

Why AI slop is bad and people who post AI slop should feel bad

EUGYPIUS

JUL 24

[Substack has partnered with Pangram to enable readers to scan posts, notes and comments for machine-generated text.](#) No AI detector is flawless, but Pangram is pretty much the gold standard. Their algorithms detect AI-generated text more than 95% of the time, with a false positive rate approaching zero. I was very happy to read about this feature, and I think more social media platforms should do something like this before their content descends into a morass of platitudinous computer-generated nonsense.

[Subscribed](#)

1. A lot of things are true of AI chatbots at the same time: They are enormously useful and powerful tools. They can melt your brain if you're not careful with them. They have changed fundamentally, and forever, how people work in a wide variety of professions, from journalism to medicine to law. They have filled social media with low-effort engagement bait, spam and reprehensible slop of all kinds. I don't think anybody has a problem with the background use of AI, for research purposes, for error-checking, and the like. AI as a content production tool – in this case for prose composition – is where the problems begin.

2. Some argue that the transparent use of AI can serve essentially salutary purposes, perhaps by increasing the volume and quality of the very content that our best writers already produce. This is a naive view and it is totally at odds with how AI works in practice. Firstly, people overwhelmingly want to read real human content on platforms like Substack; they do not come here to converse with human-ventriloquised robots because that is not what social media is for. And secondly, it's

just obvious that the vast majority of those who use AI to write give off the strong odour of the disreputable and the insincere.

3. Despite its considerable powers, AI will probably never be a great prose stylist, and this for much the same reason that those annoying Gmail assistive writing functions will never steer your correspondence into any kind of profundity. AI chatbots produce expected, average prose with some additional tics that they owe to their preference training. In short, their style is bland and ordinary at the level of content, while being punchier and more direct than the human average at the level of style. Much AI content thus reads kind of like ad copy, which is mostly fine if want to write ads but which becomes pretty bizarre if not obnoxious in other contexts.

4. At a deeper level, LLMs do not have a fixed perspective or any kind of self-identity. They are always trying to play a certain role, to lend whatever impression the user demands of them, and in this way whatever meaning they produce is subordinated to the project of seeming. The more tightly you constrain the AI with specific demands, the more useful it becomes, but the open-ended prompts necessary to get LLMs to compose whole essays often produce very strange content. Superficially these pieces may seem perfectly fine – exactly the analytic or the scientific or the philosophical content you wanted – but upon closer examination you begin to notice that they are crafted merely to reproduce the aesthetic of those things.

5. With careful prompting and even more careful editing you can reduce or eliminate many of these problems, but this is harder than you'd think and capable authors will find it much easier simply to write their own goddamn blogposts. Thus AI has become the preferred composition tool not of talented essaysists, but of lazy writers, subliterate people and scammers of all kinds. It proliferates in a variety of social media subgenres that have long been suspect for other reasons, among them the cryptosphere, self-help posters, life philosophers, deepity aphorists and ragebait political hacks. Wherever the naive or inexperienced proliferate and the game is to appeal to the lowest common denominator, AI dominates.

6. AI makes up for its stylistic and communicative shortcomings by excelling in the production of mindcrack – shallow content that attracts eyeballs primarily by being gratuitously loud, outrageous, superficially profound, controversial, enraging or

disturbing. Above all, AI chatbots are good at telling people what they want to hear, and if you know something about what people want to hear and you're willing to spend a few hours optimising your prompts, you too can produce mass quantities of mindcrack for your 30 Euro/month Claude subscription. The ensuing enshittification becomes self-reinforcing as discerning writers and readers take their attention to other platforms. This phenomenon is presently strangling X.

7. For all these reasons, social media platforms should discourage AI text algorithmically and clearly label obvious AI content. Ideally, Substack would slap neon LLM tags on comments, notes and posts flagged as having been written with AI assistance. For now they've chosen the gentler path, of reader-initiated AI scanning.

8. Now finally, here at the bottom, my own transparency statement: I have never posted, and I will never post, AI content as my own, but that does not mean I don't use AI. I use it quite a lot and for a range of tasks. From the very beginning, I have used the AI translation tool DeepL to translate German and other European languages for my Anglophone audience. This is standard practice, and I always review and edit these machine translations to avoid infelicities and errors. Like most other people, I also use AI (primarily Claude and Grok) as elaborate, sophisticated search engines. They can pull together articles and other sources going back many years; things that once took me multiple hours or days to find, I can now unearth in a matter of minutes. More rarely, I use AI to summarise books and scientific papers, particularly when I'm not sure of their relevance to whatever it is I'm working on. (Thus far I've never published anything based on AI summaries.) A while ago I experimented with using AI to identify weaknesses in my own argumentation, but I gave up on that when I realised the practice was pushing me towards bloodless content that felt workshopped and boring. Instead, I will occasionally use the most sophisticated AI models to stress-test my own theses before I begin to write. Now and again I've tried to get Grok to produce punchy post titles for me, but I've never used the results because they've been terrible.