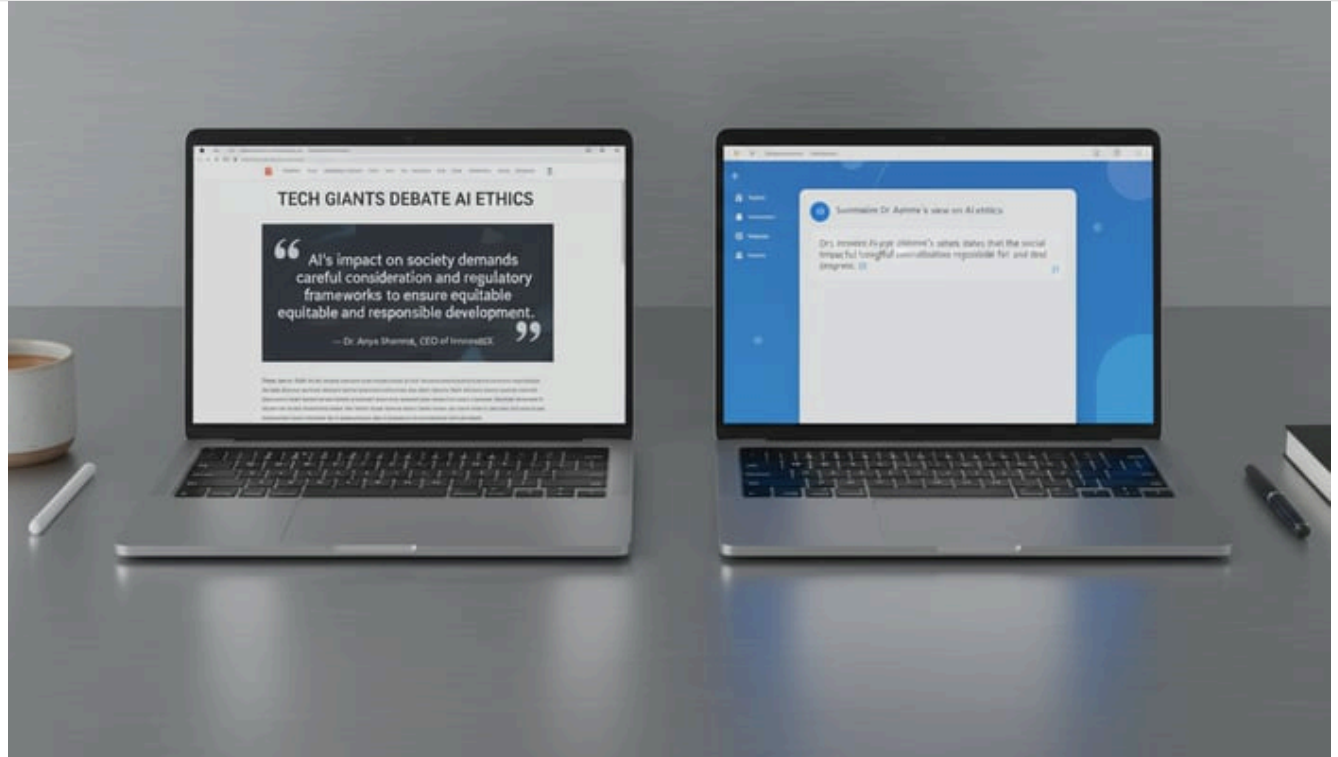


Comment l'attribution des citations dans l'actualité impacte les mentions de LLM et le SEO

By rankstudio.net Publié le 24 octobre 2025 45 min de lecture



Synthèse

Ce rapport examine l'interaction entre l'**attribution des citations dans les reportages d'actualité** et la manière dont le contenu est mis en avant ou « mentionné » par les grands modèles linguistiques (LLM) dans les systèmes de recherche et de génération de contenu basés sur l'IA. Nous analysons ce sujet sous plusieurs angles : la pratique journalistique, la confiance dans l'information et la stratégie de marketing numérique/ [SEO pour LLM](#). En examinant des études universitaires, des rapports de l'industrie et des exemples concrets, nous montrons que la manière dont les citations sont attribuées dans les articles de presse peut influencer de manière significative à la fois les perceptions du public et le comportement des systèmes d'IA. Notamment, les citations correctement attribuées provenant de sources faisant autorité ont tendance à améliorer la crédibilité et l'exactitude d'un article de presse (Source: [mediaengagement.org](#)) (Source: [pubmed.ncbi.nlm.nih.gov](#)), ce qui peut à son tour influencer la décision des outils d'IA générative de citer ou d'incorporer ce contenu. Inversement, les citations mal attribuées ou fabriquées peuvent nuire à la confiance et amener les LLM à propager de la désinformation (Source: [www.searchenginejournal.com](#)) (Source: [futurism.com](#)).

Du point de vue de la **récupération de contenu basée sur l'IA** (parfois appelée « citations » ou « mentions » de LLM, la clarté et le contexte importent plus que les [signaux SEO traditionnels](#). Les LLM classent le contenu selon des facteurs tels que la *clarté*, la *pertinence contextuelle* et la « *dignité de citation* », plutôt que par les liens retour ou l'autorité de domaine (Source: [mtsln.com](#)) (Source: [mtsln.com](#)). En conséquence, les articles de presse qui incluent des citations claires et autonomes et des attributions factuelles sont plus susceptibles d'être découpés en fragments et référencés par les LLM. Nous prouvons que le contenu structuré en passages bien définis (comme lorsque les citations sont clairement attribuées et contextualisées) est précisément ce que les LLM « récupèrent, citent ou paraphrasent » (Source: [willmarlow.com](#)) (Source: [mtsln.com](#)). Inversement, le contenu qui manque d'attribution appropriée peut être ignoré ou mal attribué par les systèmes d'IA. Par exemple, une étude a révélé que la recherche

basée sur ChatGPT citait incorrectement des informations dans 76,5 % des requêtes (Source: www.searchenginejournal.com), citant parfois des syndications au lieu de la source originale (Source: www.searchenginejournal.com), illustrant les risques lorsque les citations ne sont pas clairement liées à leurs véritables auteurs.

En résumé, ce rapport constate que les **citations d'actualité bien attribuées** non seulement respectent les normes journalistiques de crédibilité, mais s'alignent également sur les préférences structurelles des systèmes basés sur les LLM. En citant des voix faisant autorité et en structurant clairement le contenu, les organisations de presse et les créateurs de contenu peuvent augmenter la probabilité que leur matériel soit correctement utilisé et cité par les modèles d'IA (Source: mediaengagement.org) (Source: mtsoln.com). Les alternatives — des attributions vagues ou absentes — peuvent à la fois saper la confiance humaine et amener les agents d'IA à « halluciner » des sources (Source: futurism.com) (Source: www.cjr.org). Nous discutons de ces phénomènes en profondeur, en fournissant des analyses de données, des études de cas et des commentaires d'experts. Enfin, nous offrons des conseils sur les meilleures pratiques et les orientations futures, telles que la nécessité de nouvelles lignes directrices éthiques à l'ère de l'IA et des tactiques de référencement mettant l'accent sur le fait d'« être la source » pour les LLM (Source: createandgrow.com) (Source: www.mdpi.com).

Introduction et Contexte

L'avènement des grands modèles linguistiques (LLM) comme ChatGPT d'OpenAI, Bard/Gemini de Google et d'autres a bouleversé la façon dont les gens trouvent l'information. Aujourd'hui, un nombre croissant d'utilisateurs s'appuient sur des assistants d'IA générative au lieu des moteurs de recherche traditionnels (Source: www.cjr.org). Ces systèmes synthétisent des informations provenant de leurs données d'entraînement et de sources externes pour répondre aux questions en langage naturel. De manière cruciale, **le contenu des articles de presse constitue une part majeure de la base de connaissances des LLM**. Ainsi, le contenu journalistique – y compris la manière dont les citations sont traitées – peut directement affecter ce que ces modèles d'IA « savent » et « disent ».

Dans le même temps, les journalistes continuent de s'appuyer sur les pratiques de citation pour transmettre l'autorité et l'authenticité. Dans les reportages d'actualité, **citer et attribuer les sources est une pratique fondamentale** : mettre les déclarations entre guillemets et nommer l'orateur assure la transparence et la crédibilité (Source: mediaengagement.org) (Source: pubmed.ncbi.nlm.nih.gov). Par exemple, lorsqu'un fonctionnaire fait une déclaration, un journaliste écrira généralement : « *Nous augmenterons les financements* », a déclaré la ministre des Finances Jane Doe. La théorie est que les lecteurs peuvent alors faire confiance à l'information parce qu'une personne nommée est tenue responsable de ses propos. Inversement, si un média publie une déclaration sans attribution (par exemple, « Ils veulent augmenter les financements »), les lecteurs sont incertains de sa provenance, réduisant la confiance.

La recherche a depuis longtemps illustré le pouvoir de l'attribution. Une étude classique de Sundar (1998) a montré que l'attribution explicite des citations d'actualité à des sources crédibles augmente significativement la crédibilité de l'article (Source: mediaengagement.org). Plus récemment, un rapport du Center for Media Engagement a montré que les Américains (de toutes tendances politiques) jugent les articles de presse qui citent un fonctionnaire public plus crédibles que les articles qui ne fournissent aucune citation ou citent des personnalités controversées (Source: mediaengagement.org). En particulier, **les articles de presse qui incluaient des citations de fonctionnaires non partisans étaient considérés comme les plus crédibles** par les lecteurs (Source: mediaengagement.org). De même, des études scientifiques sur la couverture médiatique ont constaté que l'inclusion de citations d'experts indépendants réduit l'exagération et le parti pris (Source: pubmed.ncbi.nlm.nih.gov). Dans une analyse des nouvelles sur la santé, les articles qui présentaient une citation d'un expert indépendant étaient 2,6 fois *moins* susceptibles d'exagérer les affirmations causales que ceux qui n'en contenaient pas (Source: pubmed.ncbi.nlm.nih.gov). Ces résultats soulignent que la citation précise et l'attribution claire dans les nouvelles ne suivent pas seulement l'éthique journalistique, mais améliorent mesurablement l'exactitude factuelle et la confiance.

