف المحتوى البشري على أنه AI ليس نادرا (هناك أكثر من 20 شهادة على منصة ريديت تثبت ذلك، مثل الصورة رقم 14).

الصورة رقم 14: توضح رسالة تخرج أشير إليها بالدكاء الاصطناعي خطأً عبر "Turnitin"



المصدر: فريق البحث، بتاريخ 12 شتبر 2025

فما ورد في التجارب السابقة، يظهر قصور تلك الأدوات التي تدعي التمييز بين النصوص؛ وذلك بسبب تباين أسلوب الكتابة. فالكتاب وإن جمعهم لغة واحدة، فإنهم يختلفون في طريقة الكتابة، وأحياناً قد يتشابهون مع نماذج اللغة في الكتابة، إذ هناك من يكتب بلغة إخبارية أو تقريرية وهذا النوع من الكتابة يتسم بلغة جافة، وتكون في الغالب محايدة، وهي الأقرب إلى لغة الذكاء الاصطناعي، وفئة من الناس نجدتها تكتب بلغة أدبية تعبر عن الذات والعواطف والآراء، في حين نجد من يكتب بلغة تفسيرية تعليمية أو علمية، وأحياناً يتسم هذا النوع من الكتابة بالجفاف في التعبير، أو يكون متحذلقاً، إلخ. وكل من هذه الأساليب تختلف فيما بينها حتى في جودة النصوص، فالنصوص الأقل جودة، والتي يغيب فيها التنوع اللغوي يمكن اعتبارها ذكاء اصطناعياً. وبالتالي يصعب على هذه البرامج أن تكون دقيقة أمامها، فقد وجدت دراسة (Doughman et al., 2024) "أن هذه المصنفات أي أدوات الكشف عن المحتوى، تبدي حساسية عالية لبعض السمات اللغوية، مثل: توزيع الظروف (adverbs)، وطول الجملة، وسهولة القراءة. وظهر أنها تفرط في التعويل على علامات الترقيم والفجوات البيضاء. كما تدل النتائج إلى أن قواعد البيانات الحالية للنصوص المولدة ليست حصينة بما يكفي أمام التبدلات الأسلوبية أو تغير المجالات، وأنها أوهن بشكل خاص حين تطبيقها على الكتابة البسيطة، مثل: الواجبات المدرسية" (p.5). إلى جانب أن "التصنيفات الآلية لا تعمل في فراغ اجتماعي-ثقافي، ولا هي مجرد نتائج إحصائية منفصلة عن الواقع الاجتماعي. فالعقل العملي لأنظمة التعلم الآلي هو في حد ذاته نتيجة لتصنيفات بشرية متكررة: تلك التي يضعها منشقو الآلات والمجسدة كـ(إله في الآلة)، وأيضا تلك التي يضعها مدرسو الآلات والمحفورة في الأنماط البيانية التي تشكل ما يسميه ماسيمو إيرولدي بـ الهايتوس الآلي (Machine Habitus) (Airoldi, 2022, pp.81-82).

4) الذكاء الاصطناعي بين أثره في اللغة ومستقبله الأكاديمي

منذ القدم والتقنية قوة اجتماعية وتاريخية، تحدث تأثيراً في اللسان والثقافة بدءاً بالكتابة والمطبوعة، ومروراً بتكنولوجيا الاتصال الجماهيري ووسائل التواصل الاجتماعي (Castells, M., 2010). واليوم فإن الخوارزميات والذكاء الاصطناعي ولا سيما النماذج اللغوية الكبرى تعيش مرحلة فارقة في التطور التكنولوجي، فلم تعد وظيفتها مقصورة على تلقي اللغة البشرية، فقد تجاوزت ذلك من خلال إعادة بنائها وصياغتها، وردها إلينا في أشكال جديدة عبر التنشئة الخوارزمية التي تخلق أنماطاً جديدة في التواصل والتفاعل عليها. وهذا معناه بأن "التحكمية الخوارزمية تقوم على اقتراح نماذج إدراكية لتوجيه سلوك المبحرين السيبرانيين؛ لإعادة إنتاجها من خلال التفاعلات والمضامين التي يقومون بنشرها على

شاشات الفضاء السيبراني ("شكدالي، 2024، ص 77)، أو من خلال التفاعل مع نماذج ذكاء الاصطناعي. إذ أفضى هذا المسلك إلى ما يشبه حلقة ثقافية لا تنتهي، إذ ترتوي فيها التكنولوجيا من مصادر الإنسان، ثم تعود إلى الإنسان فتخلق بذلك ملامح لغة هجينة، تتراوح بين سجية البشر وألفاظ الخوارزميات.

