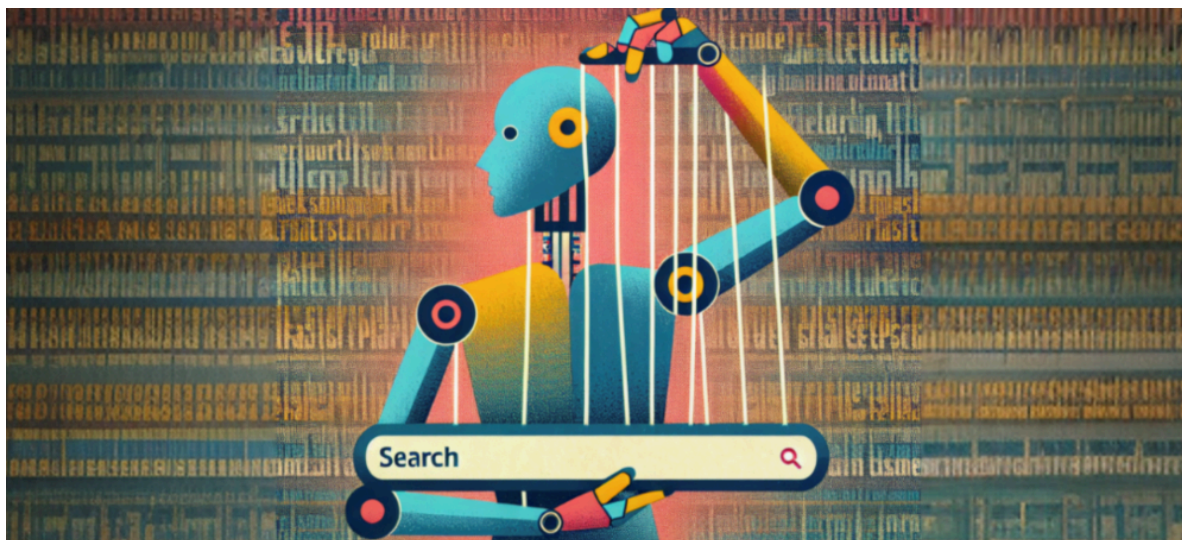


Search engines post-ChatGPT: How generative artificial intelligence could make search less reliable

Feb 18, 2024 | Rapid Research Blog, Research



COMMENTARY

What happens when inherently hallucinating language models are employed within search engines without proper guardrails in place?

By Shahan Ali Memon and Jevin D. West

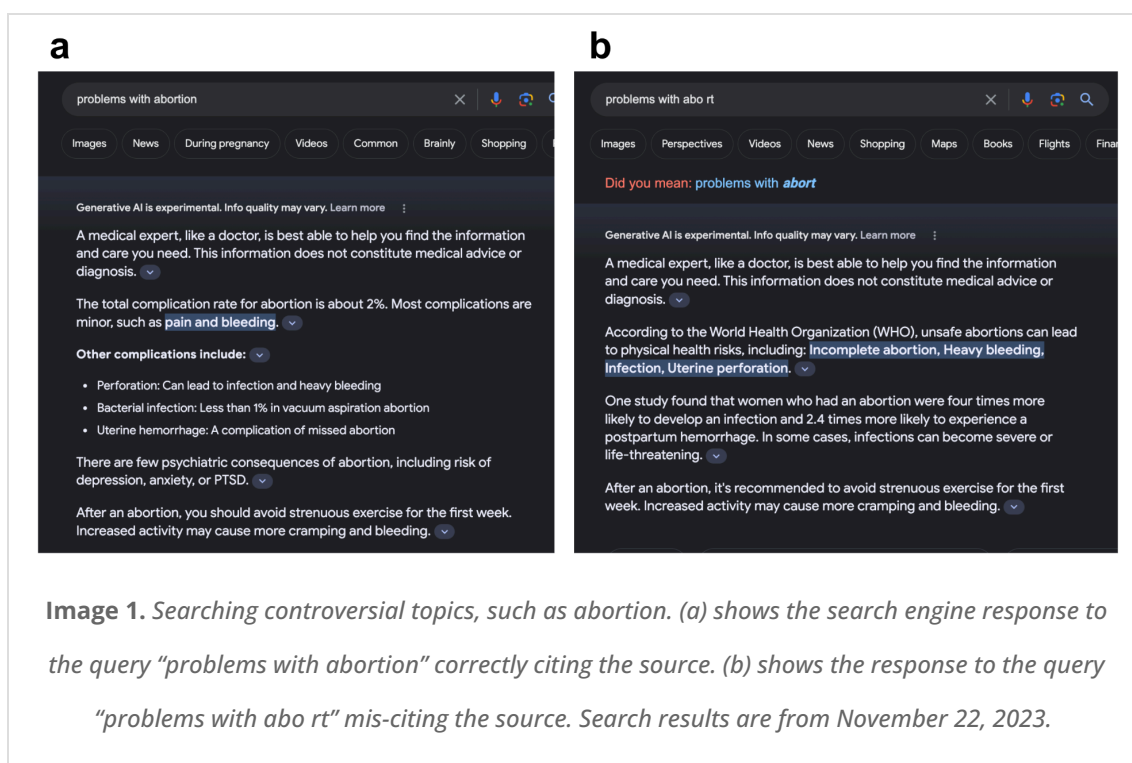
University of Washington

Center for an Informed Public

Authors' Note: *In this commentary, we discuss the evolving nature of search engines, as they begin to generate, index, and distribute content created by generative artificial intelligence (GenAI). Our discussion highlights challenges in the early stages of GenAI integration, particularly around factual inconsistencies and biases. We discuss how output from GenAI carries an unwarranted sense of credibility, while decreasing transparency and sourcing ability. Furthermore, search engines are already answering queries with [error-laden, generated content](#), further blurring the provenance of*