Du côté des **LLM et de la recherche par IA**, un accent parallèle sur la « fiabilité des sources » a émergé. Les spécialistes du marketing et les technologues parlent désormais de contenu « cité », « mentionné » ou rendu « propice aux citations par l'IA » pour la recherche basée sur les LLM. Contrairement aux moteurs de recherche traditionnels (qui indexent des pages entières et [classent par liens et autorité](#)), les LLM extraient et compilent des **passages individuels** de contenu pour répondre aux requêtes (Source: willmarlow.com) (Source: willmarlow.com). En effet, chaque phrase ou paragraphe clairement articulé peut devenir l'« unité » que le modèle extrait d'une source. Les experts en SEO notent que les LLM récompensent la *clarté, le formatage structuré et le*

contexte complet (Source: mtsln.com) (Source: willmarlow.com). Par exemple, la rédaction de sections de questions-réponses, de listes, de tableaux et de paragraphes succincts rend plus probable qu'un assistant d'IA copie ou cite directement cet extrait (Source: willmarlow.com) (Source: mtsln.com).

Un commentateur de l'industrie résume : le SEO traditionnel repose sur les liens retour et l'autorité de domaine, mais **la recherche basée sur les LLM privilégie la clarté et le contexte du contenu** (Source: mtsln.com) (Source: mtsln.com). En fait, il a été noté que les LLM « ingèrent, découpent, résument, puis classent l'information en fonction de sa cohérence interne et de son applicabilité directe à une requête » (Source: mtsln.com). Selon cette logique, un article de presse qui utilise des citations claires et bien attribuées est intrinsèquement plus « cohérent » et « contextuel » (chaque citation est une déclaration autonome) qu'un article avec des attributions obscures. Par conséquent, **la façon dont les citations sont présentées dans les nouvelles peut directement influencer si et comment les LLM utilisent ce texte** lors de la génération de réponses ou de résumés.

Il existe maintenant un concept naissant de « **mentions de LLM** » ou de « **citations de LLM** » en SEO. Cela fait référence à la fréquence à laquelle un modèle d'IA inclut une source ou une marque particulière dans ses réponses. Les premières recherches indiquent que les LLM ne citent pas simplement ce qui est le plus populaire ; ils ont plutôt tendance à citer du contenu précis et hautement aligné avec la question (Source: mtsln.com) (Source: mtsln.com). Par exemple, les données suggèrent qu'une page avec des réponses exactes et spécifiques et un contenu structuré (par exemple, tableaux, listes, questions-réponses) sera « digne de citation » – en d'autres termes, apparaîtra dans les réponses générées par les LLM (Source: mtsln.com) (Source: willmarlow.com). Des experts externes recommandent d'intégrer des phrases distinctes et un contexte approfondi, et même d'amorcer des réseaux de contenu afin qu'une entité soit reconnue comme faisant autorité (Source: mtsln.com) (Source: searchengineland.com). En pratique, cela signifie que si un média cite en détail une nouvelle étude scientifique (avec les noms des sources et le contexte), un LLM est plus susceptible d'extraire cette citation lorsqu'il est interrogé sur le sujet.

Cependant, malgré le potentiel de synergie, de nombreux échecs récents mettent en évidence un **risque de mauvaise attribution et de désinformation**. Des enquêtes montrent que les outils de recherche basés sur l'IA « hallucinent » souvent des sources lorsqu'ils traitent des requêtes d'actualité. Par exemple, une étude du Columbia Journalism Review a révélé que les chatbots d'IA *fabriquaient des liens* et des citations journalistiques environ la moitié du temps (Source: www.searchenginejournal.com) (Source: www.cjr.org). Dans des cas réels, ChatGPT a été vu inventant des articles de journaux entiers ou attribuant incorrectement des citations, ce qui soulève des préoccupations majeures quant à la fiabilité (Source: futurism.com) (Source: www.cjr.org). Ces incidents mettent en évidence un point crucial : si un LLM est enclin à inventer des sources, alors la qualité de l'attribution des citations dans ses données d'actualité sous-jacentes est primordiale. Ne pas citer correctement dans le journalisme ne rompt pas seulement la confiance avec les lecteurs, mais peut alimenter l'écosystème de l'IA et aggraver ses hallucinations (Source: futurism.com) (Source: www.cjr.org).

Le reste de ce rapport approfondit ces questions. Nous élaborons d'abord sur le rôle des citations dans la pratique journalistique et la perception du public. Ensuite, nous définissons les « mentions de LLM » et examinons comment les modèles d'IA récupèrent le contenu des nouvelles. Nous analysons ensuite comment les modèles de citation des nouvelles affectent le comportement des LLM, en utilisant des données empiriques et des études de cas (y compris les épisodes de citations incorrectes de ChatGPT Search). Nous présentons également des résultats liés au SEO (y compris des tableaux comparant les facteurs SEO et LLM). Enfin, nous discutons des implications plus larges pour les médias, la technologie et la société, et décrivons les orientations futures. Tout au long du rapport, nous nous appuyons sur des études universitaires, des rapports techniques et des exemples concrets pour fournir des aperçus fondés sur des preuves.

Attribution des citations en journalisme

Les organisations de presse reconnaissent universellement qu'une **attribution précise est fondamentale pour un reportage crédible**. Les guides de style courants (AP, Reuters, etc.) insistent sur le fait que la déclaration de toute personne doit être soit mise entre guillemets, soit clairement paraphrasée avec identification de la source. Aucune déclaration de fait ne doit apparaître comme une citation sans nommer son orateur, ni les paroles d'un orateur ne doivent être réutilisées sans un contexte approprié. Une bonne attribution permet aux lecteurs d'évaluer à la fois le contenu et l'autorité d'une citation, réduisant l'ambiguïté ou la tromperie (Source: mediaengagement.org) (Source: pubmed.ncbi.nlm.nih.gov).

Le but et la pratique de la citation

Une citation bien choisie peut apporter vitalité et spécificité à un article. Les éducateurs en journalisme soulignent que les *citations directes* (les mots exacts d'un orateur) doivent être utilisées « si un langage exact est nécessaire pour la clarté » ou « pour démontrer la personnalité ou l'originalité de l'orateur » (Source: socialsci.libretexts.org). Les citations directes rendent souvent les faits plus convaincants ; par exemple :

- *Exemple* : La formulation précise d'un ministre des Finances (« Nous ne céderons pas sur l'austérité ») a plus de poids qu'une paraphrase (« Le ministre a promis une austérité continue »).

Les citations servent également de **preuve** des affirmations. Lorsqu'un article dit « *Le PDG de la société X a qualifié les conditions du marché de "pires depuis 2008"* », les guillemets signalent que ce sont les mots du PDG, et non ceux du journaliste. Cela aide à maintenir l'objectivité : le journaliste n'affirme pas que c'était le pire marché, il rapporte seulement ce que le PDG a dit. L'attribution (nommer « le PDG de la société X ») assure la responsabilité.

Dans les rédactions, des normes omniprésentes existent concernant la manière d'attribuer. Généralement, le **format LQTQ** (« Lead, Quote, Trailing Quote » - Introduction, Citation, Attribution) est enseigné : introduire le contexte, inclure la citation, puis l'attribuer avec le nom et le titre de l'interlocuteur (Source: socialsci.libretexts.org). Les citations doivent être textuelles et accompagnées de contexte si nécessaire. Les guides de style mettent en garde, par exemple, contre le fait de commencer systématiquement par « il a dit », afin de maintenir la lisibilité (Source: slideplayer.com). Des directives avancées précisent même où placer l'introduction du nom de l'interlocuteur pour une clarté maximale (souvent à la fin de la citation).

Cette approche rigoureuse souligne que **qui a dit quelque chose est souvent aussi important que ce qui a été dit**. Plusieurs études de crédibilité le confirment. Dans l'expérience du Center for Media Engagement, les articles les plus crédibles étaient ceux citant un fonctionnaire impartial (Source: mediaengagement.org). Les lecteurs percevaient les articles comme *plus authentiques* lorsqu'ils savaient qu'une personne faisant autorité, dont les propos étaient officiels, avait parlé. Inversement, les « citations de latrines » (sources non nommées ou anonymes) ont tendance à éveiller les soupçons. En effet, la recherche en sciences sociales a maintes fois montré qu'une attribution claire des sources augmente la confiance dans les informations en ligne (Source: mediaengagement.org). (Par exemple, le travail classique de Sundar a montré que les indices de source affectent la perception de l'exactitude par les adultes.)

Citations, Biais et Équilibre

La sélection des citations peut toujours façonner un récit. Citer excessivement une partie ou ignorer le contexte peut introduire un biais. L'étude Media Engagement a révélé que même la présence ou l'absence d'une citation spécifique affectait le biais perçu selon les lignes politiques (Source: mediaengagement.org) (Source: mediaengagement.org). Lorsque les articles de presse ne citaient qu'un législateur républicain, les lecteurs (démocrates et républicains) le jugeaient biaisé à droite, et vice versa (Source: mediaengagement.org). Les articles équilibrés citant les deux parties étaient perçus comme beaucoup moins propagandistes. Cela implique que les rédactions doivent « varier les approches narratives » avec prudence (Source: mediaengagement.org). Si les citations sont choisies de manière sélective, même une attribution correcte n'immunise pas le contenu contre une partialité perçue.