في هذا المضمار، تلموح في الأفق نقطة مهمة لا يستجاز إغفالها، إذ تؤكد دراسات حديثة تعنى بالذكاء البشري والذكاء الاصطناعي بأن الاستخدام المتواصل لمنصات التواصل الاجتماعي وتطبيقات الذكاء الاصطناعي يجعل من الأفراد أسرى لتأثير الخوارزميات (Fisher, M., 2022). فالانغماس اليومي في استخدامها، يجعل الأفراد يكتسبون لغة الذكاء الاصطناعي بصورة غير واعية، وهو ما يترك بصمته العميقة عليهم على مستوى أنماط التفكير وأساليب الكتابة والتعبير والكلام أيضاً؛ مما يتشكل لديهم في النهاية ما يمكن تسميته بالهايتوس الخوارزمي الذي تخلقه تلك الخوارزميات. فقد كشف باحثون بجامعة فلوريدا الأمريكية في دراسة نشرت مؤخراً بأنه في الحوارات العلمية والتكنولوجية اليومية، نحن فعلاً نقتبس لغة الذكاء الاصطناعي أو بالأحرى، نفس الكلمات التي نسمعها من ChatGPT؛ فصارت هذه اللغة تدخل في كلامنا كأسلوب جديد. وعلى سبيل المثال، تظهر دراسات جديدة أن ألفاظاً معينة غدت أكثر شيوعاً في الخطاب البشري بعد ظهور أدوات، مثل: ChatGPT. إذ لاحظ ياكورا وزملاؤه (2025) زيادة بيئة في استخدام كلمات، مثل: (delve)، و (comprehend)، و (boast) في الخطاب الأكاديمي والمحاورات، الشيء الذي يظهر أن تفضيلات النموذج اللغوي قد انسابت إلى اللغة البشرية اليومية (Yakura et al., 2025, p. 1). ولم يقتصر هذا التأثير عند حدود النصوص المكتوبة، وإنما مس حتى الخطاب الشفوي العفوي، إذ جاء في الدراسة "أن التحولات اللغوية المدفوعة بالذكاء الاصطناعي تمتد إلى ما هو أبعد من النصوص المكتوبة لتصل إلى التواصل الشفوي العفوي" (Ibid, p.8). وقد تبين أن هذا التأثير أكثر وضوحاً في حقول العلوم والتكنولوجيا والتعليم والأعمال مقارنة بالدين أو الرياضة (Ibid, p.9)، وأن بعض هذه الكلمات يشهد نمواً سنوياً يتراوح بين 25% و 50% (Ibid, p.5). وتخلص الدراسة إلى أن الظاهرة ليست تقليداً سطحياً فحسب، بل تعبر عن عملية استبطان لأنماط اللغوية التي تنتجها الآلة داخل الوعي البشري، وهو ما يكشف عن بروز الذكاء الاصطناعي كفاعل اجتماعي، يسهم في إعادة تشكيل الثقافة، ويهدد بفرض تجانس لغوي يضعف التنوع الرمزي.

وبالمثل، تؤكد دراسة أخرى قامت بتحليل أكثر من 22 مليون كلمة من برامج بودكاست عن العلوم والتكنولوجيا؛ والنتيجة أنها وجدت تغيرات في النظام اللغوي البشري نفسه، وهناك احتمال معقول أن الذكاء الاصطناعي بدأ يؤثر بشكل مباشر في النظام اللغوي البشري (وهذا التفسير تدعمه الملاحظة بأن الكلمات التي يفرط الذكاء الاصطناعي في استخدامها هي نفسها التي تزداد في الاستخدام البشري) (Anderson, B., Galpin, R., & Juzek, T., 2025, p.7). فكلمات، مثل: "garner" (يجمع)، "delve" (يتعمق)، و "intricate" (معقد)، و "underscore" (يؤكد)، كلها أصبحت جزءاً من الكلام اليومي.

أما المولدات الذكاء الاصطناعي، مثل: ChatGPT فقد أصبحت تتبنى اللغة البشرية، وتقترب من التفوق على الأسلوب البشري. فهذه المولدات تحمل معها المجتمع بتناقضاته. إنها تقنيات تواصلية قادرة على التكيف مع التوقعات الثقافية والاجتماعية للمستخدمين بفضل تراكم المعرفة الإنسانية المؤدجلة عبر البيانات. فهي تكتسب ميولات اجتماعية متموضعة لتصنيف وتقييم العالم، موسومة بالطبقة والجنس أو الإثنية (Rama & Airoidi, 2024, p. 3)، وهو ما يشكل الهايتوس الخوارزمي. فقد أصبحت نماذج الذكاء الاصطناعي التوليدي التي يعتمد عليها الفرد في الولوج إلى

العالم العابر للحدود تؤثر في ذوقه وتفكيره واختياراته، كما تعدل ما يتمتع به من ذكاء في تعلم اللغة، وتوسيع مداركه؛ بسبب تدفق المعلومات وسهولتها، فهي تسهم في إعادة بناء شخصية الفرد، وتمكنه من إبراز معارفه وأفكاره، ومن فهم العالم الذي ينتمي إليه أيضا. غير أن ما يميز هذه النماذج التوليدية هو أنها منحازة أكثر إلى اللغة الإنجليزية بوصفها اللغة المحورية في التطوير والاستخدام التجاري؛ إذ تتوفر أفضل جودة لهذه النماذج لمتحدثي الإنجليزية، في حين تأتي مخرجات أقل جودة وكلفة أعلى عند استخدام لغات أخرى. وكما يشير ماسيمو إيرولدي بالقول: "يمكن للآلات أن تتعلم كيف تتكلم مثل البشر، وكيف تكتب مثل الفلاسفة، وكيف تقترح الأغاني مثل خبراء الموسيقى. ويمكنها أيضا أن تتعلم كيف تكون متحيزة جنسيا مثل رجل محافظ، وعنصرية مثل متفوق أبيض، وطبقية مثل متعجرف نخوي" (Airoldi, 2022, p.x). وهذا يعني أن الآلات والنماذج الذكية يمكن أن تؤثر في الأفراد في أسلوب التفكير والكلام، وحتى في طريقة الكتابة، الأمر الذي يجعلنا نستنتج بأن الذكاء الاصطناعي يغذي اللغة البشرية، واللغة البشرية تغذي الذكاء الاصطناعي.

يشعر هذا التحول في فتح باب السجال حول تخوم ما هو طبيعي وما هو مصنوع في الخطاب البشري؛ إذ تغدو النصوص التي تولدها الذكاء الاصطناعي جزءا من الممارسة اللغوية المعتادة، ما يعني أن الخطاب يصبح مزيجا بين نطق البشر بصوت الذكاء الاصطناعي إلى درجة يصعب الفصل بينهما. فهناك من "يحتاج بأننا ندلف اليوم، أو لربما غدونا بالفعل في صميمه، إلى عصر الإنسان والآلة الذي تتجسد فيه تصورات المستقبل والقصص الخيالية حول العلم بفعل ما استجد في الذكاء الاصطناعي، وتعلم الآلة والتقنيات القابلة للارتداء، وتحليل البيانات في الوقت الراهن" (Kelly-Holmes, 2024, p.5).

