

After Fluency

I. Instance

I have been working on a series of essays, including this one, that did not begin with an outline on a blank page. Each one started with questions I could not answer cleanly, worked through in conversations that wandered and doubled back before they found their shape. My interlocutor in those conversations was an AI. Which raises an obvious question: did I write this?

It is the obvious question, and probably the wrong one — or at least the small version of a much larger one. Attribution matters, for reasons that are real and worth taking seriously. But the more interesting question is not who produced the sentences. It is what writing has always been doing, and what we have been assuming it means — assumptions that turn out to be worth examining carefully. This essay is an attempt at that examination.

Writing has always done more than record thought. It also has been both a signal of thinking and a means of developing it. The immediate problem is epistemic: AI-assisted writing removes the proxy we rely on to help recognize sound thinking. The longer-term problem is developmental: it undermines the process we depend on to build the capacity to think in the first place. This essay is also, unavoidably, subject to the same scrutiny it describes.

II. Indicator

Expository writing is commonly treated as the record of thought, the place where ideas, once formed, get set down. That description is incomplete in a way that matters. Writing does not simply record thinking; it builds thinking. The effort of finding words for a half-formed idea, of arranging claims in sequence, of discovering a paragraph later that the premise was wrong — this is how thinking develops.

In expository work, where language serves the idea rather than constituting it, the result of that effort, when it succeeds, is fluency: prose that moves a reader cleanly from one

idea to the next. Because that fluency has historically required effort, when the writing effort is approached honestly, it has served as a decent indicator of the thinking behind it.

The usefulness of fluency as an indicator arises from this coupling between writing and thinking: doing one has meant doing something that looked like the other. It has never been perfect, but it has been reliable enough that institutions, educators, and readers built their practices around it. With the emergence of generative AI, that reliability is eroding.

III. Precedent

The decoupling of authorship from mechanical production is not new. Software development has been living with it for decades. But the comparison is imperfect in a way that matters.

As successive layers of abstraction accumulated — compilers, high-level languages, frameworks, libraries — lower-level work that had previously demanded real skill and attention was automated. That was genuinely productive. Effort shifted upward, away from the mechanical and toward the architectural: defining problems, structuring systems, deciding what the program should actually do. The transition was gradual enough that professional norms and educational practices could absorb it.

Crucially, software did not mistake this shift for evidence that surface success reflected underlying soundness. Quite the opposite: the field came to understand that execution and correctness were fundamentally decoupled. Code that compiles and runs has always been recognized as a minimal condition, not evidence that the right problem is being solved or that it is being solved well.

That recognition drove the development of external systems of validation: testing, code review, formal verification — practices designed explicitly to evaluate whether the underlying judgment holds up. As abstraction increased, these practices became more, not less, important.

Writing developed differently. Because the effort required to produce fluent prose was tied to the effort of thinking, fluency functioned as a workable internal proxy for judgment. Equivalent external validation systems were never widely built.

Generative AI introduces a different kind of break. In writing, it severs the connection between fluency and the thinking behind it. In software, it introduces a similar risk — but this time not as a gradual shift absorbed by stronger validation practices, but as a temptation to ignore what those practices were built to address.

The emerging practice of “vibe coding” — accepting AI-generated code because it appears to work — reflects a regression from that earlier understanding. Where the field had explicitly recognized that surface success is decoupled from soundness, it now risks behaving as though that decoupling does not exist.

Writing is entering this condition for the first time. Software, which had already adapted to it, may now be partially unlearning that adaptation. In both cases, surface success is being treated as though it reflects underlying soundness. The difference is that writing never built a safety net, and software may be loosening its own.

IV. Qualification

The claim that AI has made fluency cheap requires qualification. Generative systems are highly effective at producing coherent, organized prose in contexts where language functions as a vehicle: carrying information, summarizing arguments, explaining processes. In those contexts — which make up a large share of real-world writing — the cost of producing serviceable prose has dropped substantially. The indicator that fluency once provided is already eroding.