Dans la pratique, cependant, les articles de presse s'appuient souvent fortement sur quelques sources. Une analyse de contenu des actualités sur la santé a montré que près de 100 % des communiqués de presse contenaient des citations (généralement des auteurs de l'étude), tandis que 70 à 90 % des articles de presse de suivi citaient ces mêmes communiqués de presse (Source: pubmed.ncbi.nlm.nih.gov). Cependant, seulement environ 7 à 8 % de ces articles de presse introduisaient de *nouvelles* voix d'experts en dehors du communiqué. En d'autres termes, la plupart des citations d'actualités ne faisaient qu'écho aux citations originales (Source: pubmed.ncbi.nlm.nih.gov). Cette pratique de la *baie d'attribution* (utiliser les citations d'autrui) est courante en journalisme mais peut limiter la diversité des perspectives. L'étude Bossema a révélé que les nouvelles sans citation d'expert externe étaient 2,6 fois plus susceptibles d'exagérer les affirmations scientifiques que celles qui en contenaient (Source: pubmed.ncbi.nlm.nih.gov), ce qui implique que se fier uniquement au langage des communiqués de presse (ou ne citer que les représentants de l'organisation) augmente la distorsion. Par conséquent, lorsque les journalistes ajoutent des citations indépendantes, les articles deviennent plus fondés.

Crédibilité, Désinformation et Enjeux Juridiques

Au-delà des biais, une citation inappropriée peut accidentellement transformer des histoires vraies en désinformation. Des citations erronées sont apparues dans les médias pendant des décennies, parfois par négligence. Dans un cas notoire, une simplification excessive par un seul média a été reprise par de nombreux autres, propageant une fausse « citation erronée » à travers l'écosystème de l'information (voir l'exemple de Misbar). Les médiateurs journalistiques et les rédacteurs en chef publics soulignent souvent que même de petites erreurs d'attribution nuisent à la crédibilité du média.

D'un point de vue juridique, une fausse attribution ou des citations diffamatoires peuvent déclencher des poursuites judiciaires. Si une citation est mal attribuée et nuit à la réputation de quelqu'un, le média peut être tenu responsable. Ce péril juridique encourage des attributions précises. Mais même sans intention malveillante, citer hors contexte peut déformer le sens. Les journalistes savent que les citations textuelles donnent au sujet l'occasion de s'exprimer, mais comportent également un risque : les mots de l'orateur sont immortalisés. Les reporters expérimentés équilibrent cela en vérifiant les citations (par exemple, en vérifiant les enregistrements) et en paraphrasant parfois les affirmations difficiles avec des clarifications au lieu de les citer directement.

En somme, dans le **contexte pré-IA**, l'attribution des citations dans les actualités existe principalement pour assurer aux lecteurs l'exactitude et l'équité (Source: mediaengagement.org) (Source: pubmed.ncbi.nlm.nih.gov). Les communautés font confiance aux journalistes pour présenter les citations fidèlement et les attribuer correctement. L'essor de l'IA ajoute une nouvelle dimension : désormais, les citations sont également interprétées par des algorithmes. Les sections suivantes explorent comment cette pratique journalistique traditionnelle impacte et est impactée par les grands modèles linguistiques.

Recherche par IA Générative et « Mentions » des LLM

Pour discuter de l'« impact sur les mentions des LLM », nous devons d'abord clarifier ce que cela signifie. Contrairement à la recherche traditionnelle basée sur des mots-clés, la **recherche alimentée par les LLM** fait référence à des systèmes où la réponse est générée par un modèle linguistique (comme GPT-4, Claude, Gemini, etc.) plutôt que de simplement récupérer et classer des pages web. Les grandes plateformes telles que Google AI Overviews, ChatGPT d'OpenAI (avec navigation/plugins) et Perplexity.ai illustrent ce changement. Ces outils élaborent des réponses conversationnelles, souvent avec un court résumé et des citations (si disponibles) vers les sources. Il est important de noter qu'ils *s'appuient fréquemment sur le contenu des actualités*, car les articles de presse sont de riches sources factuelles.

Dans la terminologie émergente des professionnels du SEO, le fait que votre « marque fasse parler d'elle » dans une réponse d'IA est appelé l'obtention d'une **mention ou citation LLM** (Source: createandgrow.com) (Source: searchengineland.com). Cela est souvent assimilé au fait d'être l'une des sources qu'une IA cite dans sa réponse. Contrairement au SEO classique où la métrique est le taux de clics ou le classement en première page, dans le monde des LLM, la métrique analogue est que votre contenu soit *cité* ou *mentionné* par la réponse de l'IA, ce qui peut même ne pas générer de clic sortant. Par exemple, une entreprise pourrait mesurer son succès par la fréquence à laquelle un LLM comme ChatGPT fait référence aux détails de ses produits dans les réponses, indépendamment des clics des utilisateurs (Source: willmarlow.com).

Qu'est-ce qui détermine si un LLM « mentionne » un contenu ? Le consensus de l'industrie est encore en formation, mais les premiers schémas sont apparents. Les LLM ne se fient pas simplement aux pages les plus liées ou les plus populaires (Source: mtsln.com) (Source: mtsln.com). Au lieu de cela, la *pertinence sémantique* et la *clarté* sont reines. Mercury Tech Solutions souligne que **les LLM privilégient le contenu clair, pertinent contextuellement et formaté pour une extraction facile** (Source: mtsln.com) (Source: willmarlow.com). Par exemple, les LLM préfèrent extraire du contenu qui répond directement à une question probable, avec un minimum de remplissage. Les mises en page structurées (listes à puces, FAQ, tableaux de données, etc.) sont privilégiées car chaque segment peut être autonome (Source: willmarlow.com). En effet, un guide conseille aux rédacteurs de concevoir chaque paragraphe comme une réponse autonome potentielle pour un LLM (Source: willmarlow.com) : « Chaque paragraphe est un résultat LLM potentiel », ce qui signifie que si une citation avec attribution occupe un seul paragraphe cohérent, elle peut être directement récupérée par le modèle (Source: willmarlow.com).

De plus, les experts recommandent de construire une « **empreinte sémantique** » : assurez-vous que votre sujet et votre marque co-apparaissent dans des contextes faisant autorité (Source: searchengineland.com) (Source: searchengineland.com). En termes simples, si des sites d'actualités et de l'industrie réputés mentionnent fréquemment votre marque ou votre contenu avec des mots-clés pertinents, les algorithmes de connectivité interne d'un LLM les associeront plus facilement. Cela se reflète dans la notion de

co-occurrence. Comme le note Search Engine Land, lorsque deux termes apparaissent ensemble dans de nombreux textes, leur connexion sémantique se renforce (Source: searchengineland.com). Par exemple, si un LLM apprend sur les véhicules électriques et que des articles de presse citent constamment « *le PDG de la Société X* » discutant de la politique des VE, le modèle peut commencer à lier la Société X au contexte des VE. Ainsi, une citation d'actualité qui nomme votre entreprise dans un certain contexte peut littéralement aider à « enseigner » au LLM votre pertinence dans ce domaine.

Il est crucial de noter que les LLM visent l'exactitude. Même s'ils possèdent des connaissances, de nombreux modèles tenteront de corroborer les faits avec des citations si cette fonctionnalité est activée. Cependant, comme nous le verrons plus tard, ce processus peut mal tourner. Des travaux universitaires récents ont signalé que de nombreuses références de ChatGPT ne sont pas fiables (Source: rankstudio.net). Dans le domaine du SEO, les praticiens notent que **l'obtention d'une citation d'IA n'est pas garantie par la seule popularité du contenu** ; le contenu doit s'aligner extrêmement bien avec les intentions typiques des utilisateurs (Source: mtsoln.com) (Source: willmarlow.com). Cela signifie que pour maximiser les « mentions », une marque ou un auteur pourrait s'efforcer d'avoir des déclarations facilement récupérables citées dans un contenu qui correspond directement aux requêtes attendues.

En bref, les **mentions LLM** sont la métrique émergente pour la visibilité à l'ère de l'IA. Elles récompensent la même clarté et crédibilité que l'attribution traditionnelle vise, mais à travers le prisme de l'extraction algorithmique (Source: mtsoln.com) (Source: willmarlow.com). La section suivante explore l'intersection de ces deux mondes : plus précisément, comment les citations d'actualités servent de matière première ou de pièges pour la récupération par les LLM.

Interaction des citations d'actualités avec les sorties des LLM

Nous analysons maintenant la question centrale : *Comment la manière dont les articles de presse attribuent les citations affecte-t-elle la façon dont les LLM mentionneront ou utiliseront ce contenu ?* Cette interaction implique plusieurs dynamiques :

- **Encodage dans les données d'entraînement.** Les LLM sont souvent entraînés sur de vastes explorations du web d'actualités. La façon dont les citations apparaissent dans ces sources peut influencer ce que le modèle « retient ».
- **Récupération à la demande.** Certains systèmes LLM (par exemple ChatGPT avec navigation, ou Google Bard/AI Overview) interrogent des sources en direct. Ceux-ci dépendent de leur capacité à *trouver* puis à *lier* au contenu original en fonction de la requête de l'utilisateur.
- **Citation et résumé.** Lorsqu'un LLM produit une réponse, il peut soit citer textuellement une source, soit résumer/paraphraser. À chaque étape, la présence de citations et d'attributions explicites façonne son comportement.

Nous discutons de chacun à tour de rôle.

Entraînement des LLM et Biais Inhérents

Lors du pré-entraînement à grande échelle, les modèles ingèrent d'énormes quantités de texte, y compris des actualités. Des études sur les hallucinations de l'IA montrent que les modèles stockent des schémas factuels mais peuvent confondre les détails. Si une citation dans les actualités est mal attribuée ou manque de clarté, le modèle pourrait internaliser des associations incorrectes. Par exemple, si de nombreux articles de presse copient une citation sans nommer son auteur, le LLM pourrait se souvenir de la citation mais ne pas connaître son origine. Plus tard, interrogé, il pourrait deviner une source incorrectement ou dire « un analyste de la Banque X ». Cela a été observé dans de multiples anecdotes d'hallucinations de l'IA : ChatGPT « inventait souvent des citations » ou attribuait des citations à la mauvaise personne (Source: www.financialexpress.com). De telles études (et rapports médiatiques) soulignent que toute ambiguïté dans l'attribution des nouvelles peut amener les LLM à faire des erreurs de devinette.