وتبين ماركل بأن هذه التحولات لم تقف عند حدود الاستعمالات اليومية للسان، فقد جاوزتها لتمس حتى طرائق التفكير والممارسة ذاتها؛ إذ تقول "إن تكنولوجيا اللغة تتوغل في سياقات خطيرة كالتوظيف وشؤون الشرطة والهجرة، فلا يتبدل مسلكنا في استعمال اللغة فحسب، إنما يتبدل معه تصورنا لها. ويمكن ادراج هذه التغيرات تحت ما أسمته هي بالإدارة اللغوية الخوارزمية" (Markl, 2024, p. 2). وهو ما يفصح عن انتقال سلطة اللغة من الأطر المعهودة إلى فضاءات تقنية خوارزمية تعيد تشكيل الميادين الرمزية والثقافية. وتندر ماركل بأن القرارات المتصلة باللغة ذاتها أمست مرهونة بشكل مباشر بأيدي كبار الفاعلين في الاقتصاد قائلة: "إن قرارات جوهرية تخص اللغة من قبيل انتقاء بياناتها، ورعايتها، وضبط جودة النظام، وما يفرض عليه من قيود في سلوكه؛ تتخذها شركات خاصة فاعلة" (Ibid, p.6). وعليه، إذا كان الذكاء الاصطناعي أداة معرفية، فإنه آلية اجتماعية اقتصادية كذلك؛ لأنه يسهم في استدامة علاقات القوة والهيمنة وعدم المساواة؛ ولأنه يترك أثرا مباشرا في نسق النظام اللغوي والثقافي العالمي. وهو ما يعني بأن الإدارة التقنية اللغوية التي تؤسسها الشركات تتأسس بشكل جوهرى على أطر سياسية واقتصادية من أجل إعادة رسم الخريطة الرمزية للغات؛ ما يجعلها تعزز من مكانة لغات كبرى، مثل: اللغة الإنجليزية، مقابل إقصاء لغات مهددة أو مهملة في الفضاء الرقمي. فلا يقتصر أثر الذكاء الاصطناعي على أسلوب الكتابة والتعبير، ولكنه يسهم في إعادة هندسة الموازين الرمزية التي تنبني عليها الثقافة والتواصل في المجتمعات الإنسانية.

وهذا الواقع يجعل من الجامعات العربية والمراكز الأكاديمية مطالبة باستخدام الذكاء الاصطناعي من أجل تطوير البحث والنشر العلمي وإنجاز البحوث. ويلاحظ بأنه على الرغم من أن هذه الأدوات متوفرة للباحثين، ويمكن أن

تساعدهم في إنجاز بحوثهم، وتسريع من عمل أطاريحهم فإن منتجهم على مستوى إنتاج البحوث يبقى بطيئا وضعيفا. إذ تظهر بيانات بليومتريّة حديثة أن مساهمة البلدان العربية لا تتجاوز 2.72% من النشر الطبي العالمي (2017-2023)، مع فجوة لافتة في التمويل مقارنة بالمتوسط الدولي، ما ينعكس على التأثير والاقتباسات (Almuhaideb et al., 2024). وتؤكد اليونيسكو أن نسبة الإنفاق على البحث والتطوير تبقى منخفضة في معظم دول المنطقة، بما يحد من تحويل التعليم العالي إلى معرفة منتجة (UNESCO, 2021). في المقابل تبرهن التجارب العشوائية المضبوطة أن استخدام أدوات توليدية، مثل: (ChatGPT) يساعد في (انخفاض الزمن 40% وارتفاع الجودة 18%)، كما يحسن كفاءة إنجاز المهام الأكاديمية (Noy & Zhang, 2023). وبالتالي، فإن تعزيز الحضور العربي على مستوى الاستخدام يبدأ من إعادة تعريف "الاستخدام" بوصفه ممارسة معرفية واقتصادية وسياسية للغة والمعرفة الأكاديمية.