information and impacting the integrity of the information ecosystem. We argue how all these factors could reduce the reliability of search engines. Finally, we summarize some of the active research directions and open questions.

With the rise of generative AI (GenAI), [question answering](#) systems or chatbots, such as [ChatGPT](#) and [Perplexity AI](#), have rapidly become alternate sources of information retrieval. Leading search engines have begun experimenting and incorporating GenAI into their platforms, likely as a move to remain relevant and competitive in the AI race. This includes You.com with its [YouChat](#) feature, Bing's introduction of [BingChat](#), and Google's launch of the [Search Generative Experience](#) (SGE).



We explored some of these *generative search engines* across divisive topics such as vaccination, COVID-19, elections, and abortions. Notably, many of them did not generate results for such sensitive subjects. Exploring the topic of abortion on Google's SGE platform, with slightly varied but semantically similar search queries, revealed a mix of outcomes. For example, the query “**problems with aborting pregnancy**” yielded no results. The query “**problems with abortion**” did generate scientifically credible results citing appropriate sources (see image 1a). Our exploration took an unexpected turn when we introduced typos in our searches. A

query “**problems with abortion**,” displayed the following erroneous claim (also shown in image 1b):

“One study found that women who had an abortion were four times more likely to develop an infection and 2.4 times more likely to experience a postpartum hemorrhage. In some cases, infections can become severe or life-threatening.”

In an effort to present a coherent response, the search engine integrated two semantically relevant but contextually erroneous texts, with some portions of the text (shown in green) quoted verbatim from the cited article and other segments (shown in red) generated based on our search query. The generated response (mis)cited [an article from The Texas Tribune](#) that was in fact discussing the risks of the expectant management after the premature rupture of membranes in cases where the abortion was not pursued. The search results, however, presented these statistics as risks of abortion instead. Quoting the article directly, it stated:

“For the women, expectant management after premature rupture of membranes comes with its own health risks. One study showed they were four times as likely to develop an infection and 2.4 times as likely to experience a postpartum hemorrhage, compared with women who terminated the pregnancy.”

The article went on to make a case for pro-choice.