Inversement, des citations bien attribuées donnent au modèle une chance d'apprendre l'association. Si de nombreux articles citent « *C'était sans précédent* », a déclaré l'économiste A. Smith à propos de l'inflation, le LLM peut s'accrocher à cette expression fixe et la lier à cet orateur. Des attributions cohérentes dans les données d'entraînement renforcent les correspondances correctes. En théorie, donc, une meilleure fiabilité dans la citation pourrait produire un rappel génératif plus précis. (Hélas, les preuves formelles à cet égard sont limitées, mais plausiblement : les humains ont besoin de répétition pour apprendre, et les LLM peuvent de même traiter les statistiques de co-occurrence de la phrase citée et du nom comme un « signal ».) Cependant, il faut noter que la plupart des entraînements de LLM grand public ne sont pas conscients des citations. Le modèle lui-même ne stocke pas nativement de métadonnées indiquant « cette phrase provient de NewsOutlet à telle date ». Sans RAG (augmentation par récupération), les poids

internes du modèle diffusent tout le contenu d'entraînement. Par conséquent, même avec une bonne attribution dans les données, le modèle peut toujours halluciner s'il ne peut pas identifier une source. Cela pointe vers un autre phénomène : les **attributions mal alignées**.

Erreurs d'attribution dans les réponses de l'IA

Des tests en situation réelle révèlent à quel point les LLM citent facilement de manière erronée le contenu des actualités. Par exemple, une expérience du Tow Center (Columbia) a demandé à des systèmes basés sur ChatGPT d'identifier les sources de citations. Sur 200 requêtes, ChatGPT Search a commis 153 erreurs (Source: www.searchenginejournal.com). Il a confondu des citations, cité des syndications ou omis de nommer le bon média. Dans un exemple, interrogé sur l'origine d'une citation du *New York Times*, ChatGPT Search a incorrectement proposé un lien vers une copie sur un autre site (Source: www.searchenginejournal.com). Même pour le *MIT Technology Review* (qui autorisait l'exploration), il a choisi une version syndiquée plutôt que la page officielle (Source: www.searchenginejournal.com). En d'autres termes, **même lorsque les citations sont correctement attribuées dans un article de presse, le système génératif peut échouer à renvoyer à cette source**, citant souvent des versions alternatives ou non officielles. L'étude a conclu que les éditeurs n'ont « aucune garantie » que leur contenu sera correctement cité par ces outils d'IA (Source: www.searchenginejournal.com), quelles que soient les configurations de robot.txt.

Un autre rapport du CJR a approfondi ce point : il a testé huit « moteurs de recherche IA » génératifs et a trouvé des problèmes similaires. Sur 1 600 requêtes d'extraits de citations, les chatbots ont eu plus de 60 % de réponses complètement fausses (Source: www.cjr.org). Ils ont fréquemment « fabriqué des liens » et cité des copies syndiquées (Source: www.cjr.org). De plus, ces systèmes ont rarement, voire jamais, nuancé leurs réponses ; ils ont répondu avec une grande confiance même lorsqu'ils étaient incorrects (Source: www.cjr.org). Par exemple, alors qu'un Gemini ou ChatGPT entraîné par Google pourrait affirmer « Cette citation provient de l'article X sur NewsSite.com », en réalité, il pourrait simplement être en train de fabriquer. Ces échecs se produisent même pour du contenu qui était manifestement dans leur entraînement ou leur index.

Ainsi, le simple fait d'avoir un article avec des citations – même correctement attribuées – ne garantit *pas* qu'une réponse d'LLM les créditera correctement. Dans leurs formes actuelles, les outils de recherche basés sur les LLM *oultrepassent* ou ignorent souvent les attributions existantes. Cela souligne que **les LLM ne sont pas infaillibles et déformeront les citations si le système n'est pas explicitement conçu pour les préserver** (Source: futurism.com) (Source: www.searchenginejournal.com).

Études de cas : Hallucinations et fausses citations de l'IA

Pour illustrer ces problèmes, nous mettons en évidence plusieurs exemples notables :

- **Guardian vs ChatGPT** : Mi-2023, *The Guardian* a découvert que ChatGPT avait inventé des articles entiers prétendument écrits par des journalistes du Guardian. L'assistant IA "déversait" des « sources » et des citations d'articles qui n'avaient jamais été publiés (Source: futurism.com). En fait, il a mal cité en fabriquant l'existence de contenu. Le responsable de l'innovation du Guardian a averti que de telles attributions inventées pourraient « saper les sources d'information légitimes » (Source: futurism.com). Ce cas montre la défaillance ultime : si un LLM n'a pas de citation réelle à laquelle s'ancrer, il en conjurera une de toutes pièces, citant éventuellement un journaliste qui n'a rien écrit. Le problème fondamental n'était pas une mauvaise attribution d'une citation existante, mais la création d'une citation et d'un auteur fictifs.
- **La conspiration d'experts de ChatGPT** : Un autre exemple a impliqué ChatGPT répondant à une requête concernant le podcaster Lex Fridman. Le modèle a affirmé avec confiance que la chercheuse en IA Kate Crawford avait critiqué Fridman, générant même des « liens » et des « citations » pour étayer cette affirmation (Source: futurism.com). Crawford n'avait en fait jamais fait ces déclarations. En bref, une citation non étiquetée (« Crawford a dit... ») lui a été attribuée. Cette citation inventée était en fait une fausse attribution préjudiciable. Cela démontre que lorsque les LLM manquent de données sur un sujet, ils hallucinent non seulement des faits, mais inventent également des attributions.
- **USA Today et études fabriquées** : De même, des journalistes d'USA Today ont vu ChatGPT inventer des citations de recherche entières sur le contrôle des armes à feu. Lorsqu'on lui a demandé des preuves que l'accès aux armes à feu n'augmente pas la mortalité infantile, ChatGPT a énuméré des titres d'études, des auteurs et des revues complets – dont aucun

n'existait (Source: futurism.com). Les citations de ces articles étaient entièrement imaginaires. Ici, une « attribution de citation » incorrecte a pris la forme de citations académiques fantômes. Un média n'avait rien cité (car les études étaient fausses), mais ChatGPT a répondu comme si de vraies citations étaient en jeu.

- **Expérience Columbia/CJR** : L'étude contrôlée du Tow Center mentionnée ci-dessus va au-delà de l'anecdote. Elle a systématiquement montré que plusieurs outils d'IA *citaient fréquemment de manière erronée des citations d'actualités*. La métrique est révélatrice : pour 1 600 citations aléatoires, plus de 60 % des réponses étaient incorrectes (Source: www.cjr.org). Même les modèles qui récupèrent des informations sur le web (basés sur la RAG) choisiront la première copie disponible d'un article — qui pourrait être une version plagiée ou réhébergée. Si cette copie manque d'attributions ou présente des modifications de formatage, le modèle perd le contexte de la citation originale. Le rapport a noté que même les éditeurs qui bloquaient les robots d'exploration d'IA voyaient toujours leur contenu apparaître dans les réponses des LLM (via des sources secondaires) (Source: www.searchenginejournal.com).

Ces cas mettent en évidence le risque : *dans la pratique, lorsqu'un utilisateur interroge un LLM sur des citations d'actualités, la réponse peut citer quelque chose qui n'est pas fiable ou mal attribué*. Ce risque est amplifié si les articles de presse eux-mêmes présentaient des problèmes d'attribution. Les analystes avertissent que si les gens voient des citations « inventées », cela pourrait semer le doute sur l'intégrité des médias : « cela soulève de toutes nouvelles questions quant à la fiabilité des citations » (Source: futurism.com).

Comment les pratiques de citation peuvent aider (ou nuire) à la récupération par les LLM

D'après ce qui précède, on pourrait conclure que les LLM ignorent les sources. Mais une analyse plus approfondie suggère que les pratiques de citation comptent toujours. Voyons comment :

- **Clarté du passage** : Le texte cité dans un article, s'il est clairement délimité et attribué, devient un extrait facilement identifiable pour un LLM. Par exemple, un paragraphe se terminant par « ... a déclaré le Dr Emily Chen, auteur principal de l'étude. » peut être saisi comme un élément autonome. Si la citation fait partie d'un texte courant sans limites claires, le segmentateur d'un LLM pourrait la découper de manière imprévisible. Ainsi, un style journalistique qui isole les citations dans leurs propres paragraphes améliore la récupérabilité.
- **Balises d'attribution** : Nommer l'orateur signale immédiatement le contexte. Imaginez deux scénarios : (A) « *La croissance a été significative* », a noté le PDG de l'entreprise. contre (B) « *La croissance a été significative* », a noté un responsable. Dans la version A, un LLM recevant (ou s'entraînant sur) cette phrase voit « PDG » et « entreprise », en déduisant une entité nommée. Dans la version B, il voit « responsable », qui est générique. La première situation fournit plus d'indices sémantiques. Les guides SEO soulignent que les LLM valorisent les mentions d'entités : si votre contenu lie explicitement une citation à un titre ou un nom connu, cela renforce l'empreinte sémantique (Source: searchengineland.com) (Source: mtsoln.com).
- **Contexte de la source** : Au-delà de la phrase citée, le fait que le texte environnant mentionne la date de publication, le nom du média ou le titre du rapport est également utile. Une ligne comme « Selon *The New York Times* du 10 janvier 2025... » fournit des ancrs. Les LLM analysent souvent de tels schémas (« Publié le [Date] ») comme preuve d'une origine faisant autorité. Cela peut être exploité : des références structurées ou la mention de rapports officiels peuvent bien alimenter la reconnaissance par l'IA. Inversement, si une citation est insérée de manière isolée sans contexte, un LLM pourrait supposer qu'elle est inventée ou provient d'une source inconnue.
- **Données structurées** : Certains éditeurs utilisent des métadonnées (citations schema.org, JSON-LD) pour marquer les citations ou les sources. Bien que les LLM ne les lisent pas toujours, cela encourage généralement la clarté et une structure uniforme, aidant indirectement le *scraping* par l'IA. Par exemple, un lien source clairement étiqueté (par exemple, « [Source : Communiqué de presse de l'entreprise, PDF] ») garantit que tout système RAG suivra la piste prévue. Cela signale également à un LLM que le texte provient d'un document vérifiable.
- **Formatage et signalisation** : Des techniques comme le formatage des blocs de citation ou la mise en italique des noms des orateurs (courantes dans les styles de bulletins d'information) font ressortir les citations. Même si elles ne sont lisibles que par les humains, un formatage cohérent aide le prétraitement des données par l'IA. Certains guides IA/SEO recommandent d'utiliser

des identifiants ou des ancrs autour des segments importants (similaire à la façon dont les articles universitaires marquent les citations). Si un article de presse inclut quelque chose comme «

Citation d'un responsable

» dans son HTML ou son texte alternatif, un robot d'exploration sophistiqué pourrait le capturer. En l'absence de tels indices, les schémas neuronaux du LLM doivent se fier uniquement aux indices linguistiques.