ويجب على سياسات البلدان العربية في مجال التعليم أن تفهم بأن التحول اللغوي الذي يصاحب الإنسان في علاقته بالذكاء الاصطناعي لن يتوقف، إذ إنه سيمس حتى اللغة العربية على مستوى التفكير والكتابة وطريقة الكلام، لأنها لن تكون قادرة على منع استخدام الذكاء الاصطناعي، وما عليها إلا أن تساعد المتعلمين والباحثين والإداريين على استخدامه وتوظيفه من أجل تعزيز الإبداع وتشجيع الابتكار. وذلك بالرغم من أن الالتباس "والارتباك المتكرر مع هذه التقنيات في سياقات متعددة يمكن أن يترك آثارا مسطحة على فهمنا لذواتنا وبنائنا لها" (Romele, 2024, p.8). بيد أن الذكاء الاصطناعي أضحى لكي يبقى موجودا وحاضرا في تفاعلاتنا وممارساتنا؛ لأن التغيرات الجذرية التي أحدثتها أدوات الذكاء الاصطناعي في المشهد المعرفي واللغوي تستلزم تفكيراً جديداً في أسلوب الكتابة، وتعامل المؤسسات الأكاديمية معها. إذ نتوقع أنه ابتداء من عام 2026 ستجد الجامعات والمراكز العلمية نفسها أمام واقع متحول، يتطلب التوازن بين دور الإنسان والذكاء الاصطناعي في إنتاج المعرفة. فقد أصبحت هذه الأدوات جزءاً لا يتجزأ من البيئة الأكاديمية، وعلى سبيل المثال، جاء في تقرير أصدره مجلس التعليم الرقمي، وهو عبارة عن استطلاع شمل طلاباً من 16 دولة تحت عنوان (Digital Education Council Global AI Student Survey 2024)، أن 86% من طلاب التعليم العالي يستخدمون الذكاء الاصطناعي في دراستهم، ويتصدر "شات جي بي تي" قائمة الأدوات الأكثر استخداماً بنسبة 66%. حيث تتنوع الاستخدامات بين البحث عن المعلومات 69%، والتدقيق اللغوي 42%، وتلخيص النصوص 33%، وقد أشار أكثر من نصف الطلاب 54% إلى استخدامهم الأسبوعي لهذه الأدوات، و24% بشكل يومي (Digital Education Council. 2024). كما تبين إحصاءات عام 2025 في المملكة المتحدة أن استخدام الطلاب لمنصات الذكاء الاصطناعي بلغ مستوى شبه شامل، إذ أفاد نحو 92% منهم بأنهم يستخدمونها بشكل منتظم في دراستهم وأبحاثهم، وتتمثل الاستخدامات الرئيسة للذكاء الاصطناعي التوليدي في شرح المفاهيم، وتلخيص المقالات، واقتراح أفكار بحثية؛ غير أن نسبة ملحوظة من الطلاب 18% أدرجوا نصوصاً مولدة بالذكاء الاصطناعي مباشرة في أعمالهم (Freeman, J. 2025, p.1). وهذا يؤكد بشكل أو بآخر أن المجتمعات الأكاديمية في تفاعلها اليومي مع الذكاء الاصطناعي سوف يعزز لديها الهايتوس الخوارزمي، وسوف يحول جزءاً من اللغة البشرية إلى لغة هجينة. إذ يشير إيرولدي لهذا الأمر عندما يتحدث عن الاستخدام المتكرر للآلات والتقنيات التي تقوم على الخوارزميات، ويقول: "إذا حدث هذا بانتظام، فإن ذوقي الفعلي وممارساتي الاستهلاكية تصبح، في الواقع، نواتج حسابية، تملئها ميول خوارزمية مستخلصة من آثار سلوكي السابقة" (Airoldi, 2022, p.84).

كما يؤكد الخبراء والقائمون على سياسات التعليم العالي على أن استعانة الطلبة بالذكاء الاصطناعي قد أضحت أمراً محتوماً، وإذا فائدة في كثير من المجالات؛ وهو ما يوجب على المؤسسات أن تتحول عن منطق الحظر إلى منطق الإرشاد والمواءمة الإيجابية (Freeman, J. 2025, p.2). وبمعنى آخر، يجدر على الجامعات ومراكز البحوث العلمية أن تكون منفتحة على أدوات الذكاء الاصطناعي، وأن تحذر من التعامل السيء معها، مثل: الحظر على الباحثين والطلبة؛ لأن منعها سيؤدي إلى ضرر كبير على الحقل الأكاديمي برمته. وهو ما يستوجب من القائمين على السياسات التعليمية إعادة التفكير في إستراتيجيات جديدة مبنية على إدماج هذه الأدوات في عملية التعلم والبحث العلمي على نحو مسؤول.

لقد بدأ الباحثون يرصدون تحولاً في معايير الكتابة الأكاديمية، وهوية المؤلف بفعل ظهور ما يسمى "المؤلف الخوارزمي". ففي دراسة حديثة أجريت عام 2025، حلل الباحثان (ماريا غريتشكي وجدةون ديشون) الكيفية التي تعيد بها النماذج اللغوية الضخمة تشكيل إنتاج المعرفة في الجامعات، وتوصلاً إلى أن هذه النماذج لا تعمل مجرد أدوات مساعدة على الكتابة، وإنما تقوم بدور بنيوي في توحيد معايير الممارسة الأكاديمية، وخلق مواقع وأدوار جديدة في عملية إنتاج المعرفة (Gretzky, M., & Dishon, G., 2025). كما أن هناك إطارين يسودان نظرة الأكاديميين إلى نماذج الذكاء الاصطناعي: يرى البعض هذه النماذج بمثابة "مكتبة بابل" حديثة، أي مستودع معرفي شمولي، وهو ما يعزز منظورها تكنوقراطية للمعرفة؛ في حين يتعامل آخرون معها كأداة عامل تفاعلي ديناميكي ضمن عملية الكتابة، ويضيفون عليها صفات شبه بشرية في وصف أدوارها (Ibid, p. 338). ورغم تباين هذه التصورات، إلا أنها تشترك في الإقرار بأن دخول الذكاء الاصطناعي في نسيج الثقافة الأكاديمية قد غير قواعد اللعبة. فقد قماهى التأليف البشري مع الخوارزمي ضمن علاقة تشاركية، يتخللها قيام البشر بمراجعة وتصويب ما تنتجه الآلة بشكل مستمر (على نحو ما وصفه بعض الباحثين بـ "التأليف المشترك" بين الإنسان والآلة) (Gretzky, M., & Dishon, G., 2024, p.E38).

إن الذكاء الاصطناعي في عصرنا يعيد صوغ مفهومي الإبداع والتحرير في الحقل الأكاديمي عبر ما يعرف بالتنشئة الخوارزمية (algorithmic socialization). إذ تغدو النماذج اللغوية وسائط تثقيفية تصوغ وعي الكاتب والناشي من جديد، وتلقنهم في طيات آليتها أسلوب الكتابة وسمت الكاتب وما هو القول الجيد أو المقبول في عالم التأليف. وكما يشير كلٌّ من غريتشكي وديشون إلى أن "نماذج الذكاء الاصطناعي تقوم بتعميم التوقعات الاجتماعية في الوسط الأكاديمي، وتلعب دور مترجمين اجتماعيين الذين يجسرون الهوة بين فهم الأفراد والشفرة الاجتماعية الضمنية للأكاديمية" (Gretzky, M., & Dishon, G., 2025, p.347). وهو ما يشير إليه ألبرتو روميل بالقول: "إن الآلات الرقمية هي آلات للعادة (habitus machines)؛ لأنها تنتج بجد واستقلال ضرباً من التصنيفات والفئات الاجتماعية؛ تلك التي تعود في غالبيتها إلى ما أوجده البشر من قبل. فعندما تتحول إلى أشكال خوارزمية، فإنها تصبح أمراً يستقر في أعماق الأفراد، ويتجسد في سلوكهم" (Romele, A., 2024, p.3).

