

Essential or Familiar? Generative AI in Writing, Journalism, and Schools

A way to tell which of our institutions' rules must be protected, and which are merely habits worth rethinking, as generative AI enters writing, news, and education

THE BOTTOM LINE Generative AI is prompting a wave of new rules across newsrooms, schools, and workplaces, and some are essential, some merely defensive. The useful question is not “should we allow AI?” but a sharper one: which of our institutions’ principles must remain unchanged, and which of our familiar mechanisms can adapt? Peer review, bylines, original student writing, and fact-checking are mechanisms. Accountability, transparency, and rigor are the principles they serve. When a newspaper printed an AI-generated reading list of books that do not exist, the failure was not that AI wrote it; no human was accountable for verifying it. Distinguishing principle from mechanism is how Virginia’s schools, agencies, and employers can write GenAI policy that protects what matters without banning what helps.

The Distinction That Cuts Through the Noise

Scholarly and civic institutions exist to serve a few durable **principles** — accountability, transparency, and rigor — that they pursue through familiar **mechanisms** such as peer review, authorship, fact-checking, and original student work. The principles are the ends; the mechanisms are the means we happen to use, designed for an era of human labor. Generative AI does not break the principles. It forces a question we have not had to ask before: which of our familiar mechanisms are *essential* to a principle, and which are merely *familiar*?

Two terms carry the argument, and only two. A **principle** is what an institution is trying to protect — accountability, transparency, rigor. A **mechanism** is a concrete practice adopted to protect it — a byline, a handwritten essay, a peer-review form, a reporter’s notebook. Principles are set by professions and institutions; mechanisms are enforced by specific, accountable people. Generative AI rarely threatens a principle; it destabilizes the familiar mechanism, and often by quietly removing the person who used to enforce it. Policy that confuses the two either bans useful tools (protecting the mechanism) or lets real harms through (losing the principle).¹

Writing: Provenance Is Not the Same as Value

Anxiety about AI-written text often conflates two distinct questions. One concerns *value* — is the reasoning sound, the evidence real, the analysis honest? The other concerns *provenance* — who or what produced the words? These can come apart. A well-reasoned, honestly sourced document is not worthless because AI drafted the prose; a fluent, confident document is not trustworthy because a human typed it.

The essential principle is accountability: someone must be answerable for whether the claims are true. The familiar mechanism, “a human wrote every word,” was only ever a proxy for that accountability. When AI language tools help a non-native English speaker or a time-pressed expert communicate clearly, the mechanism changes, but the principle remains intact, provided the human remains answerable for the content. The failure mode is not AI assistance; it is AI assistance with no one accountable for verifying the result.

Journalism: A Cautionary Case

In May 2025, the Chicago Sun-Times and the Philadelphia Inquirer published a syndicated “Summer reading list for 2025.” Ten of the fifteen recommended books did not exist — real authors (Isabel Allende, Andy Weir, Percival Everett) were paired with plausible, entirely invented titles and plot summaries generated by AI and printed without verification. Both papers confirmed that the content violated their own policies and had bypassed the newsroom through a syndication partner.²

Viewed through the principle-vs-mechanism lens, the lesson is clear. The failure was *not* that AI wrote the list. The failure was that the mechanism that normally guarantees accountability (a human editor who verifies before print) had been removed from the chain, and no rule required anyone to check. The principle (nothing goes to readers

unverified) was sound; the mechanism (assume the syndicated insert was human-checked) had quietly stopped serving it. A policy response that says “ban AI in the newsroom” protects the old mechanism. A better one says “every published claim has a named human accountable for its accuracy, whatever tool produced the draft.” That protects the principle.³

The stakes rise quickly beyond a reading list. In *Mata v. Avianca* (2023), a lawyer filed a federal brief citing six court decisions that did not exist; a chatbot had invented them, and no one verified the citations before filing. The court imposed sanctions under Rule 11.⁴ Courts have since sanctioned a growing series of similar filings, and a peer-reviewed study found that even dedicated legal-research AI tools fabricate citations a substantial fraction of the time.⁵ The pattern is identical to the newspaper case, only costlier: the principle (a filing’s authorities must be real and verified) is untouched, but the mechanism that enforced it (a lawyer who reads the cases before signing) was removed from the chain. The same structure governs higher-stakes domains still — a diagnosis or a treatment plan drafted with AI assistance is safe only if a clinician remains accountable for verifying it. In every case the question is not whether AI was used, but whether a named human still stands answerable for the result.

