

Human Creativity and the Integration of Artificial Intelligence into the Education System

Siddharth Roy
siddharth.rroy@outlook.com

ABSTRACT

This paper examines the neurological basis of human creativity, the pattern-based limitations of large language models, and the cognitive risks that frequent AI reliance may pose to critical thinking and long-term retention. It investigates how the uneven distribution of AI tools may reinforce educational inequities, how student data privacy remains insufficiently protected, and how AI detection tools carry significant statistical limitations. This paper is a narrative review: it synthesizes existing research rather than reporting original experimental data. Section 4 explains how sources were selected and what the scope of the analysis is. While artificial intelligence is gradually transforming education, its irresponsible and unchecked integration appears to conflict with essential academic integrity policies, legal guidelines, and the foundational principles of authentic learning.

KEYWORDS

artificial intelligence, large language models, academic integrity, neuroplasticity, cognitive atrophy, educational equity, data privacy, AI detection, generative AI, critical thinking, learning retention, educational policy, cognitive load theory, desirable difficulties, retrieval practice, testing effect, generation effect, fluency illusion, metacognitive miscalibration, schema construction, cognitive offloading, working memory

INTRODUCTION

In the early 2020s, artificial intelligence platforms like ChatGPT and Gemini changed how students and teachers interact with technology in the classroom. Since 2022, generative AI platforms including ChatGPT, Gemini, Claude, Perplexity, DeepSeek, and Grok have become widely accessible to students. Each one is trained on large text corpora and produces fluent written responses on demand. These platforms differ in their underlying models and interfaces, but they all share the core mechanism of next-token prediction described in Section 6.2. The cognitive effects discussed in this paper apply to every one of them. The evidence, however, is strongest for ChatGPT, since it is the platform the cited studies actually tested. Within weeks of their release, students were using these tools to generate essays, solve complex problems, and summarize articles.

Aug 2026
Vol 10, No 5.

Educational policies at the time enforced strict rules against plagiarism and academic dishonesty, specifically cheating off other students. No policies accounted for AI. School districts across the United States responded in different ways. Some banned AI entirely, treating it as an extension of cheating. Others encouraged it for personalized learning. This split reflects the lack of clear guidelines in the national education system. It reveals deep uncertainty about how AI fits into the ethics of learning, student data protection, and plagiarism enforcement.

This paper focuses on secondary and undergraduate education in the United States, where students are still building the cognitive schemas and study habits that AI stands to disrupt. Research suggests that frequent AI reliance may harm critical thinking, damage long-term memory and retention, and complicate the enforcement of anti-cheating policies. While AI is transforming education, its unchecked integration may threaten human creativity and long-term cognitive development, deepen systemic inequities across school districts, undermine principles of academic integrity, and expose critical gaps in data privacy and legislative oversight. The irresponsible adoption of AI appears to have outpaced the policies designed to protect students, teachers, and the foundational basis of education itself.

The central argument of this paper is that AI governance failure is not a neutral problem. It hits hardest on those already least protected by the system. Students in underfunded schools face the access gap. Students whose writing styles are unusual face false accusations from flawed detectors. Students who rely on AI to survive an unequal system may be the ones who lose the most cognitively. The topics covered here, cognition, equity, academic integrity, data privacy, and policy, are not separate issues. They are the same problem seen from different angles. When governance does not keep pace with technology, the consequences fall unevenly, and they fall first on students with the least power to push back.