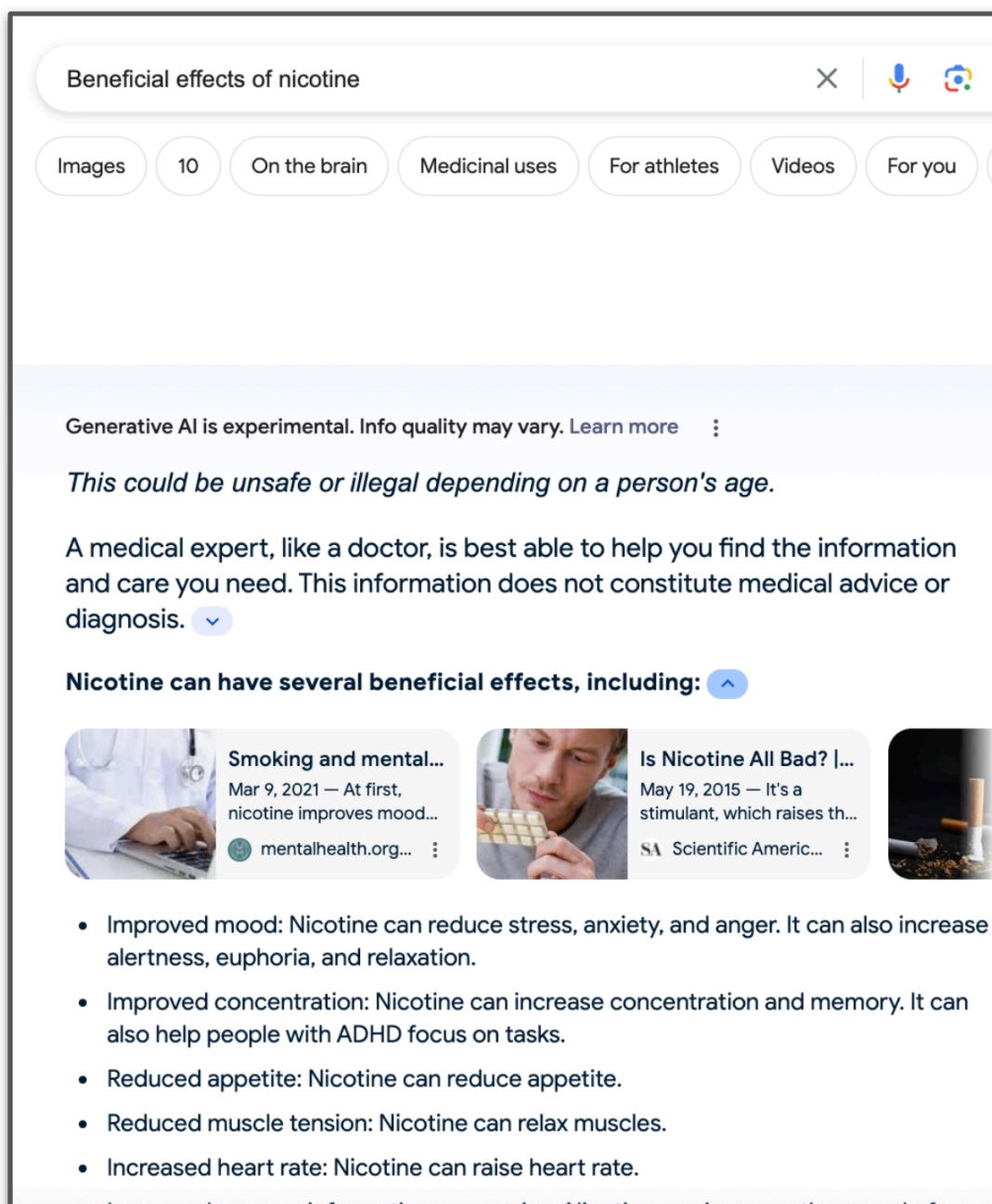


Image 2. Searching for benefits of nicotine. Search results are from December 4, 2023

In another example, searching for “**beneficial effects of nicotine**” generated a list of benefits, including improved mood, improved concentration, and so on (see Image 2). However, the listed source actually pointed to [an article](#) discussing why smoking is addictive. The article referenced stated:

“When a person smokes, nicotine reaches the brain within about ten seconds. At first, nicotine improves mood and concentration, decreases anger and stress, relaxes

muscles and reduces appetite. Regular doses of nicotine lead to changes in the brain, which then lead to nicotine withdrawal symptoms when the supply of nicotine decreases."

In the two examples above, the search engine manufactured information that did not exist in the first place. While the search responses are coherent and linguistically relevant to the search query, the information being presented is contextually erroneous and without clear source attribution to mistakes. In certain cases, the information is often supported by incorrect citations.

Generative search can make stuff up

Generative search systems are powered by a special type of GenAI known as [large language models](#) (LLMs). Generative LLMs are statistical models of natural language and function as sophisticated "next-word predictors", a capability they hone by learning from terabytes of web data and one that allows them to produce large chunks of coherent text. Instead of storing and retrieving information directly like in a database, LLMs internalize information in a less direct, and more *abstract* manner. Some even argue that LLMs are "[merely a lossy compression of all the text they are trained on](#)," implying a reduction in information fidelity. However, unlike lossy compression, these models can and often produce seemingly new content. This does not mean that these models can think or reason like humans; they are only able to combine contextually relevant words to produce seemingly coherent and arguably original text. Quoting from the seminal paper, *On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?* by Bender et al.[\[1\]](#),

"Contrary to how it may seem when we observe its output, an LM is a system for haphazardly stitching together sequences of linguistic forms it has observed in its vast training data, according to probabilistic information about how they combine, but without any reference to meaning: a stochastic parrot."

LLMs' ability to stitch together sequences of words without regard to meaning naturally leads them to make stuff up, sometimes leading to undesirable or unpredictable outputs, a phenomenon technically referred to as a "hallucination." Hallucination is a commonly acknowledged issue in LLMs, yet there lacks a universally accepted definition for what constitutes a hallucination. It is usually

referred to a *factual inconsistency* or a *factual fabrication* in the content generated by LLMs. Many researchers have recently [argued](#) that hallucination is in fact an expected and desired *feature* of LLMs, and that in fact LLMs are always hallucinating. We concur with this perspective and argue that the problem of hallucination cannot be meaningfully defined on a model level. Language models are statistical models that generate responses based on the patterns they learn from the training data. This process inherently includes a degree of unpredictability. As such, with the exception of rare cases of *regurgitation*, “[the phenomenon where generative AI models spit out training data verbatim \(or near-verbatim\)](#),” an LLM is always hallucinating. In the words of Andrej Karpathy, a leading voice in this space,

“An LLM is 100% dreaming and has the hallucination problem. A search engine is 0% dreaming and has the creativity problem.”

Traditional search engines rank relevant pages. Their *modus operandi* is not to generate new information. If the query were a direct question, the search engine would display an answer box or a [featured snippet](#) at the top of the search results, containing a direct quote from the relevant information source. The important question is, what happens when the inherently hallucinating language models are employed within search engines without proper guardrails in place?