Schools: What Are We Actually Protecting?

The classroom debate over generative AI is often framed as detection: can we catch students using it? But detection is a mechanism, and a failing one. In the most-cited study of these tools, seven widely-used AI-writing detectors misclassified more than half of essays by non-native English speakers as machine-generated (a 61% average false-positive rate) while judging U.S.-student essays with near-perfect accuracy — penalizing exactly the students least able to withstand a false accusation.⁶ Building school policy on detection protects a familiar mechanism while creating a new inequity.

The essential question is what an assignment is *for*. If the goal is to certify that a student can produce polished prose unaided, AI disrupts that mechanism. But if the goal is the underlying principle — that the student can reason, structure an argument, and understand the material, the task can be redesigned so the learning is visible in ways AI cannot silently supply: in-class work, oral defense, drafts-and-revision, or assignments that require the student to critique and correct AI output.⁷ The principle (students must learn) is non-negotiable. The mechanism (the take-home essay as proof of learning) is a habit that can adapt. Schools that conflate the two will spend their energy policing tools instead of protecting learning.

A Test Virginia Institutions Can Apply

For any proposed GenAI rule, whether in a newsroom, an agency, a school, or a contractor’s workflow, three questions separate essential from familiar:

- **What principle does this rule protect?** Name it in one word: accountability, transparency, rigor, fairness, or learning. If you cannot name it, the rule is probably protecting a habit.
- **Is the thing we are anxious about a threat to the principle or only to the mechanism we have used to serve it?** If the principle survives when the mechanism changes, the rule can adapt.
- **Does the rule keep a human accountable, or does it simply ban a tool?** Banning a tool is easy and often beside the point; preserving a named, answerable human is harder and usually the real safeguard.

Virginia in Focus

Schools and universities

Virginia’s K–12 divisions and public universities are now developing GenAI policies. The ones that will age well are built on redesigned assessment and clear accountability, not on detection software or blanket bans. The Commonwealth has an opportunity to model “protect the principle, adapt the mechanism” guidance rather than leaving each division to reinvent or misfire on its own rules.

State agencies and public communication

As Commonwealth agencies adopt GenAI for drafting public-facing materials (notices, summaries, constituent correspondence), the Sun-Times lesson applies directly: the risk is not the tool but an unverified pipeline. Agencies

should require a named human to be accountable for the accuracy of anything published, regardless of how it was drafted.

Newsrooms and syndication

The 2025 reading-list incident reached Virginia readers through the same syndication channels that serve papers nationwide. The exposure is structural: content can enter a trusted outlet without that outlet’s own verification. Transparency about provenance and a clear, accountable party at each handoff are the safeguards.

What Responsible GenAI Policy Requires

Principle	What it means in practice
Name a human, not a tool	Every published or submitted product has a named person accountable for its accuracy and integrity, whatever tool produced the draft. A rule that bans a tool but names no one protects a mechanism, not a principle.
Disclose provenance	Transparency about when and how AI was used lets readers, reviewers, and teachers calibrate trust and preserves the accountability that a byline used to carry implicitly.
Redesign, don’t just detect	Where a mechanism (the take-home essay, the assumption of human drafting) no longer guarantees the principle, redesign the task to make the principle visible again, rather than policing the tool.

What to Watch

- Whether Virginia school divisions and universities converge on assessment-redesign approaches or default to detection tools that disadvantage non-native English speakers.
- Whether newsrooms and their syndication partners adopt provenance-and-accountability standards after the 2025 fake-book-list incidents or treat them as one-off embarrassments.
- Whether Commonwealth agencies require a named accountable human for AI-assisted public communication before a visible error, rather than after.
- Whether “AI disclosure” norms in scholarly and professional writing stabilize around honest, specific statements rather than concealment or blanket prohibition.
- Whether the essential-vs-familiar distinction is adopted in policy language, protecting principles while allowing mechanisms to adapt, or whether debate remains stuck on all-or-nothing tool bans.