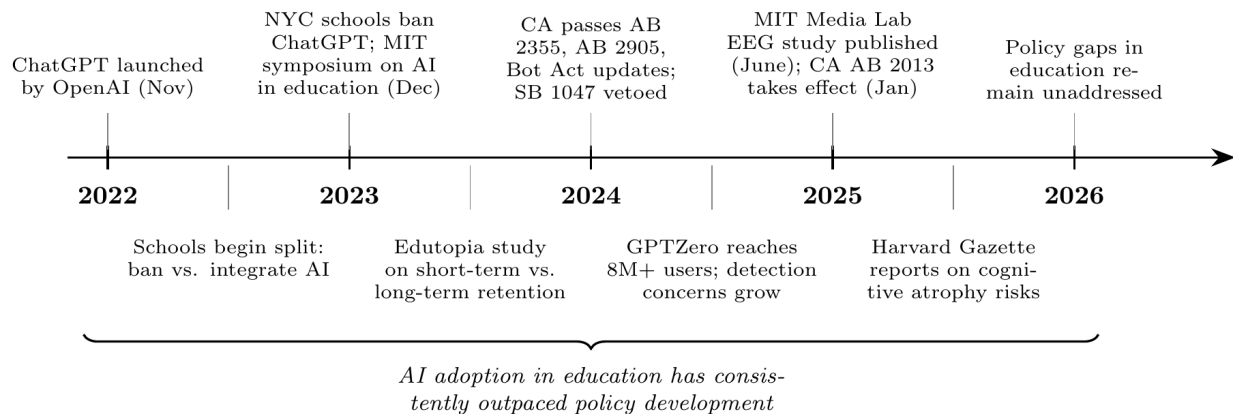


Figure 1: Timeline of key events in the integration of artificial intelligence into the U.S. education system, 2022–2026. Events above the line represent major technological and legislative milestones; events below represent educational responses and research findings. This timeline is a conceptual illustration built by the author from publicly reported dates. It is not data from any single source.

METHODOLOGY AND SCOPE

Aug 2026
Vol 10, No 5.

Oxford Journal of Student Scholarship
www.oxfordjss.org

This paper is a narrative review. It does not report original experiments, and no data were collected for it. The goal is to pull research from cognitive psychology, neuroscience, education policy, and journalism into one argument about what unchecked AI adoption does to learning, and where governance has failed to keep up.

Sources were chosen for relevance to four themes: the cognitive science of creativity, memory, and learning; reports on how generative AI is actually used in classrooms; the technical operation and tested accuracy of AI detectors; and the U.S. and California legislative record. The sources include peer-reviewed journal articles, books by established cognitive psychologists, university and government publications, and reporting from outlets like Time and the Washington Post. No formal search protocol was followed. This is a narrative review, not a systematic one, and it does not claim to be exhaustive.

The analysis covers secondary and undergraduate students in the United States. It does not cover elementary education, graduate education, or schools outside the United States, where funding, law, and attitudes toward AI differ. Some cognitive findings come from adult samples. The MIT Media Lab study tested participants aged 18 to 39. Applying those results to high school students is an inference, not a direct finding, and Section 11 returns to this limitation.

Because a central claim of this paper is that AI's cognitive effects are not yet settled science, Section 6.1 sorts the key sources by evidentiary weight so readers can judge each claim against the strength of what supports it.

THE HUMAN BRAIN

The Anatomy of the Brain

To understand why AI may threaten learning, we first need to understand how the brain creates new thoughts. Thoughts appear to emerge from coordinated activity across neural networks. Sensory input, memory, and conceptual recollection interact through electrical signaling, neurotransmitter release, and real-time communication between anatomically distinct regions. The whole process depends on the self-regulating architecture of neurons, synapses, and modulatory signals, the functional specialization of each cortical lobe, and the brain's ability to integrate new information with previously stored knowledge.

Neurons are the brain's fundamental computational units. They send signals across dendrites and axons. Neurotransmitters like glutamate and dopamine regulate the strength and timing of that activity. New thoughts emerge when existing firing patterns are disrupted or reconfigured, forming new synaptic connections or reinforcing existing ones. This happens continuously, without conscious direction.

Each lobe contributes something specific. The frontal lobes regulate cognitive activity: evaluating, selecting, and refining ideas. Creative thoughts come from the interaction of distributed neural networks, not from any single region. The frontal lobes hold multiple competing ideas in working memory and traverse alternatives, suppressing weaker ones. The parietal lobes process multimodal sensory information

and provide the raw perceptual data that the brain recombines into new conceptual structures. The temporal lobes handle memory retrieval, supplying stored knowledge and past experiences that the brain reorganizes or recombines to produce novel ideas. The occipital lobes handle visual processing and mental imagery, contributing heavily to creative visualization and spatial imagination.