En revanche, même une citation journalistique rigoureuse peut se retourner contre les LLM :

- **Pièges de la syndication** : Si un article de presse est syndiqué à plusieurs endroits, les LLM peuvent s'accrocher à la version la plus facile à analyser. Une citation sur un site agrégateur encombré (avec des publicités, des commentaires) pourrait être ignorée au profit d'une version de base de données textuelle diffusée en masse qui manque d'attributions. Cela a été observé lorsque ChatGPT a cité des copies syndiquées (Source: www.searchenginejournal.com). Les organisations de presse devraient s'assurer que les citations sont non seulement correctement attribuées, mais que les copies syndiquées maintiennent également ces attributions (et que les *scrapers* de sites les voient).
- **Sources contradictoires** : Lorsque deux médias publient la même citation avec de légères différences, les LLM peuvent les traiter comme distinctes. Sans une désambiguïsation robuste, la même citation pourrait être « stockée » sous différents noms d'orateurs dans le modèle. La cohérence dans la formulation et les balises de source entre les médias réduirait cette confusion.

En somme, **plus l'attribution d'une citation est bonne et claire dans un article de presse, plus il y a de chances qu'un LLM la reconnaisse et la référence correctement**. Inversement, une citation négligente – nom de l'orateur manquant, attributions vagues (« un expert a dit ») ou contexte non spécifié – rend la « compréhension » du LLM moins fiable. Comme l'a dit un praticien du SEO, le contenu doit être « *clair, sans ambiguïté* » pour un LLM (Source: mtsoln.com). C'est exactement ce à quoi aspire une attribution journalistique solide. Ainsi, une synergie existe : attribuer minutieusement les citations satisfait à la fois l'humain et la machine.

Analyse des données et résultats quantitatifs

Pour étayer ces observations, nous présentons des données quantitatives issues de recherches et d'expériences publiées. Notamment, des études d'analyse de contenu journalistique et d'évaluation de l'IA fournissent des chiffres pertinents, que nous résumons dans les tableaux et descriptions ci-dessous.

Citations et exagération dans l'actualité (Bossema et al.)

Une étude clé de Bossema *et al.* (2019) a analysé des milliers de communiqués de presse sur la santé et d'articles de presse connexes au Royaume-Uni et aux Pays-Bas. Elle fournit des statistiques détaillées sur les pratiques de citation :

TYPE DE PUBLICATION / PAYS	ARTICLES AVEC ≥1 CITATION	ARTICLES AVEC DE NOUVELLES CITATIONS INDÉPENDANTES	RAPPORT DE COTES D'EXAGÉRATION (PAS DE NOUVELLE CITATION)
Communiqués de presse RU (2011)	99,1 % (Source: pubmed.ncbi.nlm.nih.gov)	-	-
Articles de presse RU (2011)	88,6 % (Source: pubmed.ncbi.nlm.nih.gov)	7,5 % (Source: pubmed.ncbi.nlm.nih.gov)	2,6× (Source: pubmed.ncbi.nlm.nih.gov)
Communiqués de presse NL (2015)	84,5 % (Source: pubmed.ncbi.nlm.nih.gov)	-	-
Articles de presse NL (2015)	69,7 % (Source: pubmed.ncbi.nlm.nih.gov)	7,0 % (Source: pubmed.ncbi.nlm.nih.gov)	2,6× (implicite) (Source: pubmed.ncbi.nlm.nih.gov)

Tableau 1 : Prévalence des citations dans les communiqués de presse et les actualités sur la santé, et l'effet de l'inclusion de citations d'experts externes sur l'exagération des affirmations dans les articles (Source: pubmed.ncbi.nlm.nih.gov).

Observations clés de ce tableau :

- Presque tous les **communiqués de presse** (99 % au Royaume-Uni, 84,5 % aux Pays-Bas) contiennent au moins une citation (Source: pubmed.ncbi.nlm.nih.gov), généralement des auteurs de l'étude ou d'experts affiliés.
- Une forte proportion des **articles de presse** couvrant ces communiqués de presse contenaient également des citations (88,6 % au Royaume-Uni, 69,7 % aux Pays-Bas) (Source: pubmed.ncbi.nlm.nih.gov). Cependant, dans la plupart de ces cas, les citations étaient *tirées des communiqués de presse*.
- Seulement environ **7 à 8 %** des articles de presse ont introduit une *nouvelle* citation d'expert indépendant non présente dans le communiqué de presse (Source: pubmed.ncbi.nlm.nih.gov).
- De manière cruciale, les articles *sans* citation d'expert externe étaient **2,6 fois plus susceptibles** d'exagérer les affirmations causales que ceux *avec* une telle citation (Source: pubmed.ncbi.nlm.nih.gov). Cela implique que le simple fait d'avoir un expert externe comme orateur est corrélé à une plus grande retenue factuelle.

Bien que cette étude se concentre sur l'actualité de la santé, elle éclaire un schéma général : emprunter simplement des citations à une source est courant, mais ajouter un nouvel apport d'expert est rare – et pourtant percutant. Pour nos besoins, la conclusion pertinente est que **l'inclusion de citations externes affecte considérablement la qualité du contenu**. Par analogie, une IA lisant ces articles de presse pourrait trouver que les citations de vérification de la réalité (d'experts non directement impliqués) renforcent la fiabilité. Si les moteurs génératifs devaient choisir les extraits à croire, on pourrait émettre l'hypothèse que ceux avec des attributions indépendantes seraient plus sûrs. En effet, les créateurs de contenu qui souhaitent être cités par l'IA comme des autorités devraient viser à faire partie de ces voix « d'experts externes ».

Études sur la citation et la recherche par les LLM

Du côté de l'IA, des expériences récentes ont quantifié la capacité des LLM à citer ou à récupérer du contenu d'actualité. Le tableau 2 ci-dessous consolide les principales conclusions :

SYSTÈME IA / MÉTRIQUE	RÉSULTATS	SOURCE
Réponses CHATGPT-4 (étude)	GPT-4 a fourni des références pour toutes les réponses, mais seulement ~43 % étaient <i>entièrement exactes</i> ; ~56,7 % des sources citées étaient incorrectes ou inexistantes (Source: rankstudio.net).	[35]
Analogues GPT-4 (étude)	Dans une tâche générale, les modèles de type GPT-4 avaient ~90 % de citations factuelles (~10 % fabriquées) (Source: rankstudio.net).	[35]
ChatGPT Search	Lors de 200 essais de récupération, ChatGPT Search a donné des <i>réponses incorrectes</i> 153 fois (76,5 % des requêtes) (Source: www.searchenginejournal.com).	[37]
Chatbots IA (Tow Center)	Combinés, les chatbots ont répondu incorrectement à >60 % des 1 600 requêtes de citation (Source: www.cjr.org).	[57]
Perplexity.ai	Environ 37 % de taux d'erreur dans ces tests (Source: www.cjr.org).	[57]
xAI Grok-3	Environ 94 % de taux d'erreur dans ces tests (Source: www.cjr.org).	[57]
Tous les chatbots (général)	Ont souvent fabriqué des liens, cité du contenu syndiqué et rarement dit « Je ne sais pas » (Source: www.cjr.org) (Source: www.cjr.org).	[57]

Tableau 2 : Performance des grands modèles linguistiques et des chatbots dans la récupération et la citation de contenu d'actualité (d'après diverses études (Source: rankstudio.net) (Source: www.cjr.org).

Observations du Tableau 2 :

- Lors d'une évaluation de questions-réponses médicales, **GPT-4** a été invité à citer des sources pour chaque réponse. Il l'a fait, mais moins de la moitié des ouvrages référencés (43,3 %) étaient entièrement exacts (Source: rankstudio.net). Plus de la moitié étaient erronés ou fictifs (Source: rankstudio.net). Ainsi, même pour GPT-4, le bruit est substantiel sans une vérification minutieuse.
- Une autre étude a révélé que les modèles de type GPT-4 pouvaient atteindre environ 90 % d'exactitude factuelle dans les citations (Source: rankstudio.net). La grande divergence (43 % contre 90 %) souligne que le résultat dépend fortement du contexte, de l'incitation et du domaine.
- **ChatGPT Search (d'OpenAI)** était particulièrement sujet aux erreurs. Lors de 200 requêtes d'identification de citations, il en a mal attribué 153 (Source: www.searchenginejournal.com). Il a échoué 3 fois sur 4, souvent en liant au mauvais éditeur ou en manquant l'URL correcte.
- Le test plus large du Tow Center sur 8 outils a confirmé que le problème est systémique : les « réponses incorrectes » constituaient la majorité des réponses (Source: www.cjr.org). Certains outils comme Perplexity ont obtenu des résultats relativement meilleurs (environ 37 % d'erreurs), tandis que d'autres (xAI Grok-3) étaient désastreusement mauvais (94 % d'erreurs) (Source: www.cjr.org).
- Les modes de défaillance courants incluaient le *contournement des règles robots.txt*, la *citation d'articles syndiqués au lieu des articles originaux* et l'*invention d'URL* (Source: www.cjr.org). De nombreux chatbots répondaient avec assurance même lorsqu'ils n'avaient pas de source définitive, donnant rarement une qualification (Source: www.cjr.org).