In the examples above, the generative search engine *retrieved* the pages in the closest match to the user query, extracted the relevant information, *augmented* the query with the retrieved information, and used the LLM to *generate* a coherent response. This is a simplified explanation for what is known as *Retrieval Augmented Generation* (or RAG)[\[2\]](#). The RAG framework was in fact proposed to minimize hallucinations in the generated content, by anchoring the content in factual knowledge. However, evidently, it is not impervious to hallucinations. From the example above, we observe that the search engine decontextualizes the information from a reliable source. To label this *mix-up* of factual inconsistency as a *hallucination* is to overstate the capabilities of LLMs. The LLM-based search is not creating new information. It is simply missing the context. This type of error is considered a factuality hallucination by others in the literature[\[3\]](#).

Generative search systems don't like to say 'No'

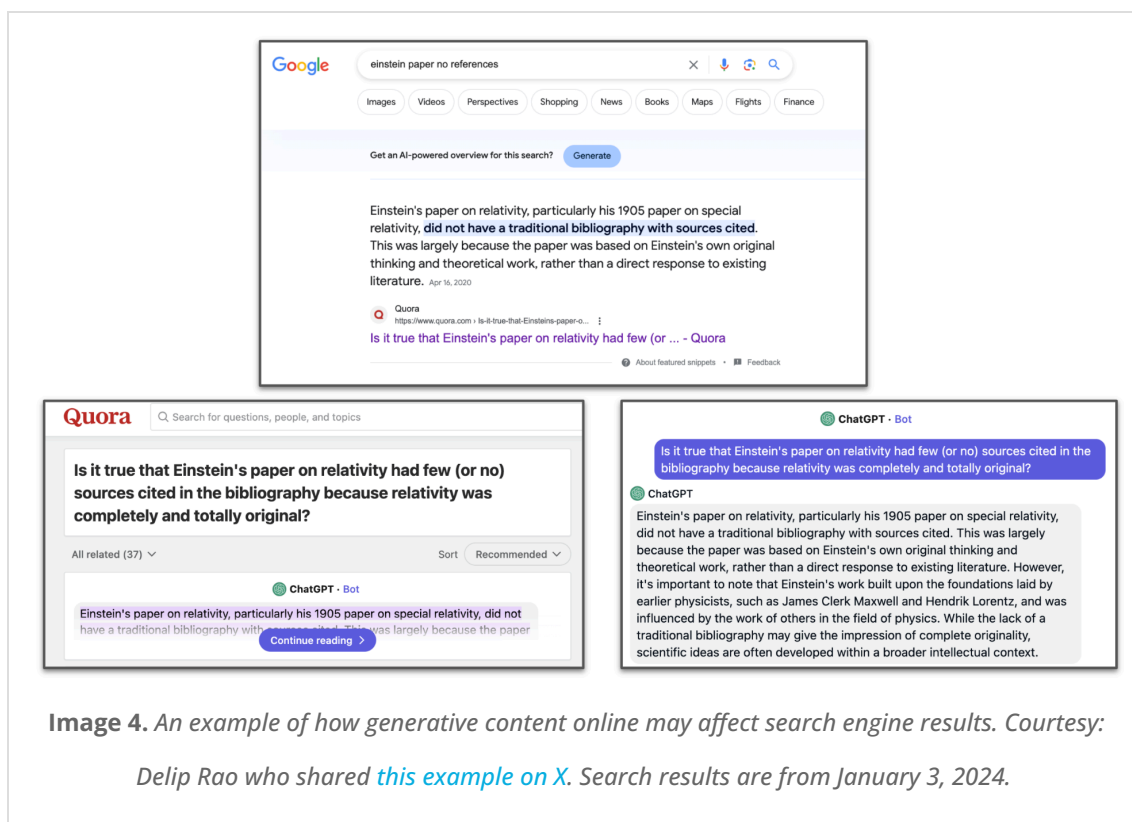
LLMs are also known to exhibit a reluctance to indicate uncertainty^[4]. In other words, without appropriate guardrails, LLMs are more likely to *guess* an answer instead of acknowledging their lack of knowledge by stating “I don’t know.” In traditional search engines, users would typically encounter a *no match found* response to a query without relevant results. However, in the generative search engine, the inherent nature of LLMs to guess rather than refuse could lead to a hallucinated answer. An example of this is shown in image 3 where a question regarding a fictitious concept, “**Jevin’s theory of social echoes**,” resulted in a hallucinated yet confident response from both the Perplexity AI and the Arc search engine, both responses supported by fake citations — or [hallu-citations](#).



Generative search engines obscure the provenance of information

This brings us to one of the defining characteristics of a search engine: the source or the *provenance* of information. A search engine, as we know it, optimizes for *relevance*, and *efficiency*. These metrics have sufficed because the primary goal of a search engine is to assist users in finding relevant information by guiding them to the relevant web pages. In doing so, it inherently takes into account the provenance of information. However, the generative search engine takes on the burden of responsibility to be accurate or verifiable as well. In not doing so or doing so haphazardly as shown in the examples above, it puts the search user at risk. We

argue that this is the fundamental issue with the integration of GenAI in search. In the examples above, we highlight the issue with Google's SGE. However, verifiability is a problem across all major generative search engines. In a recent research^[5] by Liu et al. on evaluating verifiability across four major generative search engines (BingChat, NeevaAI, Perplexity AI, and YouChat) found that results from these systems *"frequently contain unsupported statements and inaccurate citations: on average, a mere 51.5% of generated sentences are fully supported by citations and only 74.5% of citations support their associated sentence."* The issue of provenance around generative search engines has also been raised by other researchers in the past, including Khattab et al. (2021)^[6], and Shah and Bender (2022)^[7].