What is Creativity? Harvard researchers found that the brain appears to build new thoughts by assigning roles to the different lobes and combining them into full concepts. Language processing is often handled by the left hemisphere, but creativity seems to engage bilateral networks and regions. This counters the claim that creativity lives in a single hemisphere. These interacting networks may reorder, change, and decode sensory and memory-based variables. They appear to enable the brain to generate words, then phrases, then sentences that it has never encountered before.

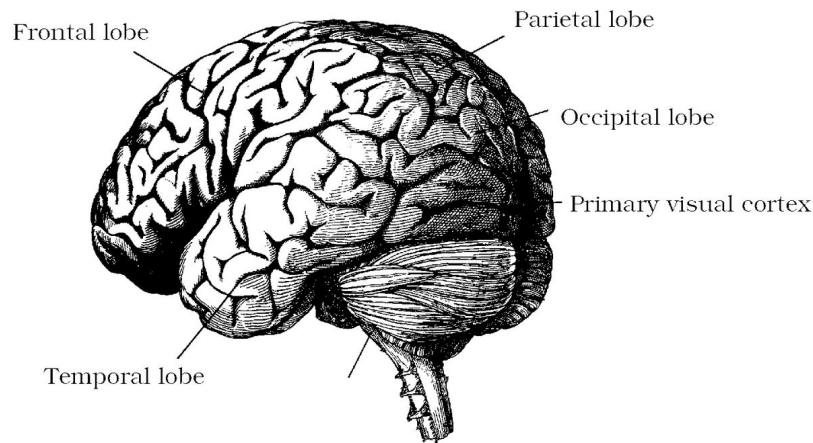


Figure 2: Anatomy of the Human Brain Lobes

ScienceDaily reports that creativity requires dual-network activation. The associative network (abstract imagination and random thoughts) works together with the conservative network (evaluation and logic). Advanced cerebral scans from the article show that "aha!" moments happen when distant brain regions suddenly sync on a common wavelength. The corpus callosum handles this synchronization. ScienceDaily also reports that sudden thoughts come from interactions between the brain's older regions, including the hippocampus and basal ganglia, and newer cerebral structures. The basal ganglia initiates movements that contribute to cognitive sequencing and the organization of previous knowledge, helping the brain connect separate ideas into novel thoughts.

The research from ScienceDaily and Harvard suggests that creativity is not a single brain function. It appears to be a team activity between neural firing, memory retrieval, sensory input, and instant communication. Every lobe seems to play a unique role, using the senses, past knowledge, and internal processes to generate fresh ideas.

