

AI Critics Have an Ignorance Problem

Artificial intelligence desperately needs critics. It needs people willing to puncture the extravagant claims of technology companies, expose failures, investigate the consequences for workers and artists, and ask uncomfortable questions about privacy, surveillance, misinformation and the concentration of power. But criticism has one elementary requirement: the critic should understand, at least broadly, the thing being criticized.

Increasingly, however, that requirement seems to have been abandoned. Too much commentary about artificial intelligence is produced by people who appear embarrassingly unfamiliar with the technology they are confidently condemning. Worse, they often seem to have made remarkably little effort to inform themselves before publishing their judgments.

There are now at least two distinct forms of intellectual laziness surrounding AI. The first is criticizing the technology without bothering to learn what it can actually do. The second is using the technology without bothering to learn what it cannot reliably do. Both are becoming alarmingly common.

Consider a recent [article about Wispr Flow](#) by [Amy X. Wang](#), an editor and writer at *The New York Times Magazine*. Wang opens by confessing that she dictated the entire piece using the software. She then expresses, with apparent surprise, her disappointment that the resulting prose is not especially elegant. She finds herself wanting to move sentences, reorganize paragraphs, improve transitions and, as she puts it, “make [herself] sound smarter.”

Fair enough. Raw speech often makes for bad prose. But then comes the astonishing part. The article seems largely unaware that the very capabilities the writer wishes she had are already among the most commonplace uses of generative AI. If one wishes to polish their spoken

words and “sound smart,” all that one needs to do is to take the accurately transcribed text and ask an AI system to clean up the text, reorganize the argument, eliminate repetition, improve transitions, tighten the prose , preserve the author's voice, and so forth. These are not hypothetical capabilities awaiting some distant technological breakthrough. They are already here.

Astonishingly, the writer then expresses uncertainty about how AI is involved in the application under her scrutiny in the first place, seemingly unaware that one of the first material breakthroughs in AI, and the beginning of its mainstreaming, took place in the field of speech recognition with the advent of Siri in 2011 and the Amazon Echo in 2015.

But even more serious is her apparent ignorance that another form of AI would be needed to do what she wants, and that this technology falls under the broad field of Natural Language Processing. To complain that a transcription tool does not automatically produce a polished literary essay betrays such a basic misunderstanding of the technology she is examining that one is left to ask: Did anyone at *The Times* ask the obvious questions? Where were the editors?

A journalist can misunderstand a complicated technology; that is human. But journalism is not supposed to depend upon one person's understanding alone. Editors exist partly to challenge assumptions, expose gaps and ask the question every writer should fear: *Are we sure this is true?* A single editor asking, *Have you tried putting your dictated text into a generative AI system and asking it to revise it?*, might have produced a very different article.

The second form of AI laziness is the mirror image: people who use AI without understanding its limitations and then fail to exercise the human judgment they insist remains indispensable. Steven Rosenbaum's [*The Future of Truth*](#) provides an almost too perfect [example](#). The book warns about the dangers AI poses to truth, trust and our epistemological foundations. Yet readers discovered fabricated citations, invented quotations and other AI generated errors in its pages.

The irony is magnificent, but the failure is mundane. Nobody is forced to accept an AI generated citation. A machine can produce one. It cannot force an author to publish it. Somewhere between the generation of the false information and the printing of the book, a human being was supposed to verify it. That did not happen.

The issue was not merely that AI made an error. AI does that and we should all be well aware of this basic fact of life. The issue was that the author had written a book about truth and then failed to perform the elementary act of checking whether the claims and citations bearing his own name were true.

The same pattern appeared when Jeff Hancock, a Stanford researcher and expert on misinformation, submitted a declaration in support of Minnesota's law addressing [AI generated deepfakes](#). The filing included citations that did not check out, and Hancock later acknowledged using ChatGPT in his research process. Again, the lesson is not that nobody should use AI. It is precisely the opposite: if you use a tool, understand the conditions under which it can fail.

And then there is the institutional version of the same problem. A major Australian report used in the debate over restrictions on children's use of social media came under scrutiny for apparently [fabricated or incorrect references and citations](#). Questions were raised about the use of AI in its production. Whether the immediate error originated with an AI system or a human being is almost beside the point. The report passed through an entire chain of human responsibility: researchers, authors, reviewers, editors and institutions. Where were they?

This is the question that increasingly hangs over our arguments about AI. Journalists, writers and editors understandably worry that artificial intelligence may replace parts of their professions. They should worry. But the case for the irreplaceable superiority of human labor becomes considerably weaker when humans fail to perform the very functions they claim machines cannot perform: judgment, context, verification, skepticism and responsibility.

Nobody is suggesting that AI should replace journalists because a journalist misunderstands AI. Nor should editors be replaced because an editor misses a factual problem. But neither profession is entitled to immunity from comparison simply because its members are human.

The answer to AI is not human mediocrity demanding special protection. It is better human work. The best journalists should learn enough about the subjects they cover to criticize them intelligently. The best editors should ask the questions their writers forgot to ask. And anyone using AI should understand that the existence of an impressive machine does not absolve the human being using it of responsibility for the final product.

AI will make mistakes. That much is certain. But humans are increasingly making a different mistake: demonstrating, through ignorance and laziness, that they are willing to surrender precisely the qualities they claim make them indispensable. The problem is not that machines sometimes fail to exercise judgment. It is that, increasingly, the humans criticizing them do not either.