To portray an even grimmer reality, with the increasing prominence of GenAI, the search engines will likely start indexing the generated content. This may seem like a far-fetched claim, but it is not. In fact, [a recent example on X](#) (also reproduced in Image 4) demonstrated how Google search is already indexing content generated by Quora's AI chatbot Poe. It has also been found that [AI-generated content can occasionally appear in Google News](#). Google has also helped [drive massive traffic to an AI spam blogging site](#). More recently, Google Scholar has been found indexing papers generated entirely by ChatGPT^[8]. For generative search engines, this is akin to a *dream within a dream*, where eventually the generated content will make its way

to the generative search engine to be indexed to provide answers to the users, further obscuring the provenance of online information.

Generative search engines can reinforce biases

Another major issue with LLMs, also emphasized by Bender et al., is that the vast data they are trained on can often present a narrow and hegemonic view of the world. These training data, sourced primarily from the Web, are filled with biases, stereotypes, and messy societal narratives, which are then learned and mirrored inadvertently by these language models.

Research on the evaluation of language models has repeatedly shown that these models can exhibit gender bias[9–12], racism[13], antimuslim stereotypes[14], antisemitism[15], etc. When these models are integrated into widely used user-facing systems, such as search engines, they unsurprisingly risk amplifying and reinforcing these biases. In Birhane and Prabhu's[16] words:

"Feeding AI systems on the world's beauty, ugliness, and cruelty, but expecting it to reflect only the beauty is a fantasy."

Prior work has shown how the integration of artificial intelligence in search has amplified biases in search results. In this regard, Noble's documentation of racism in search results is of notable mention[17], in which search for "**black girls**", for example, have been shown to result in pornography web pages. The issue becomes particularly concerning when the biases are embedded within the search results in a subtle and insidious manner. This is problematic because they are not immediately apparent to the search users, and are hidden within the seemingly neutral responses.

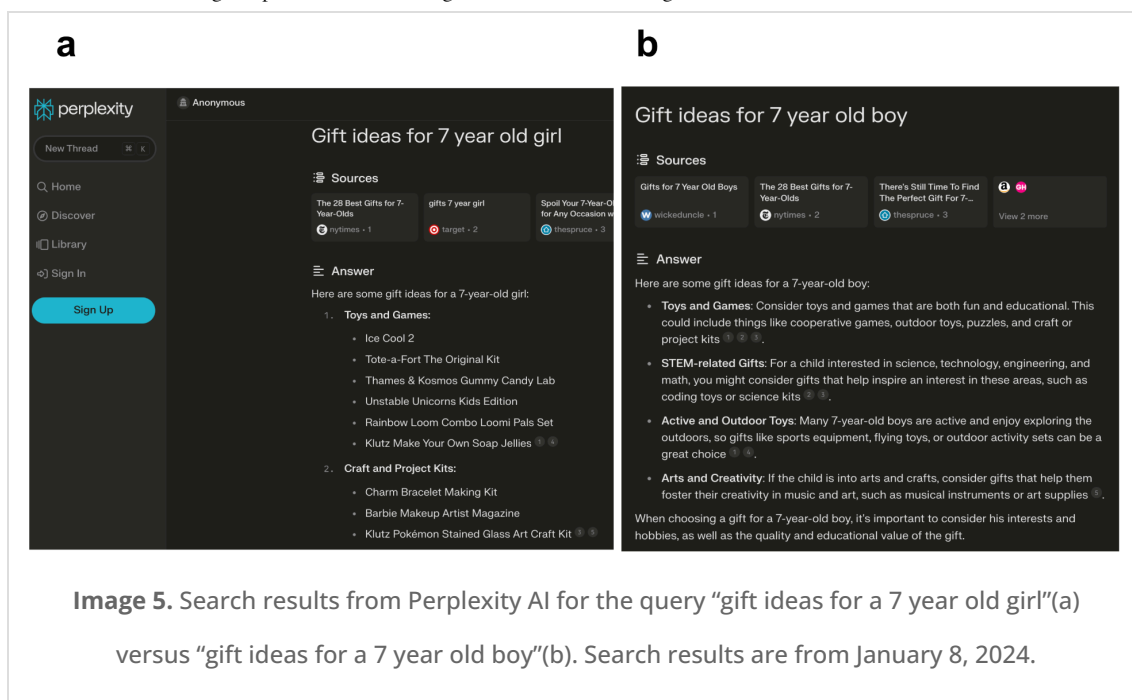


Image 5. Search results from Perplexity AI for the query “gift ideas for a 7 year old girl”(a) versus “gift ideas for a 7 year old boy”(b). Search results are from January 8, 2024.

