

ظاهرة هلوسة الذكاء الاصطناعي

عندما "يخترع" العقل الآلي الحقيقة

جمال مراد قيس

2026-05-05

في ظل الطفرة المتسارعة في تقنيات الذكاء الاصطناعي، وخاصة النماذج اللغوية الضخمة، باتت هذه الأنظمة قادرة على إنتاج نصوص تحاكي الكتابة البشرية بدرجة مذهشة من الطلاقة والدقة الظاهرية. غير أن هذه القدرة تخفي خلفها ظاهرة معقدة تُعرف بـ "هلوسة الذكاء الاصطناعي"، وهي من أبرز التحديات المعرفية والتقنية التي تواجه هذا المجال اليوم.

تشير هذه الظاهرة إلى ميل النماذج إلى توليد معلومات تبدو صحيحة ومتناسكة لغويًا، لكنها في الواقع غير دقيقة أو مختلقة بالكامل، وهو ما يطرح تساؤلات جوهرية حول موثوقية هذه الأنظمة وحدود استخدامها في السياقات الحساسة.

تكمّن جذور هذه الظاهرة في طبيعة عمل النماذج اللغوية نفسها، إذ إنها لا تمتلك فهمًا حقيقيًا للمعرفة أو الواقع، بل تعتمد على التنبؤ الإحصائي بالكلمات بناءً على الأنماط التي تعلمتها من كميات هائلة من البيانات. ونتيجة لذلك، فإنها قد تنتج إجابات تبدو منطقية، لكنها لا تستند إلى مصدر حقيقي أو تحقق معرفي. هذا النمط من "التخمين المتقدم" يجعل النماذج عرضة للوقوع في أخطاء لا يمكن تمييزها بسهولة، خاصة عندما تُعرض بثقة لغوية عالية توهي بالدقة.

وتزداد خطورة هذه الظاهرة في ظل الضغوط التصميمية التي تدفع النماذج إلى تقديم إجابات في جميع الأحوال، حتى في الحالات التي لا تتوفر فيها معلومات كافية. فبدلاً من الاعتراف بعدم المعرفة، تميل هذه الأنظمة إلى توليد إجابات محتملة، وهو ما يعزز من احتمالية إنتاج معلومات غير صحيحة. كما أن جودة البيانات المستخدمة في التدريب تلعب دورًا محوريًا، إذ إن أي تحيز أو نقص أو خطأ في البيانات الأصلية قد ينعكس بشكل مباشر على مخرجات النموذج، مما يفاقم من ظاهرة الهلوسة.

وقد أظهرت [دراسات حديثة](#) أن هذه الظاهرة ليست مجرد خلل عرضي، بل قد تكون مرتبطة ببنية النماذج نفسها، مما يجعل القضاء عليها بالكامل أمراً بالغ التعقيد. فحتى مع التحسينات المستمرة، لا تزال النماذج تعاني من صعوبة في التمييز بين الحقيقة والاحتمال، خاصة في المجالات التي تتطلب دقة معرفية عالية مثل الطب والقانون والبحث العلمي. وفي هذا السياق، تشير مراجعة علمية حديثة نُشرت عام 2025 إلى أن هلوسة الذكاء الاصطناعي تمثل أحد أبرز التحديات أمام اعتماد هذه التقنيات في التطبيقات الحرجة، نظراً لما قد تسببه من تضليل أو أخطاء ذات تبعات خطيرة.

ويمتد تأثير هلوسة الذكاء الاصطناعي بشكل خاص إلى المجال العلمي والبحثي، حيث يمكن أن يشكل تحدياً حقيقياً لسلامة المعرفة وإنتاجها. فاعتماد الباحثين أو الطلبة على هذه النماذج دون تدقيق قد يؤدي إلى إدراج مراجع غير موجودة، أو الاستناد إلى معلومات غير دقيقة تُعرض بثقة عالية، مما يهدد مصداقية الأبحاث العلمية. كما أن هذه الظاهرة قد تُربك عملية مراجعة الأدبيات العلمية (literature review)، إذ يصعب التمييز بين المصادر الحقيقية والمختلقة.