وتكشف دراسة غريتشكي وديشون أن هذه النماذج اللغوية يتعدى أثرها توليد المحتوى فقط؛ لأنها تنقل المعرفة الاجتماعية الضمنية، وتقوم مقام التنشئة البشرية في عملية تكوين الأكاديميين، ولا سيما في بعض المهارات، مثل: تحرير الرسائل المهنية، وتقديم التوصيات. وكذلك تدبير الحياة الأكاديمية بين الزملاء، فالمعرفة المكتسبة التي كانت تتم عبر التدرج في الثقافة الأكاديمية، قد غدت في يومنا هذا متاحة من خلال وسيط خوارزمي (Gretzky, M., & Dishon, G., 2025, p.349). ولهذا فإن المعضلة الحقيقية تنصرف عن كيفية منع هذه التقنيات إلى البحث عن سبيل

لاستثمارها وتوجيهها وجهة تخدم الإنصاف المعرفي والتنوع الثقافي في ميدان العلم. خصوصاً وأننا نقف على أعتاب عصر جديد متداخل، يتضمن قدرات الإنسان والخوارزمية في كتابة المستقبل، وهو ما أكدته تقرير شركة تيرنيتين نشر بتاريخ 20 نوفمبر 2024 على موقعها بأنه في سنة 2025، ستكون الكتابة المعاصرة مزيجاً بين الكتابة الآلية والكتابة البشرية، وأنها لن يصبحا خصمان بل متعاونان. وأن الكتابة بالذكاء الاصطناعي سوف تعزز من جودة الكتابة البشرية، من حيث التنظيم والدقة والجاذبية الأسلوبية وغيرها، وقد لوحظ أن استخدام AI مع الأسلوب البشري يجعل من النص أصيلاً ومتميزاً (Turnitin, 2024, November 20).

إن هذا التصور الجديد حول المقاربة الأمثل للكتابة الأكاديمية سيكون له تأثير كبير في الجامعات والمراكز العلمية؛ إذ سيدفعها إلى العمل أكثر من أجل تحقيق توازن أكبر بين البشر والذكاء الاصطناعي؛ نظراً لما تشهده منصات الذكاء الاصطناعي التوليدي من تطور وابتكار مستمر، مع ارتفاع كبير في عدد المستخدمين لهذه النماذج اللغوية الذكية، وخصوصاً من طرف الباحثين والطلبة، الأمر الذي يفيد بأن فكرة حظر هذه الأدوات لم تعد ممكنة على الإطلاق؛ لأن هناك تحوفاً من أن يؤثر ذلك في البحث العلمي، خصوصاً على مستوى الإنتاج والنشر العلمي. وكما جاء في الدراسات السابقة حول تأثير الذكاء الاصطناعي في اللغة البشرية، والتقرير الأخير لشركة تيرنيتين، فإن العالم سيتجه نحو التأليف المشترك لا محالة، إذ سيتعين على الكاتب أن ينجز مقالته أو بحثه أو أطروحته بالاعتماد على النماذج اللغوية للذكاء الاصطناعي؛ بهدف تطوير مهارات جديدة، والبحث عن المصادر العلمية، وتحسين مستوى اللغة لديه، مثل: تصحيح الأخطاء النحوية واللغوية، وغيرها.

خلاصة الدراسة ونتائجها

تدعونا هذه الدراسة المعززة نظرياً وتجريبياً إلى توسيع رؤيتنا للذكاء الاصطناعي واللغة البشرية، فما يحصل اليوم من تغيرات في استخدام الذكاء الاصطناعي عند الأفراد والمجتمعات الأكاديمية، يبين بأن العالم يعيش تعقيداً على مستوى الثقة في أدوات الذكاء الاصطناعي وبرامجه، والتي تعدنا بتحقيق العدل والمساواة للجميع. خصوصاً وأن الخوارزميات في هذه الأنظمة التي تتعامل مع المنتج البشري لا تفكر ولا تفهم القصدية، وإنما تعتمد على التعلم الآلي. وهو ما يضع المجتمعات الإنسانية ومؤسساتها الأكاديمية أمام مسؤوليات فكرية وأخلاقية كبيرة، تتمثل في إعادة بناء العلاقة بين الإنسان وأنظمة الذكاء الاصطناعي من حيث مساءلة القيم والمعايير التي تبنى عليها هذه الأنظمة. ولهذا فإن الدراسة تأتي بمجموعة من النتائج هي كالآتي:

أولاً، تكشف الدراسة عن قصور وخلل في أدوات الكشف عن المحتوى المولد بالذكاء الاصطناعي، فقد أبانت مراجعتنا للأدبيات والشواهد على منصة ريديت، وما تم إصداره من بيانات رسمية لمواقف الجامعات الكبرى من هذه البرامج أن هذه الأدوات كثيراً ما تقدم نتائج إيجابية غير صحيحة، فلا تميز بين النصوص البشرية والنصوص المولدة بالذكاء الاصطناعي، وهو ما حدا بالجامعات الأمريكية والكندية والأسترالية إلى وقف العمل بميزة الكشف أو تعطيلها، وأن هذه الأدوات تمارس العنف الرمزي الذي يكرس هيمنة اقتصادية ومعرفية تستتر تحت عباءة الدقة العلمية والموضوعية، وهو ما يستوجب إعادة نظر في الاعتماد عليها من طرف الأكاديميين بشكل جذري.