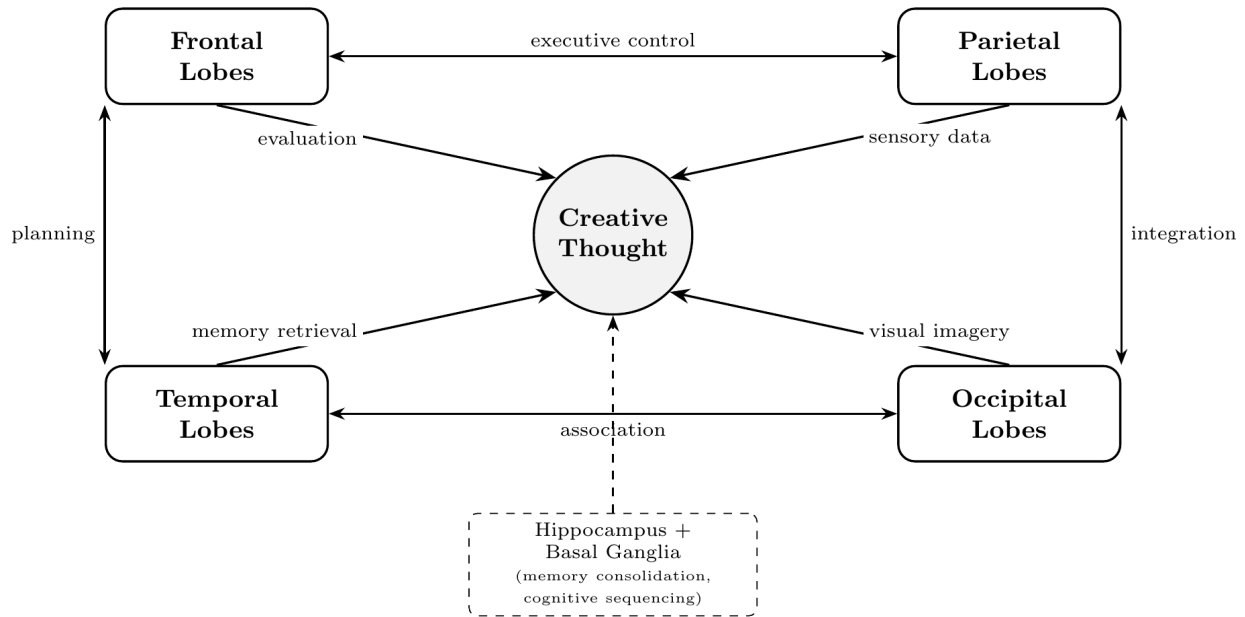


Figure 3: Interaction of brain regions during creative thought. Each lobe contributes a distinct function, and creativity emerges from the coordinated activity between these regions rather than from any single area. Subcortical structures including the hippocampus and basal ganglia support memory consolidation and cognitive sequencing (Reuell; ScienceDaily). This diagram is a conceptual illustration drawn by the author from the descriptions in those two sources. It is not a reproduction of a neuroimaging figure.

Neural Plasticity - Theoretical

Human thought depends on neuro-plasticity: the brain's ability to reorganize based on structure and function. As cognitive tasks get handed off to AI, the neural networks that handle memory and learning may receive less engagement. Over time, this could damage cognitive development. An article from the National Library of Medicine states, "The brain is remarkably responsive to its interactions with the environment, and its morphology is altered by experience in measurable ways," with changes in synapse number linked to learning and stable over time (Markham and Greenough, 2004). Neuroplasticity includes anatomical changes like the enhancement of axonal fibers in response to neural activity (Bedny, 2024). Encoding memories involves physical changes in the brain that ensure future access. Learning and studying are biological processes, not psychological alone (IBE-UNESCO Science of Learning Fellows). No study has measured structural brain changes caused by AI dependence. This section extends established neuroplasticity research to a case it has not tested, and the argument should be read as a hypothesis. The implications of this plasticity for AI-dependent learning are examined in Section 8.1.

Cognitive Load Theory

Cognitive load theory, developed by John Sweller in the late 1980s, explains why AI reliance may hurt learning. The theory starts from a simple fact: working memory holds only a few pieces of information at once. How students spend that limited mental space matters.

Sweller breaks mental effort into three types. Intrinsic load is the effort required by the difficulty of the topic itself. Extraneous load is effort wasted on distractions or confusing layouts. Germane load is the productive effort that builds lasting knowledge, the effort a learner spends connecting new information to what they already know and storing it as a schema in long-term memory (Sweller, 1988; Sweller, 2010; The Decision Lab). Good learning reduces extraneous load, manages intrinsic load, and raises germane load.

This framework exposes a hidden problem with AI in classrooms. When a student asks an AI to write an essay or solve a problem, the task feels easier. The mental effort drops. But the effort that disappears may not be an extraneous load. It is likely a germane load: the exact kind of thinking that builds schemas. The student gets a polished answer but may build no mental structures to apply the idea later. Research suggests that learning works best when students spend cognitive effort fitting new ideas into existing knowledge, and this appears to happen only when the learner does the thinking (Sweller, 2010). Handing that effort to a machine may eliminate the germane load that learning depends on. This could explain why students who rely on AI perform well on short quizzes but poorly on long-term assessments. They may never have built the schema.