وفي بيئات البحث المتقدمة، قد يؤدي ذلك إلى تضليل مسارات البحث أو بناء فرضيات على أسس غير صحيحة، وهو ما ينعكس سلباً على جودة الإنتاج العلمي. ومع تزايد استخدام الذكاء الاصطناعي كأداة مساعدة في الكتابة والتحليل، تبرز الحاجة إلى تطوير معايير واضحة لاستخدامه في البحث العلمي، وتعزيز ثقافة التحقق والتدقيق، لضمان أن يبقى الذكاء الاصطناعي أداة داعمة للعلم، لا مصدرًا للالتباس أو الخطأ.

ولا تقتصر آثار هذه الظاهرة على الجانب النظري، بل تمتد إلى الواقع العملي، حيث تم توثيق حالات متعددة لتوليد مراجع علمية غير موجودة، أو تقديم معلومات تقنية خاطئة، أو حتى اقتراح حلول برمجية تعتمد على مكتبات وهمية. وتكمن الإشكالية في أن هذه الأخطاء لا تُقدّم على أنها احتمالات، بل تُعرض بصيغة تقريرية واثقة، مما يجعل اكتشافها أكثر صعوبة، خاصة لدى المستخدمين غير المتخصصين.

في مواجهة هذه التحديات، يعمل الباحثون على تطوير مجموعة من [الحلول التقنية](#) والمنهجية للحد من ظاهرة الهلوسة. من أبرز هذه الحلول ربط النماذج بمصادر بيانات خارجية موثوقة، بحيث يتم دعم الإجابات بمعلومات قابلة للتحقق، بالإضافة إلى تطوير آليات للتقييم الذاتي تسمح للنموذج بمراجعة مخرجاته قبل تقديمها. كما يجري العمل على تحسين أنظمة التدريب بحيث تُكافئ النماذج على الدقة والصدق، وليس فقط على الطلاقة اللغوية. ومن الاتجاهات الواعدة أيضاً تعزيز قدرة النماذج على التعبير عن عدم اليقين، بدلاً من تقديم إجابات حاسمة في جميع الحالات.

ورغم هذه الجهود، فإن هلوسة الذكاء الاصطناعي تظل تذكيرًا مهمًا بأن هذه الأنظمة، مهما بلغت من التطور، لا تزال أدوات تعتمد على أنماط إحصائية، وليست كيانات واعية تمتلك فهمًا حقيقيًا للعالم. ومن هنا، فإن الاستخدام المسؤول لهذه التقنيات يتطلب وعيًا بحدودها، وعدم الاعتماد عليها كمصدر وحيد للحقيقة، بل كأداة مساعدة ضمن منظومة أوسع من التحقق والتحليل البشري.

وفي هذا الإطار، يمكن القول إن هلوسة الذكاء الاصطناعي لا تمثل فقط تحديًا تقنيًا، بل تفتح أيضًا أفقًا فلسفيًا جديدًا حول طبيعة المعرفة في العصر الرقمي. فهي تضعنا أمام مفارقة عميقة: أنظمة قادرة على محاكاة الفهم دون أن تفهم، وعلى إنتاج المعرفة دون أن تتحقق منها. وربما يكون التحدي الأكبر في المرحلة القادمة ليس فقط في تقليل هذه الظاهرة، بل في إعادة تعريف علاقتنا بالمعرفة ذاتها في زمن أصبحت فيه الحقيقة قابلة للتوليد الآلي.

تواصل مع الكاتب: mohamedmouradgamal@gmail.com

References

1. [/https://arxiv.org/abs/2311.05232](https://arxiv.org/abs/2311.05232) — A Survey on Hallucination in Large Language Models

2. — Frontiers in Artificial Intelligence (2025) – Hallucinations in AI Systems
<https://www.frontiersin.org/journals/artificial-intelligence>