ثانياً، تظهر نتائج الدراسة انحيازاً لغوياً وعرقياً أصيلاً في خوارزميات الكشف عن محتوى الذكاء الاصطناعي، إذ إنها تميل إلى وسم كـ تابيات غير الناطقين الأصليين أو النصوص ذات البنية البسيطة بأنها مولدة آلياً، في حين يقل خطؤها مع

النصوص التي تدنو من المعايير الأنجلوساكسونية الأصلية، ومن خلال الأدبيات والتقارير تؤكد هذه البرامج انخيازاً ضد جماعات عرقية، مثل: السود واللاتينيين. وتدل هذه النتيجة على أن أدوات الكشف لا تعمل بمعزل عن محيطها الاجتماعي، فهي تعكس البنى الاجتماعية والثقافية القائمة، وتعيد إنتاجها في الحقل الأكاديمي، فتغدو أداة لإدامة اللامساواة، لا لتحقيق العدل والمساواة.

ثالثاً، قُمِيط الدراسة من خلال التجارب التي أجريتها على برنامج آي إن تي كيت التابع لشركة ترينيتي وما شابه من برامج هشاشة في تصنيف النصوص وتمييزها، كما أظهرت ضعفاً كبيراً أمام أساليب التحايل. فقد تم تصنيف نصوص بشرية لناطقين أصليين على أنها آلية بنسبة كبيرة تصل حتى 100%. في حين استطاعت نصوص آلية أعيد صياغتها أن تمر دون أن تثير الشك. وحتى التحديثات التي طرأت في 27 غشت 2025 لم تفلح في أمام البرامج التي تؤنس النصوص. وهذه النتائج برهان على أن التطوير التقني وحده قاصر عن ضمان موثوقية النصوص ودقتها.

رابعاً، أبانت الدراسات والبحوث في مجال الذكاء البشري والاصطناعي، وبلاستناد إلى بعض البحوث والدراسات في السوسيولوجيا بأن الذكاء الاصطناعي أضحى من الكيانات الاجتماعية المؤثرة فعلاً في الصعيدين الأكاديمي والثقافي؛ وذلك بسبب عملية التنشئة الخوارزمية التي تعيد صوغ بنية التعبير الإنساني. فهناك زيادة ملحوظة على مستوى استخدام المفردات، والجمل، والتراكيب ذات الصلة بنموذج ChatGPT في اللغة البشرية اليومية، بما فيها اللغة الأكاديمية وغيرها؛ لأن العالم التكنولوجي، ولا سيما عالم الذكاء الاصطناعي هو عالم اجتماعي بالدرجة الأولى، ويقوم على التفاعل بين الأفراد والجماعات وحتى المؤسسات. وهذا ما يؤكد من فرضية الهايتوس الخوارزمي الذي يقوم على صياغة طبائع الكتابة والتفكير والتعبير بأسلوب الذكاء الاصطناعي.

من هنا فإن سياسات الحظر والمجاجة لهذه الأدوات لن تجدي نفعاً؛ نظراً للتداخل اللغوي بين الإنسان والآلة، بالمقابل فإن السبيل الأمثل هو بناء الشراكة مع هذه الأدوات الذكية لا الانفصال عنها، وذلك من خلال العمل على دمجها في الممارسة الأكاديمية بما يحقق من عملية الابتكار وتعزيز الإبداع، ومراعاة القيم والمعايير العلمية. إذ نستطيع القول إن مآل البحث العلمي والكتابة الأكاديمية في الأعوام القادمة سوف يتجه نحو شراكة معرفية خلاقة بين الإنسان والذكاء الاصطناعي في إطار الكتابة المشتركة، الشيء الذي يستدعي سياسات تعليمية وتشريعية تتسم بالمرونة والوعي.

المراجع المعتمدة

بورديو، بيير. (2012). مسائل في علم الاجتماع، ترجمة هناء صبحي، ط 1، أبوظبي: هيئة أبوظبي للسياحة والثقافة، مشروع كلمة.

شكدي، مصطفى. (2024). سيكولوجيا الإنسان المرقمن: التنشئة الخوارزمية وبناء الهوية في عصر الثورة الرقمية، ط 1، تازة. المغرب: مؤسسة باحثون للدراسات، الأبحاث، النشر والاستراتيجيات الثقافية.

Airoidi, M. (2022). Machine habitus: Toward a sociology of algorithms, Polity Press.

Almuhaidib, S., Alqahtani, R., Alotaibi, H. F., Saeed, A., Alnasrallah, S., Alshamsi, F., Alqahtani, S. A., & Alhazzani, W. (2024). Mapping the landscape of medical research in the Arab world countries: A comprehensive bibliometric analysis. Saudi Medical Journal, 45(4). (pp : 387-396) <https://doi.org/10.15537/smj.2024.45.4.20230968>