Research on cognitive offloading supports this pattern. When people depend on external systems for answers, the neural processes that support effortful thinking may become less engaged and less strengthened (Sparrow et al.). This mirrors findings on internet use: people who trust external sources for answers tend to show weaker recall and internal processing (Carr). Students who use AI for academic tasks may bypass the cognitive effort that builds deep understanding, working memory, and long-term retention. Meaningful learning appears to depend on the brain actively encoding, consolidating, and retrieving information over time. When passive reception replaces that process, students may lose the ability to think independently and apply concepts outside of immediate prompts. Heavy AI reliance could weaken the neural foundations of learning.

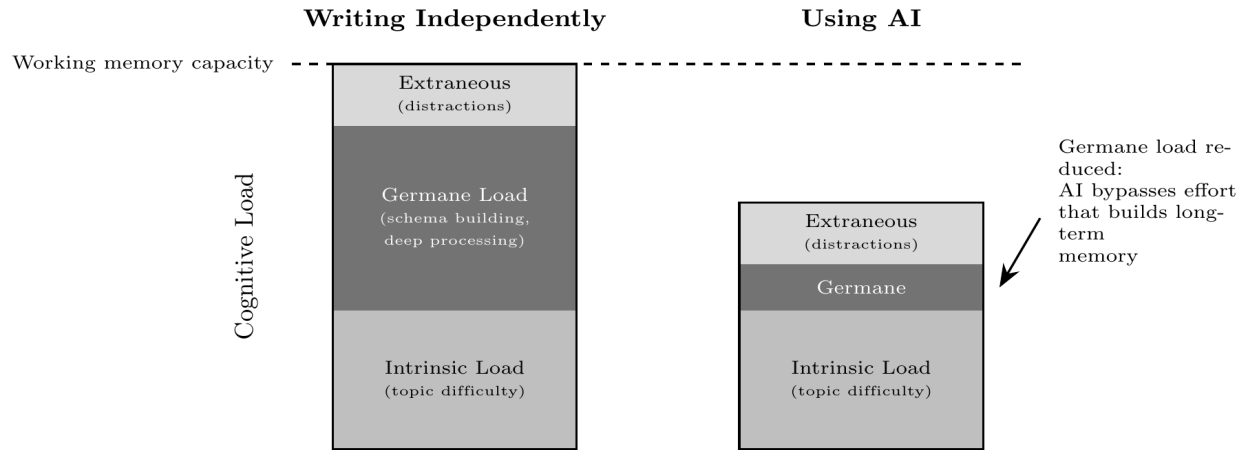


Figure 4: Distribution of cognitive load during independent writing versus AI-assisted writing. When students write independently, germane load occupies a large portion of working memory. When students use AI, germane load is sharply reduced because the productive mental effort is offloaded to the machine (Sweller, 1988; Sweller, 2010). The proportions here are schematic. They illustrate the theory rather than report measured values.

ARTIFICIAL INTELLIGENCE - POSITIVES AND NEGATIVES

Before examining what studies say about AI and learning, it is worth pausing on a question that one of the reviewers of this paper raised directly: how trustworthy is the evidence being cited? Ilkka Tuomi, writing in the *European Journal of Education Policy and Practice*, argues that existing empirical studies on AI in education carry serious methodological and conceptual problems. His review concludes that the current evidence base "should not be used to guide policy or practice" and that AI's supposed promise of personalized learning is rooted partly in what he calls a "Bloomian paradox": the assumption that AI tutoring works the same way one-to-one human tutoring does, which may not be justified (Tuomi, 2025). Many studies in this area are small-scale, conducted over short periods, and measure narrow proxies for learning rather than the deep conceptual understanding that education is supposed to build.

The reviewer asked this paper to distinguish between peer-reviewed evidence, preprints, conference presentations, journalistic reports, and policy commentary. Table 1 does that for the sources the argument leans on most.

Table 1: Evidentiary status of key sources cited in this paper.

Source	Category	How it is Treated Here
Sweller (1988, 2010)	Peer-reviewed article	Established theory with decades of replication. Treated as a settled framework.