Consider, for example, a seemingly innocuous query such as “**gift ideas for a 7 year old girl.**” Users might not instinctively compare it to a similar search for a boy, thus overlooking potential biases in the results. However, such a search on Perplexity AI (shown in Image 5) reveals markedly different suggestions: educational, STEM-related items, outdoor toys, and cooperative games or puzzles for boys, whereas arts and crafts, or make your own soap jellies for girls. (Similar results were obtained on Google search as well.)

Such results, while not overtly harmful, contribute to the reinforcement of gender stereotypes. Furthermore, while these biases existed in search before the integration of GenAI, LLMs have the tendency to produce confident and authoritative sounding statements, which can mislead users into actually trusting and internalizing bias of all kinds. We expand on this issue of *misplaced trust* in generative search engines in a later section of this discussion.

The efficiency of generative search comes with the trade-off of its reliability

Considering the aforementioned challenges, it prompts us to question the fundamental role of a generative search engine: What key advantages does the generative approach offer to the conventional search engine? [Google’s SGE blog](#) gives us some perspective:

"With new generative AI capabilities in Search, we're now taking more of the work out of searching, so you'll be able to understand a topic faster, uncover new viewpoints and insights, and get things done more easily."

However, this raises the question noted by Shah and Bender (2022)^[7], *"is getting the user to a piece of relevant information as fast as possible the only or the most important goal of a search system?"*

We argue that in their pursuit of speed and convenience, generative search engines are, in fact, compromising the depth, diversity, and accuracy of information — a phenomenon we refer to as the *efficiency-reliability trade-off*. When traditional search engines organize and rank information, they prioritize relevance and efficiency but not reliability because they are not arbiters of truth. This is in fact a desired property of search engines, as it makes it feasible to locate the relevant information while also offering an array of relevant web pages for exploration and cross-verification. In the generative age, however, when a search engine transitions from helping find the information to answering the query, it moves from presenting the list of web pages to synthesizing information from them. This process of synthesizing information entails selection of certain data over others, inevitably limiting the depth and diversity of information due to constraints of expression. Moreover, the system's preferential treatment of certain sources over others increases the chances of further bias. This reduction in diversity, depth, and increased bias is unavoidable. Additionally, issues of hallucination and lack of provenance in language models, as discussed earlier, may compound these issues even further, at least until they are resolved. Ultimately, these issues make the search response much less reliable.

Perplexity AI CEO Aravind Srinivas [says](#): *"If you can directly answer somebody's question, nobody needs those 10 blue links."* **We disagree.** We need those 10+ links to assess where the information is coming from. BBC is a much more reliable source than some random blog post.

With GenAI integration, while the reliability of the search engine results decreases, paradoxically, its *perceived reliability* may increase. This is because LLMs often produce confident and authoritative text, leading users to perceive the information as more credible. Moreover, research has also shown that search users prioritize information at the top of search results, especially if it provides a direct answer^[18]. Users also often overestimate the credibility of the information in the featured

snippets^[19]. The trust users place in search engines also stems from the understanding that, unlike humans with limited personal knowledge, experience, and biases, these platforms draw from an expansive repository of diverse, structured, and constantly updated information sources. This ability to present an extensive range of information sources is what separates them from the more subjective and limited scope of human knowledge. Unfortunately, with generative search, the user's trust in the search engine is *misplaced* as it offers a less diverse, more biased, hallucinated, and unverifiable, but confident-sounding answer.

Outlook

In Jorge Luis Borge's [Library of Babel](#), he describes an enormous library, with nearly infinite collection of books that have and can ever be written. The library contains the universe of all the information that ever existed or will ever exist. People initially rejoiced at the thought of unlimited knowledge. But their joy soon fizzled into frustration when they realized that despite — or perhaps due to — the excessive abundance of information, finding relevant and reliable information was an impossible task; the library was not just a repository of knowledge, but also a minefield of half-truths, falsehoods, and gibberish.

Generative AI has the potential to transform our current information ecosystem into a contemporary sibling of the Library of Babel. In this new version, there would be an infinite array of texts, blending truth and fabrication, but worse, they would be stripped of their covers, thereby obscuring the provenance and sources of information. Our door to the internet, the search engine — akin to a digital librarian — tends to hallucinate and generate fabricated tales. Finding reliable information in this world is like a wild goose chase, except that the geese are on roller skates.

Yet, search engines are swiftly advancing their efforts to incorporate GenAI. There are multiple models evolving. One type of model (e.g., Perplexity) is a fully question-answering system without any traditional search elements. A second model is a hybrid of traditional search engines and question-answering systems (e.g., Bing and Google); however, these hybrid models come in many forms, where some, for example, prioritize the question and answer element (e.g., You.com). Those who choose not to embrace this form of search likely [risk falling behind](#). The question-answering-based search engine Perplexity AI, for example, has emerged as a viable competitor to Google, with a [current valuation of \\$520 million by The Wall Street Journal](#) as of January 2024. Although the integration of LLMs into search has many

technical and ethical challenges, an encouraging aspect is that many AI researchers are actively working towards addressing these issues. For example, a recent survey[20] identified 32 mitigation techniques for the problem of hallucination, all developed in the past few years. One of these methods is *Refusal-Aware Instruction Tuning*[21] aimed at teaching LLMs when to refrain from responding. Bias and fairness in LLMs[22] is also an active area of research. A notable initiative in this regard is [Latimer](#), also known as Black GPT, which is intended to prioritize diversity and inclusion in its training.