- Anderson, B., Galpin, R., & Juzek, T. S. (2025). Model misalignment and language change: Traces of AI-associated language in unscripted spoken English (Version 1). arXiv. <https://arxiv.org/abs/2508.00238v1>
- Ávila, F., Hannah-Moffat, K., & Maurutto, P. (2021). Machine learning algorithms and risk assessment. (pp : 87-103). In M. Schuilenburg & R. Peeters (Eds.), The algorithmic society: Technology, power, and knowledge. Routledge.
- Castells, M. (2010). The rise of the network society (2nd ed., with a new preface). Wiley-Blackwell.
- Charmaz, Kathy. (2006). Constructing grounded theory. Los Angeles-SAGE Publications.
- Cheng, Y., Sadasivan, V. S., Saberi, M., Saha, S., & Feizi, S. (2025, June 8). Adversarial paraphrasing: A universal attack for humanizing AI-generated text. arXiv preprint arXiv:2506.07001.
- Cheng, Y., Sadasivan, V. S., Saberi, M., Saha, S., & Feizi, S. (2025). Adversarial Paraphrasing: A Universal Attack for Humanizing AI-Generated Text (Version 1). arXiv. <https://arxiv.org/abs/2506.07001v1>
- Coglianese, C. (2021). Algorithmic regulation : Machine learning as a governance tool. (pp :35-52). In M. Schuilenburg & R. Peeters (Eds.), The algorithmic society: Technology, power, and knowledge. Routledge.
- Coley, M. (2023, August 16). Guidance on AI detection and why we're disabling Turnitin's AI detector. Brightspace Support. Vanderbilt University. <https://www.vanderbilt.edu/brightspace/2023/08/16/guidance-on-ai-detection-and-why-were-disabling-turnitins-ai-detector/>
- Common Sense Media. (2024). The dawn of the AI era: Teens' and parents' experiences with AI. Common Sense. Retrieved from <https://www.common Sense media.org/research>
- Curtin University. (2025, September 04). Update on Turnitin AI detection tool. Curtin University News. Retrieved from https://www.curtin.edu.au/news/oasis-news/update-on-turnitin-ai-detection-tool/?utm_source=chatgpt.com
- Digital Education Council. (2024). Global AI student survey 2024: AI or not AI—What students want. <https://digitaleducationcouncil.com>
- Doughman, J., Afzal, O. M., Toyin, H. O., Shehata, S., Nakov, P., & Talat, Z. (2024). Exploring the limitations of detecting machine-generated text. arXiv preprint arXiv:2406.11073. <https://arxiv.org/abs/2406.11073>
- Faucher, K. X. (2018). Social capital online. University of Westminster Press.
- Fisher, M. (2022). The chaos machine: The inside story of how social media rewired our minds and our world. Little, Brown and Company.
- Freeman, J. (2025, February). Student generative AI survey 2025 (HEPI Policy Note 61). Higher Education Policy Institute. <https://www.hepi.ac.uk/wp-content/uploads/2025/02/HEPI-Kortext-Student-Generative-AI-Survey-2025.pdf>
- García Mathewson, T. (2025, June 26). Costly and unreliable: AI and plagiarism detectors wreak havoc in higher ed. CalMatters & The Markup. <https://calmatters.org/education/higher-education/2025/06/ai-plagiarism-detectors-turnitin>
- Gegg-Harrison, W. (2024, August 22). Here I am, still shouting from the rooftops about “AI-detection” and why it’s so harmful to students. Medium. <https://writerethink.medium.com/still-shouting-from-the-rooftops-about-ai-detection-and-why-its-so-harmful-to-students-c1608bb57b20>
- Giray, L. (2024). The problem with false positives: AI detection unfairly accuses scholars of AI plagiarism. The Serials Librarian. Advance online publication. (pp :181-189). <https://doi.org/10.1080/0361526X.2024.2433256>

- Górska, A. M., & Jemielniak, D. (2025). AI racial bias: How text-to-image artificial intelligence generators construct prestigious professions. (pp :211-226). In D. Brzeziński, K. Filipek, K. Piowar, & M. Winiarska-Brodowska (Eds.), *Algorithms, artificial intelligence and beyond: Theorising society and culture of the 21st century*. Routledge. <https://doi.org/10.4324/9781032646930>
- Gretzky, M., & Dishon, G. (2024). Algorithmic authors in academia: Blurring the boundaries of human and machine knowledge production (short paper). In D. Olenik-Shemesh, I. Blau, N. Geri, A. Caspi, Y. Sidi, Y. Eshet-Alkalai, & Y. Kalman (Eds.), *Proceedings of the 19th Chais Conference for the Study of Innovation and Learning Technologies: Learning in the Digital Era*. The Open University of Israel.
- Gretzky, M., & Dishon, G. (2025). Algorithmic-authors in academia: Blurring the boundaries of human and machine knowledge production. *Learning, Media and Technology*, 50(3), (pp : 338-351). <https://doi.org/10.1080/17439884.2025.2452196>
- Henman, P. (2021). Governing by algorithms and algorithmic governmentality. (pp :19-34). In M. Schuilenburg & R. Peeters (Eds.), *The algorithmic society: Technology, power, and knowledge*. Routledge.
- Herbold, S., Hautli-Janisz, A., Heuer, U., Kikteva, Z., & Trautsch, A. (2023). A large-scale comparison of human-written versus ChatGPT-generated essays. *Scientific Reports*, 13(18617). (pp :1-11) <https://doi.org/10.1038/s41598-023-45644-9>
- Hong, S. (2020). Fetishizing Algorithms and Rearticulating Consumption. (pp :35-48). In : *Algorithmic culture: How big data and AI are transforming everyday life*, Edited by Stefka Hristova, Soonkwan Hong, and Jennifer Daryl Slack, Lexington Books.
- Jessop, B. (2018). On academic capitalism. *Critical Policy Studies*, 12(1), (pp :104–109). <https://doi.org/10.1080/19460171.2017.1403342>
- Kelly-Holmes, H. (2024). Artificial intelligence and the future of our sociolinguistic work. *Journal of Sociolinguistics*, 28(3). (pp :3-10). <https://doi.org/10.1111/josl.12678>
- Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4(9), 100811. <https://doi.org/10.1016/j.patter.2023.100779>
- Markl, N. (2024). Algorithmic language management: How do language technologies affect linguistic practices and beliefs? SSRN. <https://doi.org/10.2139/ssrn.5042410>
- Matzner, T. (2024). *Algorithms: Technology, Culture, Politics*. Routledge.
- Mintz, S. (2025, April 2). Apparently, I'm a bot: Writing in the age of AI suspicion. Inside Higher Ed. <https://www.insidehighered.com/opinion/columns/higher-ed-gamma/2025/04/02/apparently-im-bot>
- Noy, S., & Zhang, W. (2023). Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, 381(6654). (pp : 187-192). <https://doi.org/10.1126/science.adh2586>
- OpenAI. (2023, July 20). New AI classifier for indicating AI-written text [Blog update]. OpenAI. https://openai.com/index/new-ai-classifier-for-indicating-ai-written-text/?utm_source=chatgpt.com
- Pratama, A. R. (2025). The accuracy-bias trade-offs in AI text detection tools and their impact on fairness in scholarly publication. *PeerJ Computer Science*, 11, e2953. <https://doi.org/10.7717/peerj-cs.2953>
- Quintana, C. (2024, February 9). Professors proceed with caution in using AI detectors. Inside Higher Ed. <https://www.insidehighered.com/news/tech-innovation/artificial-intelligence/2024/02/09/professors-proceed-caution-using-ai>