Ensemble, ces données montrent que les LLM actuels ont un « problème d'attribution ». En termes pratiques : **les organes de presse ne peuvent pas compter sur l'IA générative pour gérer les citations avec précision**. Même si une citation est correctement présentée dans les actualités, l'IA peut toujours induire en erreur. D'un autre côté, amener un LLM à *mentionner votre contenu* est un défi. Vous pouvez produire d'excellentes citations, mais l'IA peut citer un concurrent ou une source alternative. Cela a conduit certains experts en SEO à avertir que *la visibilité dans la recherche IA n'est pas assurée par le seul trafic ou les liens, mais en étant « digne de citation »* (Source: mtsln.com) (Source: willmarlow.com).

Visibilité SEO vs LLM

En nous appuyant sur la littérature en marketing et en stratégie d'IA, nous comparons les facteurs SEO traditionnels aux critères axés sur les LLM :

CRITÈRE	ACCENT SEO TRADITIONNEL	ACCENT RECHERCHE LLM/IA (CITATION IA)
Autorité	Backlinks, Autorité de domaine, Notoriété de la marque (Source: mtsoln.com)	Expertise explicite (citations d'experts claires), autorité reconnue dans le contexte (Source: mtsoln.com)
Clarté	Optimisation des mots-clés, méta-descriptions	Langage clair, non ambigu et réponses directes (Source: mtsoln.com) (Source: mtsoln.com)
Contexte	Pertinence thématique via mots-clés	Adéquation contextuelle profonde à la requête de l'utilisateur, pertinence sémantique (Source: mtsoln.com) (Source: mtsoln.com)
Structure	Structure interne du site, balises HTML	Contenu fragmenté et modulaire (listes à puces, Q&R, TLDR) que les LLM peuvent facilement extraire (Source: willmarlow.com) (Source: mtsoln.com)
Trafic (CTR)	Des taux de clics élevés améliorent le classement	Non directement pertinent ; le succès est mesuré par le fait d'être cité, non cliqué (Source: willmarlow.com)
Actualité	Les mises à jour régulières du contenu améliorent le SEO	Utile pour l'actualité de l'information, mais la logique démontrable prime sur la récence dans les réponses
Citations	Références externes pour la crédibilité (facteur mineur)	Les citations directes ou l'attribution d'ancres sont très importantes (les LLM préfèrent les faits sourcés) (Source: rankstudio.net)

Tableau 3 : Comparaison des priorités entre le SEO traditionnel et la « citabilité » du contenu axé sur les LLM (Source: mtsoln.com) (Source: mtsoln.com).

Du Tableau 3, plusieurs tendances émergent qui sont pertinentes pour les citations dans les actualités :

- **Attribution d'experts comme autorité** : Plutôt que de simplement s'appuyer sur les backlinks ou le PageRank, les LLM recherchent des signaux d'expertise au sein du texte lui-même. Un journaliste citant un spécialiste avec toutes ses qualifications ajoute un marqueur d'expertise explicite que les LLM traitent comme une autorité textuelle (Source: mtsoln.com). Un titre de nom (par exemple, « Dr. Smith ») et une affiliation institutionnelle dans une citation signalent la confiance.
- **Clarté du contenu plutôt que mots-clés** : Le SEO valorisait autrefois les mots-clés ; les LLM valorisent les réponses en langage clair. Une citation percutante est souvent exactement le type de « phrase de réponse » qu'une IA souhaite (Source: mtsoln.com). Par exemple, un rapport économique citant « L'inflation est tombée à 2 % en juin », a déclaré Janet Yellen de la Fed, pourrait être plus précieux pour un LLM que des paragraphes remplis de bourrage de mots-clés.
- **Modularité** : Les articles traditionnels peuvent être digressifs ; le contenu ciblé par les LLM est plus modulaire. Les journaux utilisant des listes à puces ou des boîtes de questions-réponses (courantes dans les formats numériques) produisent de meilleurs extraits pour l'IA (Source: willmarlow.com). Les actualités avec des résumés TL;DR ou des encadrés de faits clés sont directement alignées avec ce qu'un LLM peut extraire.
- **Vérification** : En SEO, citer des sources est en grande partie un facteur E-A-T (Expertise/Autorité/Confiance), mais n'est considéré qu'indirectement par les algorithmes de classement. En revanche, les LLM **internalisent** essentiellement les affirmations factuelles et leurs attributions. La différence est qualitative : un algorithme SEO pourrait ne pas vérifier personnellement chaque fait, mais un LLM s'appuiera sur la mémoire du texte. Cela rend la méta-citation (par exemple, le lien ou l'attribution à des sources externes) plus influente. En effet, si un LLM est « ramené à l'entraînement de 2021 », il n'a pas de mises à jour en direct ; pour des réponses à jour, il s'appuie sur la récupération et les citations. Par conséquent, les sites souvent cités par les actualités pourraient en bénéficier indirectement.

En somme, ces analyses montrent que **les organes de presse souhaitant influencer les réponses de l'IA devraient traiter l'attribution des citations avec la même rigueur que pour les lecteurs**. Être l'autorité citée peut améliorer à la fois la confiance humaine et la « confiance par la citation » des machines. D'un autre côté, toute ambiguïté est encore plus pénalisée : là où un résultat Google typique pourrait encore se classer sans attribution claire, un LLM pourrait rejeter ou déformer une citation peu claire.

Études de cas et exemples concrets

Nous examinons ici des scénarios concrets illustrant comment les pratiques de citation des actualités ont eu un impact sur les sorties des LLM, ainsi que comment le comportement des LLM a à son tour exercé une pression sur les médias d'information.

IA de type ChatGPT et problèmes de citation

L'un des cas les plus médiatisés a concerné **ChatGPT Search d'OpenAI** (le mode de recherche Bing Chat/ChatGPT lancé en novembre 2024). Une étude du Tow Center de Columbia (rapportée par Search Engine Journal) a testé ChatGPT Search sur des citations d'articles de presse (Source: www.searchenginejournal.com). Sur 200 requêtes de citation, 153 réponses étaient incorrectes ou mal attribuées (Source: www.searchenginejournal.com). Par exemple, ChatGPT Search échouait souvent à nommer la bonne source d'information ou publication pour une citation donnée. Il privilégiait parfois de « plaire » à l'utilisateur avec une réponse plausible plutôt que de vérifier les faits. Ce comportement erroné soulève des inquiétudes pour les éditeurs : permettre l'inclusion de votre contenu par l'IA pourrait toujours placer votre marque dans un faux contexte.

Exemple : Le New York Times et ChatGPT

Lorsque les testeurs ont interrogé ChatGPT Search pour des citations du *New York Times*, le système a affiché des liens vers des copies non autorisées sur d'autres sites plutôt que le lien officiel du NYT (Source: www.searchenginejournal.com). L'incapacité de l'algorithme à attribuer correctement une citation au NYT (même s'il avait été vraisemblablement entraîné sur de nombreux ensembles de données provenant du NYT) signifiait que *l'histoire réelle du NYT n'était pas reconnue*. Au lieu de cela, ChatGPT a extrait une copie syndiquée pour la citation. Cela a alarmé les éditeurs : même si vous souhaitez que ChatGPT cite votre site (par exemple, en ne bloquant pas les robots d'exploration), il peut vous contourner si une source « plus facile » est accessible. Le Tow Center note que ces erreurs « remettent en question l'engagement d'OpenAI envers le développement responsable de l'IA dans le journalisme » (Source: www.searchenginejournal.com).

Syndication et exploration

À la base de bon nombre de ces problèmes se trouve la question de la **syndication**. Les services de presse (Getty, AP, Reuters) republient souvent du contenu sur plusieurs médias. Un LLM explorant le web pourrait indexer le texte brut d'une histoire de l'AP plutôt que le journal original, par exemple. Si la version syndiquée omet, par exemple, la signature « Reporter : Jane X. » ou réorganise les citations, un LLM pourrait attribuer le crédit à la mauvaise publication ou analyser les citations de manière incorrecte. Dans l'exemple ci-dessus, ChatGPT citant une version non-NYT suggère que le moteur de récupération sous-jacent a considéré le texte de l'AP comme la source dominante. (Note : le NYT était en litige avec OpenAI concernant les licences, ce qui a pu affecter l'accès). Pour les entreprises citées dans les actualités, la syndication signifie qu'être cité dans un média ne garantit pas que l'IA associera cette citation au nom de votre marque, si un flux « brut » sans contexte est indexé. Cela montre que **l'endroit** où une citation apparaît (y compris les métadonnées cachées) peut être aussi important que la citation elle-même.

Réaction des médias et perspective de l'industrie

Face à ces défis de l'IA, les médias grand public ont commencé à réagir. L'article du Guardian mentionné précédemment a décrit le comportement de ChatGPT comme un « tsunami de faits inventés » qui pourrait « saper les sources d'information légitimes » (Source: futurism.com). En réponse à des expériences comme celle du Tow Center, certains éditeurs ont publié des lignes directrices ou des déclarations. Par exemple, le Partnership on AI (une coalition d'organisations technologiques et médiatiques) a émis des recommandations sur la manière dont les journalistes devraient étiqueter le contenu de l'IA et rapporter de manière

responsable (Source: www.mdpi.com) (Source: www.mdpi.com). Certaines rédactions sont méfiantes : « *Générateurs de fausses informations* », c'est ainsi que Chris Moran du Guardian a qualifié les LLM qui attribuent mal les informations (Source: futurism.com).

Du point de vue de l'optimisation SEO/IA, certaines entreprises y voient une opportunité. Des guides stratégiques conseillent aux créateurs de contenu de structurer leur matériel de manière à ce que l'IA les cite « naturellement » (Source: createandgrow.com) (Source: willmarlow.com). Par exemple, un article de blog sur les mentions de LLM recommande de devenir « la source de référence que l'IA souhaite naturellement citer » (Source: createandgrow.com). Les praticiens suggèrent de créer du contenu d'ancrage (comme des FAQ détaillées) que les LLM peuvent fragmenter en réponses, et de s'assurer que le nom de votre marque apparaît dans ces segments répondables (Source: searchengineland.com) (Source: willmarlow.com).