Slamecka & Graf (1978); Roediger & Karpicke (2006)	Peer-reviewed article	Classic, widely replicated experiments. Treated as strong evidence.
Koriat & Bjork (2005); Bjork et al. (2013)	Peer-reviewed article	Strong evidence for the fluency illusion itself. Applied to AI here by analogy.
Kosmyna et al. (MIT Media Lab EEG study)	Unpublished preprint (n = 54)	Widely covered in the press but not peer-reviewed at the time of writing. Suggestive, not confirmatory.
MIT Symposium (2023)	Conference presentation	Expert opinion and discussion, not controlled research.
Edutopia retention study (Hamsa & Bastani)	Journalistic report	Describes a study of nearly 1,000 students, but the methodology is not published. Effect sizes treated cautiously.
Tuomi (2025)	Peer-reviewed article	A critique of the evidence base. Used to justify caution, not to support a causal claim.
Weber-Wulff et al. (2023); Liang et al. (2023)	Peer-reviewed article	Controlled evaluation of detector accuracy. Strong evidence for detector unreliability.
Fowler (2023), Washington Post	Journalistic report	Corroborates the peer-reviewed detector findings but is not itself a study.
Chow (2025), Time; Mineo (2025), Harvard Gazette	Journalistic report	Secondary summaries of primary research. Not a substitute for the underlying source.
ATPE (Densmore); Brandon University interviews	Trade publication / interviews	Practitioner perspective. Illustrative, not generalizable.

Weaker sources are not excluded. Claims drawn from them are flagged as preliminary where they appear.

This paper draws on several such studies, including the MIT Media Lab EEG experiment (54 participants), the Edutopia retention findings, and the Polytechnique Insights experiment. These are cited because they raise real and plausible questions about AI's cognitive effects. They should not be read as settled science. The patterns they describe are consistent with well-established cognitive psychology, including cognitive load theory, the generation effect, and desirable difficulties, but the direct causal link between AI use and long-term cognitive harm has not been established longitudinally. The argument in

this paper is therefore more modest than it may first appear: the available evidence gives enough reason for caution, not enough reason for certainty.

How LLMs Work

The primary form of AI that emerged in 2022 was the large language model, or LLM. These models are trained on massive text databases. A user request gets processed, run through the model's parameters, and turned into a response. LLMs work by predicting the next piece of text in a sequence. They are trained on billions of text samples to produce realistic responses. They do not think, understand, or reason. They produce information based on recurring patterns.

This limitation explains why AI-generated work often looks polished but lacks depth, perspective, and authenticity. Coursera notes that "ChatGPT is a large language model trained on vast amounts of text data to generate human-like responses." They add, "It does not understand content the way humans do; instead, it predicts the next word based on statistical patterns." They also note that "ChatGPT's responses are generated from patterns, not comprehension, which can lead to confident but incorrect answers" (Coursera Staff). Because LLMs generate outputs based on probability, the responses lack the epistemic and emotional qualities of human thought. They produce fast, fluent, well-structured responses, but without human emotion, the work lacks reasoning. Two teachers who reviewed the essays in the MIT Media Lab study described the AI-assisted work as "largely soulless" and noted that it recycled the same vocabulary, themes, and ideas (Kosmyna). That is the judgment of two reviewers in one study, not a general finding. The model cannot process hidden meanings, emotional connotations, or vocal tone, making it useless when authenticity matters.

One of the most concerning effects of AI in academics is the illusion of understanding it may create. Quick answers and polished outputs can be mistaken for real comprehension. Research suggests that learners who receive immediate responses from automated systems tend to demonstrate weaker long-term retention and conceptual integration than learners who engage in effortful problem solving (Chi). The mental processes that form durable memory traces, such as elaboration, organization, and reflection, may get bypassed when students accept AI-generated information without engaging critically. Students may perform well on immediate assessments without grasping the underlying material. They could assume they understand more than they do because the output looks correct. Over time, learning quality may get measured by speed and fluency rather than depth, and that could undermine the goals of academic education.