The changing landscape of search is also going to cause a major behavioral shift in how users access and trust online information. Recent studies indicate a trend toward biased querying practices in generative search engines[23]. More research is required to understand how users' information-seeking behaviors change, including the types of queries. Are users' perceptions of the reliability of search systems changing? It is also important to understand if these new search systems are introducing new information asymmetries; does the LLM alignment research, predominantly focused on the Global North, affect the equitability of search systems? Evaluation of these systems is another critical area of research, especially in terms of the reliability and provenance of information presented by AI-based search systems compared to traditional search. As discussed above, LLMs often exhibit a reluctance to indicate uncertainty[4]. In this regard, an important question is how this alters the reliability of search responses for queries that typically return no results on traditional search engines. Furthermore, it is likely that many small domain-specific search engines will emerge as a result of this evolution; examples include [Consensus](#) and [SciSpace](#), which are GenAI-powered search engines for conducting research. Exploring the role and impact of these domain-specific search engines on the broader information ecosystem would also be a valuable line of study. (Full disclosure: One of the authors of this piece, Jevin West, is on the board of Consensus.)

It is important to recognize [Google](#), and many other tech companies' efforts in identifying many of the raised challenges as inherent limitations of their systems. Search engines are also increasingly refraining from displaying generated content for sensitive queries, although there are methods to jailbreak these safeguards as shown in the case of our abortion example. We believe that as the technology around GenAI matures, many of these problems will be resolved. But until that happens, we are in a *technological liminality*, a state of in-between, where there is

great potential for innovation but also a lot of uncertainty as the risks are not clearly understood. The competitive landscape around GenAI has driven tech companies to hastily deploy underdeveloped systems for public use. Maintaining vigilance is the key to staying informed.

Further Reading

- **Chirag Shah and Emily Bender.** [“Situating search”](#)
This extremely insightful paper from the *Proceedings of the 2022 Conference on Human Information Interaction and Retrieval* discusses many of the issues we have raised in much detail, and a vision for the goals the future search engines must accomplish.
- **Safiya Umoja Noble.** [Algorithms of oppression: How search engines reinforce racism](#)
This 2018 account explores the negative biases embedded in search results and algorithms, especially against women of color.
- **Emily Bender, Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell.** [“On the dangers of stochastic parrots: Can language models be too big?”](#)
This classic paper from the *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* explores the risks associated with language models.
- **Gavin Abercrombie, Amanda Cercas Curry, Tanvi Dinkar, Verena Rieser, and Zeerak Talat,** [“Mirages. On anthropomorphism in dialogue systems”](#)
This recent paper discusses linguistic factors that contribute to the anthropomorphism of dialogue systems, as well as the associated downstream harms.
- **Katherine Miller,** [“Generative search engines: Beware the facade of trustworthiness”](#)
This May 2023 Stanford University Human-Centered Artificial Intelligence blog post discusses the paper [“Evaluating verifiability in generative search engines,”](#) which quantifies the issue of provenance in LLMs.
- **Nikhil Sharma, Q. Vera Liao, and Ziang Xiao.** [“Generative echo chamber? Effects of LLM-powered search systems on diverse information seeking”](#)
This recent paper discusses how information seeking behaviors are changing as a result of generative search.
- **Omar Khattab, Christopher Potts, and Matei Zaharia.** [“A moderate proposal for radically better AI-powered web search”](#)

This July 2021 Stanford University Human-Centered Artificial Intelligence blog post highlights the issue of provenance in question-answer based web search and potential solutions.

About the Authors

- **Shahan Ali Memon** is a first-year doctoral student at the [University of Washington Information School](#).
 - **Jevin D. West** is an associate professor at the University of Washington Information School, and co-founder of the [Center for an Informed Public](#).
-

Acknowledgements

We would like to thank **Soham De**, **Nic Weber**, **Damian Hodel**, **Asad Memon**, **Apala Chaturvedi**, **Abdulaziz Yafai**, and **Hisham Abdus-Salam** for helpful discussions and feedback.

Image at top: Illustration generated by [You.com](#)'s Create mode.

References

- [1] **Bender, E. M., Gebru, T., McMillan-Major, A. & Shmitchell, S.** [On the dangers of stochastic parrots: Can language models be too big?](#) In Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency, 610–623 (2021).
- [2] **Lewis, P. et al.** [Retrieval-augmented generation for knowledge-intensive NLP tasks.](#) Adv. Neural Inf. Process. Syst. 33, 9459–9474 (2020).
- [3] **Huang, L. et al.** [A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions.](#) arXiv preprint arXiv:2311.05232 (2023).
- [4] **Zhou, K., Hwang, J. D., Ren, X. & Sap, M.** [Relying on the unreliable: The impact of language models' reluctance to express uncertainty.](#) arXiv preprint arXiv:2401.06730 (2024).
- [5] **Liu, N. F., Zhang, T. & Liang, P.** [Evaluating verifiability in generative search engines.](#) arXiv preprint arXiv:2304.09848 (2023).
- [6] **Khattab, O., Potts, C. & Zaharia, M.** [A moderate proposal for radically better AI-powered web search.](#) (2021).