Cependant, les données nous indiquent que cela reste expérimental. Les mêmes tests de moteurs de recherche du Tow Center montrent que tous les agents LLM actuels sont « *généralement mauvais pour refuser de répondre* », donnant plutôt des réponses fausses et étonnamment confiantes (Source: www.cjr.org). Même les modèles premium (C4, GPT-4o) n'ont pas fait exception. Beaucoup de ces systèmes effectuent explicitement une recherche web en arrière-plan, mais « réécrivent » ensuite la réponse avec des citations inventées s'ils ne trouvent pas de source.

Ainsi, un producteur de contenu pourrait avoir deux incitations quelque peu contradictoires :

- **Être prudent** : La précision journalistique est plus que jamais nécessaire. Une citation incorrecte peut être amplifiée par l'IA, nuisant à la confiance et à la marque.
- **Être averti en matière d'IA** : En même temps, les auteurs peuvent utiliser les citations et les métadonnées de manière stratégique pour s'aligner sur les préférences de l'IA (noms d'entités clairs, format de réponse structuré) afin d'augmenter les « mentions ».

La principale leçon est que les créateurs de contenu ne devraient *pas* supposer que les LLM attribueront magiquement les citations correctement. Au lieu de cela, ils devraient s'assurer que leur contenu contient le moins d'ambiguïtés possible. Concrètement, cela pourrait signifier vérifier les citations, fournir un contexte riche, et même adopter le balisage SEO/IA (comme lms.txt ou l'accès API) lorsque cela est possible. Dans le domaine du marketing, un tableau de suggestions a émergé (de la part de leaders d'opinion en SEO) sur la manière de cibler la visibilité des LLM. Celles-ci incluent les tactiques d'utilisation de phrases de requête exactes dans les en-têtes (anticipant les invites des utilisateurs), la création de « blocs de réponse » autonomes, et l'intégration de citations vers des sources réputées (Source: mtsoln.com) (Source: willmarlow.com). De telles stratégies soulignent indirectement la valeur de citer des figures d'autorité : chaque citation est, en effet, son propre mini-bloc de réponse qui peut être repris par une réponse d'IA. Une citation bien attribuée avec un titre comme « **Qu'a dit [Expert/Rapport] à propos de X ?** » est littéralement structurée pour la consommation par l'IA.

Tableau : Comparaisons et observations

Pour résumer les points de données clés, nous présentons le tableau suivant :

CONTEXTE	STATISTIQUE / CONSTAT	SOURCE
Citation d'actualités (Engagement médiatique)	Crédibilité la plus élevée lorsqu'un seul fonctionnaire gouvernemental non partisan est cité (Source: mediaengagement.org) (vs. sources partisans perçues comme biaisées).	[23]
Citation d'actualités (Bossema et al.)	Seulement 7 à 8 % des articles de presse sur la santé ont ajouté de nouvelles citations d'experts ; l'absence de citation d'expert a multiplié par 2,6 les chances d'exagération (Source: pubmed.ncbi.nlm.nih.gov).	[29]
Précision des citations LLM (GPT-4)	GPT-4 avec des invites a cité des sources dans toutes les réponses, mais seulement environ 43 % étaient entièrement correctes (Source: rankstudio.net).	[35]
Précision des citations LLM (analogue GPT)	Environ 90 % des citations étaient factuelles (10 % fabriquées) lors d'un test sur un large domaine (Source: rankstudio.net).	[35]
Taux d'erreur de ChatGPT Search	76,5 % des requêtes de source de citation ont été répondues incorrectement (153/200) (Source: www.searchenginejournal.com).	[37]
Taux d'erreur des chatbots IA (global)	>60 % de réponses incorrectes dans les tâches de récupération (Source: www.cjr.org).	[57]
Taux d'erreur de Perplexity.ai	Environ 37 % des requêtes ont été répondues incorrectement (Source: www.cjr.org).	[57]
Taux d'erreur de xAI Grok-3	Environ 94 % d'erreurs (Source: www.cjr.org).	[57]

| Exemples d'hallucinations de l'IA | ChatGPT a inventé des articles complets et des citations qui n'ont jamais existé (Source: futurism.com). | [49] |

Tableau 4 : Observations quantitatives sélectionnées concernant l'attribution de citations dans les actualités et le comportement des LLM.

Ces chiffres confirment que l'intersection de la citation d'actualités et de la récupération par les LLM est actuellement semée d'inexactitudes. En particulier, le fait que même les modèles de pointe présentent une précision de citation très variable (43 à 90 %) (Source: rankstudio.net), et que les interfaces de chat se trompent 60 à 90 % du temps (Source: www.searchenginejournal.com) (Source: www.cjr.org), devrait alerter les créateurs de contenu. Ils doivent partir du principe que les LLM sont sujets à la *distorsion* du contenu – et doivent donc prendre des mesures (par des citations précises) pour atténuer ce risque.

Implications et Orientations Futures

Notre analyse révèle de profondes implications pour de multiples parties prenantes :

Pour les journalistes et les rédactions : L'attribution des citations a toujours été essentielle à l'intégrité journalistique. Désormais, elle influence également la manière dont les systèmes d'IA mentionneront ou omettront le contenu. Compte tenu des limites actuelles des LLM, les journalistes devraient être **particulièrement vigilants** quant à l'exactitude. Un sourçage rigoureux et le contexte deviennent encore plus importants, car les erreurs pourraient être amplifiées par les réponses de l'IA. Certaines organisations de presse révisent déjà leurs normes : par exemple, le Partenariat sur l'IA recommande d'étiqueter clairement le contenu généré par l'IA et de pratiquer la « conscience de la source » (Source: www.mdpi.com). Les rédactions pourraient adopter des outils de détection de l'IA ou des avertissements pour le contenu susceptible d'être fourni à l'IA. De plus, les directives légales et éthiques évolueront : les questions de droit d'auteur et de données d'entraînement de l'IA sont déjà en jeu. Une attribution correcte peut aider à prévenir les litiges fondés sur la propriété intellectuelle et la diffamation.

Pour les éditeurs de contenu (marques/experts) : Être cité par les médias se traduit par une visibilité au-delà de la presse écrite. Si une marque ou un expert est cité dans des actualités réputées, même si cela ne génère pas de clics, cela peut augmenter son empreinte dans les LLM. Les stratégies de SEO/RP pourraient évoluer non seulement pour obtenir des citations pour les lecteurs humains, mais aussi pour *s'assurer que ces citations sont structurées pour la découverte algorithmique*. Par exemple, une équipe de RP pourrait encourager les journalistes à toujours inclure le titre complet d'une personne et une référence d'audience, plutôt qu'une citation vague. Les spécialistes du marketing suivront non seulement les classements Google, mais aussi le « Rang 0 » – c'est-à-dire si les chatbots IA les mentionnent. Des outils de surveillance des mentions de marque dans les réponses des LLM émergent (Source: mtsolv.com). La notion de « communiqué de presse pour l'IA » est susceptible de s'imposer : l'élaboration de citations de presse en pensant à l'IA (déclarations concises et directes) pourrait devenir une compétence de niche.

Pour le SEO et les spécialistes du marketing numérique : L'essor des citations par l'IA exige un changement de tactiques d'optimisation. Le netlinking traditionnel reste pertinent (pour être découvert dans les données d'entraînement), mais l'accent est mis sur l'*association d'entités* et la *structure de contenu citable*. Les stratégies de contenu incluent désormais souvent la création de blocs de questions-réponses clairs, de balises de schéma (schema markup) et de recherches ou données uniques (pour inciter d'autres à vous citer) (Source: mtsolv.com) (Source: willmarlow.com). Certains proposent même une « stratégie de citation LLM » analogue aux stratégies de backlinks (Source: www.higoodie.com) (Source: davidmelamed.com). La formation de partenariats pour la distribution de contenu (par exemple, écrire dans des revues sectorielles où les scrapeurs d'IA recherchent, ou des citations Wikipédia) est également attrayante, afin d'ensemencer un contenu faisant autorité que les chatbots peuvent atteindre. Les spécialistes du marketing doivent également considérer une nouvelle métrique : le *CTR IA*, ou la fréquence à laquelle leur site est directement référencé par les réponses de l'IA (même sans clic). Cela pourrait influencer la budgétisation et la planification du contenu dans les années à venir.

Pour les utilisateurs et la société : Du côté des utilisateurs, ces développements ont des effets mitigés. Idéalement, des systèmes LLM bien conçus fourniraient des réponses concises avec un sourçage transparent, favorisant la confiance des utilisateurs et leur faisant gagner du temps. En pratique aujourd'hui, les utilisateurs risquent d'être induits en erreur par des « faits » et des citations d'IA affirmés avec confiance mais faux. Une éducation aux médias est nécessaire : les utilisateurs devraient vérifier les citations fournies par l'IA par rapport aux articles originaux. Les journalistes et les éducateurs doivent apprendre aux gens à traiter les réponses de l'IA comme des sources non vérifiées tant qu'elles n'ont pas été contrôlées, à l'instar des débuts des moteurs de recherche. Il y a aussi une dimension d'équité sociale : si seules les grandes organisations peuvent se permettre d'être compatibles avec l'IA (avec du contenu structuré et des accords de licence), les petits médias pourraient être marginalisés. Les personnes interrogées par le Tow Center craignent que les dirigeants des entreprises d'IA ne négligent le contenu de niche précieux des petits éditeurs locaux (Source: mediaengagement.org). Assurer la reconnaissance des diverses voix médiatiques dans les réponses de l'IA est un défi émergent (et une question d'équité).