That is not the same as saying AI-generated prose is currently indistinguishable from expertly crafted writing. At the moment it often isn't. Careful readers notice the cosmetic tells: the em-dash deployed as a default connector, bullets where prose would serve better, the rule of three that makes thin analysis look comprehensive, the binary reversal that substitutes rhetorical snap for actual argument (not a flaw, a feature; not the end, the beginning). These are irritants, the kind of habits a good editor would flag

regardless of their origin. The more consequential tells are both more subtle and more substantive. Conclusions stated before they have been argued. Open questions treated as settled. A confidence unmarked by the uncertainty that genuine engagement with hard questions produces. They are the fingerprints of a process in which no one was weighing the claims.

A sufficiently trained model will learn to avoid the cosmetic tells, whether by introducing noise to break up the probabilistic regularities that produce them or by other means we haven't yet seen. But detecting AI-generated prose isn't the most important goal anyway. A piece of AI-assisted writing that is carefully judged, honestly argued, and grounded in evidence is more valuable to a reader than poorly reasoned prose produced entirely by human hand.

V. Plausibility

The loss of fluency as a reliable signal is one consequence of generative AI. A separate one operates on a different axis: the cost of producing fluent prose has collapsed, precisely in the professional writing, analysis, and documentation contexts where fluency has mattered most as a signal. The more significant consequence is not confusion about authorship, but plausibility: content that sounds authoritative without being sound.

Fluency has always been an imperfect indicator. The history of human communication includes propaganda, pseudoscience, motivated argument, and eloquent nonsense. All have always required effort, and all have always been capable of persuading readers who lacked the tools or context to evaluate them carefully. The difference now is that the cost of producing such content has dropped to nearly nothing.

Manufacturing convincing misinformation previously required significant effort because faking domain familiarity has been a costly challenge in its own right. That cost limited how much misleading content could circulate. Generative AI removes that constraint. Now anyone can create plausible-sounding text on any subject at negligible cost, while the effort required to produce well-grounded content has not dropped correspondingly.

What is shifting is not just the volume of deliberately fabricated misinformation, but the larger and less visible volume of content generated without anyone weighing the claims, careless rather than malicious but no less misleading for that.

And there is a natural governor on that second kind that generative AI is removing: the effort of writing carefully has always been a partial corrective, forcing some gaps in reasoning into view before they become confident claims. What AI-assisted writing removes is not the ability to recognize the presence of thinking, but the ability to detect its absence — fluency that once indicated someone was weighing the claims now indicates nothing of the kind.

For non-experts (which is most of us, most of the time) the traditional tools for navigating this have always been less about evaluating content directly and more about evaluating its sources. Institutional reputation, demonstrated track record, credentials earned over time — these are the types of indicators non-experts have always used to decide whose fluency to trust, and they remain the most rational tools available. A paper from a research institution with decades of reliable work is still more trustworthy than an anonymous document, regardless of how convincing either sounds. The indicators have not stopped working, at least not yet.

VI. Development

The developmental consequence is slower and less visible because it operates on the capacity that writing practice builds over time instead of on individual acts of writing. Writing is not only a signal of thinking; it is one of the primary means by which the capacity for thinking develops. The struggle toward fluency — constructing arguments, maintaining logical dependency, making reasoning precise enough that gaps become visible — is how analytical thinking is built.

Not all difficulty is worth preserving. Some friction impedes without developing. But some difficulty is generative, and removing it improves apparent output while leaving the underlying capacity undeveloped — the answers look right, but the thinking that should have produced them never happened.

Whether AI assistance removes generative friction or merely tedious friction depends on the kind of writing and the stage of the learner. Someone who depends on AI assistance while still developing the capacity to construct an argument risks bypassing the struggle that builds that capacity. Someone who has already developed that capacity may find the same assistance simply gets the work done faster. Blanket prohibition misses the distinction, but so does blanket permission.