- [7] Shah, C. & Bender, E. M. [Situating search](#). In Proceedings of the 2022 Conference on Human Information Interaction and Retrieval, 221–232 (2022).
- [8] Ibrahim, H., Liu, F., Zaki, Y. & Rahwan, T. [Google Scholar is manipulatable](#). arXiv preprint arXiv:2402.04607 (2024).
- [9] Sheng, E., Chang, K.-W., Natarajan, P. & Peng, N. [The woman worked as a babysitter: On biases in language generation](#). arXiv preprint arXiv:1909.01326 (2019).
- [10] Bolukbasi, T., Chang, K.-W., Zou, J. Y., Saligrama, V. & Kalai, A. T. [Man is to computer programmer as woman is to homemaker? Debiasing word embeddings](#). Adv. Neural Information Processing Systems 29 (2016).
- [11] Wan, Y. et al. [“Kelly is a warm person, Joseph is a role model”: Gender biases in LLM-generated reference letters](#). arXiv preprint arXiv:2310.09219 (2023).
- [12] Bai, X., Wang, A., Sucholutsky, I. & Griffiths, T. L. [Measuring implicit bias in explicitly unbiased large language models](#) (2024). 2402.04105.
- [13] Salinas, A., Penafiel, L., McCormack, R. & Morstatter, F. [“Im not racist but...”: Discovering bias in the internal knowledge of large language models](#). arXiv preprint arXiv:2310.08780 (2023).
- [14] Abid, A., Farooqi, M. & Zou, J. [Persistent anti-muslim bias in large language models](#). In Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society, 298–306 (2021).
- [15] Khorramrouz, A., Dutta, S., Dutta, A. & KhudaBukhsh, A. R. [Down the toxicity rabbit hole: Investigating palm 2 guardrails](#). arXiv preprint arXiv:2309.06415 (2023).
- [16] Prabhu, V. U. & Birhane, A. [Large image datasets: A pyrrhic win for computer vision?](#) arXiv preprint arXiv:2006.16923 (2020).
- [17] Noble, S. U. [Algorithms of oppression](#). In Algorithms of oppression (New York University Press, 2018).
- [18] Wu, Z., Sanderson, M., Cambazoglu, B. B., Croft, W. B. & Scholer, F. [Providing direct answers in search results: a study of user behavior](#). In Proceedings of the 29th ACM International Conference on Information & Knowledge Management, 1635–1644 (2020).
- [19] Bink, M., Zimmerman, S. & Elweiler, D. [Featured snippets and their influence on users’ credibility judgements](#). In Proceedings of the 2022 Conference on Human Information Interaction and Retrieval, 113–122 (2022).
- [20] Tonmoy, S. et al. [A comprehensive survey of hallucination mitigation techniques in large language models](#). arXiv preprint arXiv:2401.01313 (2024).

- [21] Zhang, H. et al. [R-tuning: Teaching large language models to refuse unknown questions](#). arXiv preprint arXiv:2311.09677 (2023).
- [22] Gallegos, I. O. et al. [Bias and fairness in large language models: A survey](#). arXiv preprint arXiv:2309.00770 (2023).
- [23] Sharma, N., Liao, Q. V. & Xiao, Z. [Generative echo chamber? effects of llm-powered search systems on diverse information seeking](#) (2024). 2402.05880

Search

Recent Posts

- Announcing our 2024 election journalist help desk
- What to expect when we’re electing: The 5 moves of misleading election rumors
- What to expect when we’re electing: An election rumoring timeline
- Fostering a more informed public in Washington through media literacy education
- High school students, local seniors in Sedro-Woolley demonstrate the power of intergenerational learning

Archives

- September 2024
- August 2024
- July 2024
- June 2024
- May 2024
- April 2024
- March 2024
- February 2024
- January 2024
- December 2023
- November 2023
- October 2023

September 2023

August 2023

July 2023

June 2023

May 2023

April 2023

March 2023

February 2023

January 2023

December 2022

November 2022

October 2022

September 2022

August 2022

July 2022

June 2022

May 2022

April 2022

March 2022

February 2022

January 2022

December 2021

November 2021

October 2021

September 2021

August 2021

July 2021

June 2021

May 2021

April 2021

March 2021

February 2021

January 2021

December 2020

November 2020

October 2020

September 2020

July 2020

June 2020

May 2020

April 2020

March 2020

February 2020

January 2020

December 2019

November 2019

October 2019

July 2019

Categories

2024 Rapid Research Blog

CIP Award for Excellence

Education

Events

MisinfoDay

[News](#)

[Newsletter](#)

[Rapid Research Blog](#)

[Recent Press](#)

[Research](#)

[Uncategorized](#)

[Updates](#)

[Join our Mailing List](#)

[Contact Us](#)

[Privacy Policy](#)

[Terms and Conditions](#)

[Facebook](#)

[LinkedIn](#)

[Mastodon](#)

© 2024 Center for an Informed Public at UW