- Rama, I., & Airoidi, M. (2024). The sociocultural roots of artificial conversations: The taste, class and habitus of generative AI chatbots. SocArXiv. <https://doi.org/10.31235/osf.io/dx9a3>
- Romele, A. (2024). Digital habitus: A critique of the imaginaries of artificial intelligence. Routledge.
- RTL – Research, Teaching, & Learning. (2023, May 1). Availability of Turnitin's artificial intelligence detection [Web page]. UC Berkeley. Retrieved September 7, 2025, from <https://rtl.berkeley.edu/news/availability-turnitins-artificial-intelligence-detection>
- Shirky, C., & Maddox, B. (2023, September 7). Disabling the AI Tool in Turnitin. NYU. Retrieved from <https://www.nyu.edu/content/dam/nyu/provost/documents/Disabling%20the%20AI%20Tool%20in%20Turnitin.pdf>
- Turnitin. (2023, April 5). Understanding false positives within our AI writing detection capabilities. Turnitin. <https://www.turnitin.com/blog/understanding-false-positives-within-our-ai-writing-detection-capabilities>
- Turnitin. (2024, November 20). AI or human? 2025 predictions say AI could boost the value of human writing. Turnitin. Retrieved September 7, 2025, from <https://www.turnitin.com/blog/ai-or-human-2025-predictions-say-ai-could-boost-the-value-of-human-writing>
- UNESCO. (2021). The Arab States. In UNESCO Science Report: The race against time for smarter development. UNESCO. <https://www.unesco.org/reports/science/2021/en/arab-states>
- University Information Services. (n.d.). Turnitin. Georgetown University. Retrieved September 7, 2025, from <https://uis.georgetown.edu/turnitin/>
- University of Cape Town. (2025, July 24). UCT scraps flawed AI detectors. UCT News. Retrieved from https://www.news.uct.ac.za/article/-2025-07-24-uct-scraps-flawed-ai-detectors?utm_source=chatgpt.com
- University of Waterloo. (2025, September). Discontinuing use of AI detection functionality in Turnitin. Retrieved from <https://uwaterloo.ca/associate-vice-president-academic/discontinuing-use-ai-detection-functionality-turnitin>
- Weber-Wulff, D., Anohina-Naumeca, A., Bjelobaba, S., Foltýnek, T., Guerrero-Dib, J., Popoola, O., Šigut, P., & Waddington, L. (2023). Testing of detection tools for AI-generated text. arXiv preprint arXiv:2306.15666.
- Yakura, H., Lopez-Lopez, E., Brinkmann, L., Serna, I., Gupta, P., Soraperra, I., & Rahwan, I. (2025). Empirical evidence of Large Language Model's influence on human spoken communication. Max Planck Institute for Human Development, Center for Humans and Machines. arXiv. <https://doi.org/10.48550/arXiv.2409.01754>

أدوار الباحثين:

- ✓ عبد الإله فرح: باحث أساسي، قاد البحث وصاغ الإشكالية والإطار السوسيولوجي وأشرف على المراجعة النهائية.
- ✓ محمد بزي: باحث مساعد، أسهم في بناء التأصيل النظري، وضبط المراجع، وصياغة الأقسام التحليلية.
- ✓ صابر بن داود: باحث مساعد، نفذ التجارب على أدوات كشف الانتحال، وحلل نتائجها تفصيلياً.
- ✓ عبد اللطيف طالي: باحث مساعد، دقق البيانات الرقمية، وراجع المنهجية التقنية؛ لضمان صرامة التحليل.

✓ حسن حبران: باحث مساعد، جمع الأمثلة الميدانية، ورصد تجارب أكاديميين، ومنصات دولية لدعم الدراسة.

الشكر والتقدير

يود المؤلفون أن يتوجهوا بالشكر إلى مراجعي هذه الدراسة؛ لتعليقاتهم البناءة، واقتراحاتهم المفيدة، كما يتوجهون بالشكر والتقدير إلى الجمعية المغربية لعلم الاجتماع على مساندتهم.

بيان حول الذكاء الاصطناعي التوليدي

يؤكد المؤلفون أنه تم الاعتماد في هذه الدراسة على برنامج "صحح لي" (<https://sahehly.com>) للتدقيق اللغوي والنحوي، وعلى نموذج الذكاء الاصطناعي التوليدي (ChatGPT) في المساعدة على البحث عن المصادر والمراجع العلمية ذات الصلة بموضوع الدراسة.

التمويل

يصرح المؤلفون بأنهم لم يتلقوا أي دعم مالي لإجراء هذه الدراسة أو لنشرها.

تضارب المصالح

يصرح المؤلفون بأن هذه الدراسة أنجزت في غياب أي علاقات تجارية أو مالية يمكن أن تفسر على أنها تضارب محتمل في المصالح.

ملاحظة الناشر

جميع الآراء الواردة في هذه الدراسة تعبر حصريا عن رأي المؤلفين، ولا تعكس بالضرورة آراء المؤسسات التي ينتمون إليها، أو آراء الناشر أو المحررين أو المراجعين. كما أن أي منتج قد يقيم في هذه الدراسة، أو أي ادعاء يورده مصنعه لا يضمنه الناشر، ولا يوصي به.