VII. Ratio

As AI-assisted writing becomes more common, the individual consequences of bypassed generative friction become a collective one. Fewer people doing the hard work of constructing arguments means more fluent content produced without the underlying thinking that fluency once required — and at a volume and rate no previous era has had to contend with. Sound judgment has not become more rare, but it is rapidly being overwhelmed. As the ratio of well-considered content to content that merely appears so shifts, assessing which is which becomes the harder and more consequential problem.

There is a genuine upside to AI's ability to generate fluent prose on demand. Facility with prose has always been unequally distributed, and the gap between analytical capacity and that facility has always left some well-considered thinking unexpressed or poorly served by its packaging. People with the ability to reason but not the ability to express themselves now have a vehicle they didn't before. The same mechanism generating fluent content of questionable validity is also, for some, finally making sound judgment expressible.

But the asymmetry is real and should not be minimized. The barrier to producing fluent nonsense has dropped to nearly zero for everyone. The barrier to producing well-considered content has dropped only for people who already had the capacity in the first place. The junk will outpace the well-judged by a wide margin, and probably already does. In many workplaces, first drafts of reports, summaries, and analyses are now routinely generated with AI assistance. The volume of fluent, serviceable text has increased dramatically, but the time spent evaluating whether those drafts are actually

sound has not increased in parallel. The result is not just more writing, but more writing whose apparent coherence exceeds the judgment applied to it.

Which means the proxy that readers have always used to identify sound writing, that is, fluency as an indicator of judgment, is failing at exactly the moment when the need for direct assessment is becoming most acute. If, as many expect it will, AI-assisted writing becomes the norm, fluency will increasingly fail as a signal of judgment. Both reader and author must engage more directly with judgment itself: the reader to find it, the author to exercise and expose it. Authorship, under these conditions, shifts from the production of language to the evaluation of the argument, from generating fluency to deciding what survives the creative process. The questions that define authorship become: Is this claim defensible or merely coherent-sounding? Does this distinction survive critical challenge? Is this argument structured around evidence or arranged around a conclusion already reached? They have always been the questions serious writers and thinkers asked. What changes is who can avoid asking them while still producing something that appears, on the surface, as though they did.

This reframing of authorship is less novel than it might appear. The distinction between writing and authorship — between producing sentences and being accountable for what they claim — has existed in practice long before AI made it urgent. Technical documents at engineering and research organizations are routinely produced through a division of labor: subject matter experts supply the judgment, technical writers shape it into prose. The document belongs to neither alone, but accountability for what it asserts belongs to the expert, not the writer.

Academic authorship formalizes a related principle: to be named as an author is not merely to have contributed sentences but to have participated substantively in the intellectual work, approved the final version, and agreed to answer for the claims, their accuracy, their integrity, their defensibility under scrutiny. The academic author is also obligated to cite sources, both to give credit where it is due and to show the foundation the argument rests on. Even here the model is imperfect: academic papers routinely carry multiple authors, sometimes many, and whose judgment the paper really represents is often less clear than the byline suggests. These questions — of credit, of

accountability, of what authorship actually means — have always existed. What AI-assisted writing does is force them into every domain at once.

In AI-assisted work, the model generates the sentences. Whether there is an author depends on what the human is doing. If the process is primarily one of prompting and accepting with little evaluative pressure applied to what comes back, the judgment that defines authorship is never really exercised at all. Authorship resides with the one making judgments about what the work should contain, what claims it should make, what survives the creative process.

If no human is exercising that judgment, the work has no author.

Responsible authors have always had the obligation to construct a sound argument. The advent of AI-assisted writing introduces a new obligation: to make that argument verifiable, not merely fluent. The reader has always had reason to look beyond fluency, but the grounding that makes verification possible is becoming harder to ignore. The academic world requires this: citations, methodology, peer review, the apparatus that allows a reader to check the work rather than merely accept it. What is changing is who should consider meeting that standard. The apparatus of academic argument was never designed for general audiences, and how that translates to expository writing more broadly is an open question. As it turns out, AI has also made the age-old obligation harder to meet: constructing a sound argument in the first place is more difficult when the tools themselves work against it.