Two harder questions sit just beyond this brief and deserve their own treatment: once a named human is accountable, **who enforces that accountability, and who bears the cost when verification fails** — questions of professional sanction, legal liability, and whether organizations will insure against AI-caused harms. A future STEP Brief will take these up. The point here is prior to them: liability and enforcement only become tractable once the principle and the accountable party are named clearly, which is what this framework is for.

Notes

1. The distinction between a practice’s essential function and its familiar form draws on a long tradition in the philosophy of technology — e.g., L. Winner, “Do Artifacts Have Politics?” *Daedalus* 109(1), 1980; P.-P. Verbeek, *Moralizing Technology*, Univ. of Chicago Press, 2011; and A. Feenberg on the difference between a technology’s essence and its contingent social form. The application of that distinction to editorial, journalistic, and educational policy — the “essential vs. familiar / principles vs. mechanisms” lens used throughout this brief — is the authors’ own synthesis.
2. Chicago Public Media / Chicago Sun-Times, public statement on the May 18, 2025 “Heat Index” special section (“recommended books that do not exist”), May 20, 2025. The syndicated section, licensed from King Features (a unit of Hearst), also appeared in the *Philadelphia Inquirer*. Reporting: NPR, May 20, 2025; Snopes, May 21, 2025.
3. M. Bell (CEO, Chicago Public Media), “Lessons — and an apology — from the Sun-Times CEO on that AI-generated book list,” *Chicago Sun-Times*, May 29, 2025. Notably, the newsroom’s own corrective independently arrives at the principle argued here: the revised policy requires third-party licensed content to state its source, not be presented as newsroom-created, and be reviewed by newsroom editors — that is, a named accountable human in the chain, not a ban on the tool.

4. *Mata v. Avianca, Inc.*, 678 F. Supp. 3d 443 (S.D.N.Y. 2023) (Castel, J.). Attorneys were sanctioned under Rule 11 for a brief citing six nonexistent, ChatGPT-generated decisions; a \$5,000 penalty issued June 22, 2023.
5. Courts have since sanctioned a growing series of AI-fabricated filings; an independent tracker maintained by researcher Damien Charlotin catalogs them. On tool reliability, see the Stanford RegLab study finding that dedicated legal-research AI tools still produce hallucinated or erroneous citations a substantial fraction of the time (Magesh et al., “Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools,” 2024).
6. W. Liang, M. Yuksekgonul, Y. Mao, E. Wu, and J. Zou, “GPT detectors are biased against non-native English writers,” *Patterns* 4(7): 100779, 2023. DOI: 10.1016/j.patter.2023.100779.
7. The redesign strategies noted here — oral defense, in-class work, drafts-and-revision, critique-the-AI assignments — are not new pedagogy; they are long-standing ways to make learning visible that GenAI simply makes urgent again. This is itself an instance of the brief’s thesis: the principle (visible learning) endures while the familiar mechanism (the take-home essay) is the part that must adapt.
8. Related STEP work: L. Medsker and S. Virmani, “The Accountability Vacuum: Agentic AI in High-Stakes Domains,” *AI Matters* (ACM SIGAI), 2026 (DOI: 10.1145/3795125.3795132); and STEP Briefs 1 (AI & Work) and 2 (Agentic AI).

About this brief • STEP Briefs are living documents maintained by the STEP Collaborative at George Mason University’s Center for Assurance Research and Engineering (CARE) that summarize fast-moving science and technology policy issues for Virginia policymakers and citizens. This brief develops an “essential vs. familiar / principles vs. mechanisms” framework for generative-AI policy, illustrated with the documented 2025 *Chicago Sun-Times* / *Philadelphia Inquirer* AI reading-list incident and published research on AI-writing-detector bias.⁸ Next scheduled review: November 2026.

Feedback or corrections: STEPair@gmu.edu • care.gmu.edu/step • TRACE Repository: care.gmu.edu/step-agentic-ai-repository