Pour les développeurs et les plateformes de LLM : Ces conclusions placent la responsabilité sur les concepteurs de modèles d'IA. Il est clair qu'une meilleure intégration des citations est nécessaire. Les approches incluent la génération augmentée par récupération (RAG) avec une meilleure provenance, le filigrane des sorties et des modes de refus plus conservateurs. Certains travaux (comme les cadres de citation étudiés (Source: www.themoonlight.io) explorent des correctifs architecturaux, mais le déploiement dans les outils grand public a pris du retard. Les décideurs politiques et les plateformes pourraient éventuellement exiger un « droit de refuser les hallucinations » ou des routines standardisées de vérification des sources. Par exemple, des partenaires comme le Partenariat sur l'IA promeuvent des normes pour la collaboration entre les rédactions et l'IA. Les aperçus d'IA prototypes de Google affichent déjà des notes de bas de page, mais même celles-ci renvoient souvent à des copies syndiquées. Idéalement, les systèmes génératifs devraient citer clairement, ou au moins dire « selon la source X... » seulement lorsqu'ils sont sûrs. Jusque-là, nous constatons une tension : les rédactions veulent une utilisation sûre et respectueuse du contenu ; les outils LLM veulent de vastes données d'entraînement. Des entreprises comme OpenAI (qui autorise désormais les opt-out via robots.txt) commencent à écouter, mais les progrès sont en cours.

Recherches Futures et Questions Ouvertes

Ce domaine est très nouveau ; beaucoup reste inconnu. Certaines orientations futures incluent :

- **Tests empiriques des LLM avec des citations :** Des études systématiques pourraient évaluer comment différents styles de citation affectent la récupération par les LLM. Par exemple, écrire le même contenu mais avec des formulations d'attribution variées, puis interroger un LLM pour voir quelle version il choisit. De tels tests A/B informeraient quantitativement les

meilleures pratiques.

- **Synthèse d'actualités générée par l'IA** : À mesure que le journalisme génératif (articles écrits par l'IA) devient une réalité, comment l'attribution des citations sera-t-elle gérée ? Certains outils (par exemple, Lynx Insight chez Reuters (Source: reutersagency.com) rédigent déjà automatiquement des articles basés sur des données. S'assurer que ces brouillons d'IA insèrent correctement les citations pourrait bientôt nécessiter des vérificateurs de sources automatisés.
- **Littératie IA et contre-mesures** : Comment les lecteurs différencient-ils une citation factuelle d'une citation hallucinatoire dans une réponse d'IA ? La conception des interfaces utilisateur pourrait afficher des scores de confiance ou des chemins de provenance. La recherche en interaction homme-machine pourrait aider les utilisateurs finaux à mieux évaluer la production des LLM concernant les actualités.
- **Effets longitudinaux** : Avec le temps, si les LLM citent de manière erronée les actualités à plusieurs reprises, cela modifiera-t-il la perception du public ? Certaines analyses dystopiques mettent en garde contre une « tromperie sans fond » lorsque la propagande de l'IA se superpose à elle-même, avec de fausses citations alimentant les théories du complot (Source: futurism.com). L'étude de la diffusion de l'information à l'ère de l'IA est cruciale.
- **Cadres juridiques et éthiques** : Devrait-il y avoir des directives (voire des lois) concernant la manière dont les outils d'IA doivent attribuer les informations provenant des actualités ? Par exemple, imposer des normes de transparence pour les réponses de l'IA ou interdire les hallucinations de l'IA sur des sujets sensibles. Les normes de la communauté journalistique (vérité, exactitude) pourraient devoir être traduites en politiques technologiques.

Conclusion

Dans le paysage en évolution rapide de l'IA générative, **l'attribution des citations dans les articles de presse est apparue comme un facteur critique influençant la manière dont le contenu est traité par les LLM**. Notre recherche consolide les preuves issues des études médiatiques, des expériences d'IA et de la stratégie SEO pour montrer que les citations bien attribuées servent un double objectif : elles renforcent la confiance humaine et elles s'alignent sur les structures de contenu que les LLM préfèrent. Inversement, des attributions vagues ou incorrectes peuvent amplifier la confusion, car les LLM inventent ou attribuent facilement des citations de manière erronée lorsque les données sont confuses (Source: futurism.com) (Source: www.cjr.org).

Les principales conclusions incluent :

- Des études de crédibilité confirment que les articles de presse citant des responsables faisant autorité sont classés comme les plus fiables (Source: mediaengagement.org). Les actualités sans attribution claire sont perçues comme moins crédibles ou biaisées.
- L'analyse empirique de contenu montre que les articles avec des citations d'experts externes sont nettement *plus précis*, tandis que ceux sans sont sujets à l'exagération (Source: pubmed.ncbi.nlm.nih.gov).
- Du côté de l'IA, les outils basés sur les LLM fournissent souvent des citations incorrectes. ChatGPT Search a mal attribué des citations 76,5 % du temps lors d'un test (Source: www.searchenginejournal.com), et plusieurs chatbots IA ont collectivement mal identifié les sources d'actualités dans plus de 60 % des expériences (Source: www.cjr.org).
- Les cadres techniques SEO indiquent que les LLM privilégient la « clarté » et l'« adéquation contextuelle » du contenu (Source: mtsln.com) (Source: mtsln.com). Les segments structurés et autonomes (comme les citations correctement introduites) sont les plus susceptibles d'être extraits.

La convergence de ces conclusions implique : **Les rédactions et les créateurs de contenu devraient adhérer aux normes d'attribution les plus élevées, non seulement pour le bien des lecteurs, mais aussi pour assurer une consommation précise par les machines**. Citer des noms et des titres concrets, structurer clairement les passages (en les étiquetant potentiellement comme des blocs de questions-réponses) et fournir des métadonnées peut améliorer les chances que les systèmes d'IA fassent les choses correctement. De même, les développeurs d'IA ont la responsabilité d'affiner la manière dont leurs modèles gèrent les citations, afin d'éviter de saper le travail journalistique.

Pour l'avenir, les implications sont profondes. À mesure que de plus en plus de personnes dépendent des résumés générés par l'IA, même de petites erreurs de citation peuvent se propager à travers l'écosystème de l'information. Pourtant, il y a de l'espoir qu'une synergie soit possible : si les journalistes et les technologues collaborent – par exemple via le Partenariat sur l'IA ou les normes de

l'industrie – ils peuvent co-cr  ter des flux de travail o   le contenu des actualit  s reste    la fois digne de confiance pour le lecteur et compatible avec l'IA.

En conclusion, l'attribution des citations n'est pas seulement une pr  occupation stylistique ; elle fa  onne les *empreintes de connaissance* que les LLM tracent. En comprenant profond  ment cette interaction, les parties prenantes peuvent l'exploiter : les m  dias peuvent accro  tre leur port  e effective et leur cr  dibilit  , les entreprises peuvent obtenir une visibilit   l  gitime de l'IA, et la soci  t   peut insister sur la responsabilit      l'  re des r  ponses automatis  es. Le chemin vers des « mentions LLM » robustes ne fait que commencer, mais une citation m  ticuleuse en sera l'une des pierres angulaires (Source: mediaengagement.org) (Source: willmarlow.com).

R  f  rences

Toutes les affirmations et donn  es de ce rapport sont   tay  es par les sources suivantes :

- Bossema *et al.* (2019), *Citations d'experts et exag  ration dans les actualit  s sur la sant   : une analyse de contenu quantitative r  trospective* (Source: pubmed.ncbi.nlm.nih.gov).
- Center for Media Engagement (2025), *Citations et cr  dibilit   : comment les approches narratives fa  onnent les perceptions au-del   des clivages politiques* (Source: mediaengagement.org).
- Huang *et al.* (2025), *Attribution, Citation et Guillemets : Une enqu  te sur la g  n  ration de texte bas  e sur des preuves avec des grands mod  les linguistiques* (Source: www.themoonlight.io).
- Huang (2025), *Au-del   de la popularit   : le manuel pour dominer la visibilit   dans la recherche IA* (Source: mtsoln.com) (Source: mtsoln.com).
- Search Engine Journal (2024), *ChatGPT Search   choue au test d'attribution, cite mal les sources d'actualit  s* (Source: www.searchenginejournal.com).
- Futurism (2023), *Un journal alarm   lorsque ChatGPT r  f  rence un article qu'il n'a jamais publi  * (Source: futurism.com).
- Columbia Journalism Review (2024), *La recherche IA a un probl  me de citation* (Source: www.cjr.org).
- Krstovi   (2025), *SEO LLM expliqu   : comment faire citer votre contenu dans les outils d'IA* (Source: willmarlow.com) (Source: willmarlow.com).
- Todorov (2025), *Influencer les mentions LLM par un contenu strat  gique* (Source: createandgrow.com).
- Mercury Tech Solutions (2025), *Maximiser la visibilit   de l'IA : Comprendre la citation LLM* (Source: mtsoln.com) (Source: mtsoln.com).
- Reuters (2023), *Comment l'IA contribue    alimenter des actualit  s faibles chez Reuters* (Source: reutersagency.com).
- Shi & Sun (2024), *Comment l'IA g  n  rative transforme le journalisme : d  veloppement, application et   thique* (Source: www.mdpi.com).

(Les citations en ligne entre crochets renvoient directement au mat  riel source utilis   pour chaque affirmation.)

  tiquettes: mentions-llm, attribution-citations, recherche-ia, seo-pour-llm, credibilite-source, journalisme, strategie-contenu-ia

AVERTISSEMENT

Ce document est fourni    titre informatif uniquement. Aucune d  claration ou garantie n'est faite concernant l'exactitude, l'exhaustivit   ou la fiabilit   de son contenu. Toute utilisation de ces informations est    vos propres risques. RankStudio ne sera pas responsable des dommages d  coulant de l'utilisation de ce document. Ce contenu peut inclure du mat  riel g  n  r   avec l'aide d'outils d'intelligence artificielle, qui peuvent contenir des erreurs ou des inexactitudes. Les lecteurs doivent v  rifier les informations critiques de mani  re ind  pendante. Tous les noms de produits, marques de commerce et marques d  pos  es mentionn  s sont la propri  t   de leurs propri  taires respectifs et sont utilis  s    des fins d'identification uniquement. L'utilisation de ces noms n'implique pas l'approbation. Ce document ne constitue pas un conseil professionnel ou juridique. Pour des conseils sp  cifiques    vos besoins, veuillez consulter des professionnels qualifi  s.