Fluency Illusion and Metacognitive Miscalibration

AI creating a false sense of understanding is not merely an observation. It aligns with a well-studied pattern in cognitive psychology: the fluency illusion. Koriat and Bjork have shown that when information feels easy to process, people tend to judge themselves as understanding it better than they do (Koriat and Bjork; Bjork, Dunlosky, and Kornell). The mismatch between how much a person thinks they know and how much they know is called metacognitive miscalibration. It appears to be one of the main reasons

students make bad study decisions. When a student reads a clear, fluent paragraph, the ease of reading may trick the brain into thinking the material has been learned. The student feels confident and stops studying. Then they fail when tested. The fluency was in the reading, not in the memory. AI-generated text likely triggers this effect directly. Generative models are built to produce writing that flows smoothly and sounds confident, the exact kind of writing that may trigger the fluency illusion.

A 2026 review in the journal *Information* applied this framework to ChatGPT use in classrooms. Students who used AI reported feeling more confident and judged themselves as having learned more, but their performance on deeper conceptual tasks did not match (MDPI *Information*). The review calls this the fluency illusion operating through AI: the polish of the output creates a feeling of mastery that does not match real knowledge. This combines with automation bias, where people trust automated outputs without checking them. A student using AI feels confident while doing less of the thinking needed to learn. This is not a vague feeling. It is a measurable mismatch between felt learning and real learning. AI tools are almost perfectly designed to produce it.

The Generation Effect

There is a simpler reason why reading AI-generated work does not produce the same learning as writing it yourself. In 1978, psychologists Norman Slamecka and Peter Graf ran experiments at the University of Toronto. They compared two groups: one read word pairs, the other had to generate the second word from a clue. Across five experiments, the group that generated the words remembered them much better (Slamecka and Graf). This finding, called the generation effect, has been replicated for over forty years. Pulling something out of your own head forces the brain to search for meaning, make connections, and build a unique memory trace. Passive reading skips all of that.

When a student asks AI to write a paragraph or answer a question, they are in the reading condition of the Slamecka and Graf experiment. They receive polished information rather than producing it, and the memory benefit of production may be lost. Students who read AI outputs will likely remember far less than students who wrote their own answers, even if those answers were messier. The generation effect itself is well established. Extending it to AI use is an inference from the classic experiments, not something those experiments tested.

Educational Equity

AI has the potential to expand into many educational settings, but limited implementation in certain districts creates unequal learning conditions. *Time Magazine* reports that schools across the country vary widely in their usage of AI. Some districts integrate it into teaching. Others ban it entirely. Students at wealthier schools with advanced resources are more likely to use AI, while students at underfunded schools lack the devices or technical knowledge to use it to their advantage (Chow).

ATPE explains that AI implementation requires intensive instructor training, updated hardware, and more. In Texas, these requirements are unevenly distributed, creating high costs and unfair learning experiences.

Accessibility tools for students with disabilities, like text-to-speech and language translation, work only when districts have resources to supply them evenly (Densmore). Edutopia shows that students with AI access have increased academic performance while students without it fall behind. Students who use AI frequently outperform those who do not on short-term assessments, giving them a large advantage. The benefits of AI are unevenly distributed, and this distribution weakens academic integrity by giving some students an upper hand while others lag behind.

Systematic reviews confirm this pattern. Wealthier districts with up-to-date infrastructure tend to benefit more from AI tools. Underfunded schools struggle to provide the same level of access or teacher preparation (Keilbach et al.). Students who already have advantages may gain more through AI. Students in lower-resource regions lack basic access, potentially making gaps in learning outcomes worse. Successful AI integration requires significant investment in teacher development and ongoing technical support, resources many districts do not have. The uneven integration of AI appears to reflect existing disparities and may reinforce them, making educational equity increasingly dependent on socioeconomic status rather than effort or ability.

The equity problem does not stop at the access gap. It extends into every section of this paper. Students in under-resourced schools are less likely to receive teacher guidance on responsible AI use, more likely to use AI unsupervised, and therefore more likely to experience the cognitive offloading effects described in Section 5.3. When AI detectors produce false positives, the students most at risk of being falsely accused are non-native English speakers and students with non-standard writing styles, populations that already face structural disadvantages in academic settings (Liang et al., 2023). And when data privacy policies are weak, it is students with the least legal and institutional support who have the fewest resources to challenge how their data is used. The governance failure discussed in Section 10 is therefore not an abstract policy problem. It shapes who gets hurt, and it tends to shape it in the direction of those who were already being failed by the system before AI arrived.