The same probabilistic mechanisms that make AI fluent also make it capable of generating confident, plausible-sounding claims that are factually wrong. The errors range from subtle to wholesale: evidence that doesn't appear in the cited literature, legitimate-sounding citations referencing sources that don't exist, facts asserted with no basis in reality and no attempt to justify them. The good news, such as it is, is that these errors are detectable in principle: assertions can be checked, sources verified, citations confirmed, but only if the author is willing to do the checking.

A less tractable problem is sycophancy, or “Glazing.” Conversational AI interfaces exhibit a bias toward maintaining engagement that produces a tendency to validate

rather than challenge. The executive developing a business strategy is at risk of being fooled by an interlocutor that tends to confirm that the argument is working, the distinctions are sharp, the structure holds — whether or not that is true. The bias is structural, and not always easy to detect. The careful author working with AI has to know, going in, that the conversation is not a neutral one, and must actively resist accepting its output at face value.

VIII. Scale

When fluency fails as a signal, readers fall back more heavily on external indicators. The proliferation of AI-generated content is arriving simultaneously with coordinated attacks on the institutions that have historically served as the most accessible shortcuts to validity assessment. Scientific agencies, professional bodies, established media — the organizations whose track records made them useful proxies for non-experts trying to evaluate complex claims — face challenges to their credibility coming from multiple directions at once. In some cases the challenge is legitimate recalibration, institutions that have earned skepticism through their own failures. In others it is coordinated delegitimization by bad faith actors working from the outside. And in others still it is the active removal of the expertise that gave the institution its credibility to begin with — which may be the most corrosive form, because the reputation outlasts the substance it was built on, and the hollowing out is not immediately visible from the outside. Distinguishing among those three is itself a judgment call requiring exactly the evaluative capacity that is hardest to develop and easiest to bypass. The timing is unfortunate. Plausibility abundance and institutional stress are arriving together, and their interaction is considerably more destabilizing than either would be alone.

What this suggests is not despair but adjustment — a recalibration of what reading carefully actually means. It has always meant more than parsing sentences: asking where a claim comes from, what would have to be true for it to be wrong, whether the source has been reliable before, whether the argument's structure serves its evidence or merely decorates a conclusion reached in advance. These questions were always the right ones. They are now more necessary, and less optional, than they have ever been.

IX. Reckoning

These changes return us to the original question.

The prose you are reading is fluent. It is organized. It makes what may sound like reasonable claims in a confident register. None of that is evidence of its validity. The judgment that would make it more than plausible — the decisions about what to assert and what to hedge, what distinctions are real and what are merely convenient, whether the grounding holds up — that judgment is not visible in the sentences themselves. It has to be evaluated by the reader, using exactly the tools the essay has been describing: not fluency assessment, but validity assessment. The essay is not exempt from the problem it describes. No fluent text is.

Not long before writing this, I had used AI to help build analytical tools that themselves leveraged AI, then used them in a consulting engagement with the advanced research computing office at a large research university — exploring how the office should evolve its research computing environment, what infrastructure to prioritize, where AI itself fit into that picture. AI all the way down — and the vertigo that comes with that recognition is real. But the questions driving that work were not questions the AI was answering. They were questions I was asking, shaped by enough experience to know which answers to challenge and which to trust. The technology accelerated the thinking. It did not supply the direction.

That is, as best I can tell, what happened here as well — making this essay its own recursive instance of the condition it describes. The conversations were long and winding. Claims got pushed back on, formulations that didn't hold got replaced, distinctions got sharpened, the sequencing changed when the logic demanded it. With all of that back and forth over framing, word choice, and structure, I can safely say I contributed to the writing. But by the terms the essay has developed, I am its author. The thinking was mine. The judgment was mine. The accountability is mine. If the argument fails, that failure is mine to answer for. Whether the thinking is sound is a question that can be answered best by readers who bring their own judgment to bear on it.

What diminishes with AI assistance is not the need for thinking, for judgment, or for accountability — but a proxy we relied on to recognize their presence.