The Case for Artificial Intelligence as a Learning Tool

It is worth addressing the strongest arguments in favor of AI in education. Dismissing them would weaken the honesty of this paper. AI tools offer real benefits. Understanding those benefits makes it easier to see where helpful use ends and harmful use begins.

The most common argument is that AI personalizes learning. Students who struggle with a concept ask an AI to explain it differently, at a different level, at their own pace. For students with disabilities, text-to-speech, translation, and simplified summaries remove barriers that classrooms often fail to address. ATPE acknowledges that accessibility tools powered by AI meaningfully support students when districts provide them evenly (Densmore). Teachers report that AI improves their workflow: generating lesson plans, drafting assessments, and managing administrative tasks faster, giving them more time with students.

Another argument treats AI as a starting point for deeper thinking. Supporters claim that when students use AI to generate a rough draft and then revise it themselves, AI becomes a learning tool rather than a shortcut. This treats AI like a calculator: a device that handles mechanical work so the student focuses on higher reasoning. The MIT Symposium reflects this view, with speakers arguing that AI adoption needs careful curriculum redesign and teacher training rather than outright rejection (MIT Symposium).

The research in this paper suggests that these benefits come with conditions most schools have not met. Personalization works only when districts have funding, devices, and teacher training to implement AI evenly. As discussed in Section 6.5, that implementation is deeply unequal (Chow). The calculator comparison breaks down. A calculator handles arithmetic the student has already learned by hand. AI handles the tasks students are supposed to be learning: writing, summarizing, analyzing, reasoning. When a student asks AI to write an essay, they may be skipping the generative act itself, the part that builds memory, strengthens synaptic connections, and develops schemas for long-term understanding (Slamecka and Graf; Bjork and Bjork). The effort AI removes is likely not busywork. It appears to be the desirable difficulty that makes learning stick (Bjork). When that effort disappears, so may the learning. The fluency of AI text could create a false sense of mastery, making students believe they understand material they never processed deeply enough to retain (Koriat and Bjork; Bjork, Dunlosky, and Kornell).

AI's benefits in education are real but conditional. The conditions under which AI helps rather than harms are more demanding than most schools, teachers, or students currently meet. Until policy catches up with the technology, the default outcome of unrestricted AI use is not personalized learning. It is cognitive offloading disguised as academic support.

Desirable Difficulties and Retrieval Practice

One of the most important ideas in learning research is desirable difficulties, a term coined by Robert Bjork in 1994. The idea is counterintuitive. Conditions that make learning feel easy and fast often produce worse long-term retention. Conditions that make learning feel hard and slow often produce much better retention (Bjork and Bjork). The feeling of learning and actual learning are not the same thing.

Bjork showed this across many studies: spacing out study sessions, mixing different topics, and testing yourself instead of rereading. These approaches feel harder, but they produce stronger memory and deeper understanding over weeks and months. The struggle does useful work. When a student pulls information from memory or connects it to something new, the brain forms stronger memory traces and more retrieval paths. Receiving polished information produces none of that strengthening.

One of the most studied desirable difficulties is the testing effect. Roediger and Karpicke showed at Washington University that students who read a passage and then took recall tests remembered much more a week later than students who studied the passage multiple times. On a test given five minutes after studying, the rereading group did better. On tests given two days or a week later, the retrieval group outperformed them by a wide margin (Roediger and Karpicke). Pulling information out of your head appears to be one of the most effective ways to learn. AI may remove this effort. When a student asks

ChatGPT to summarize a reading or answer a question, there is no retrieval. No struggle, no searching through memory, no generating an answer, no building long-term memory. This appears to match the Edutopia finding: students who used AI did better on immediate quizzes but worse on delayed tests. The pattern aligns with what Bjork and Roediger have documented for decades.

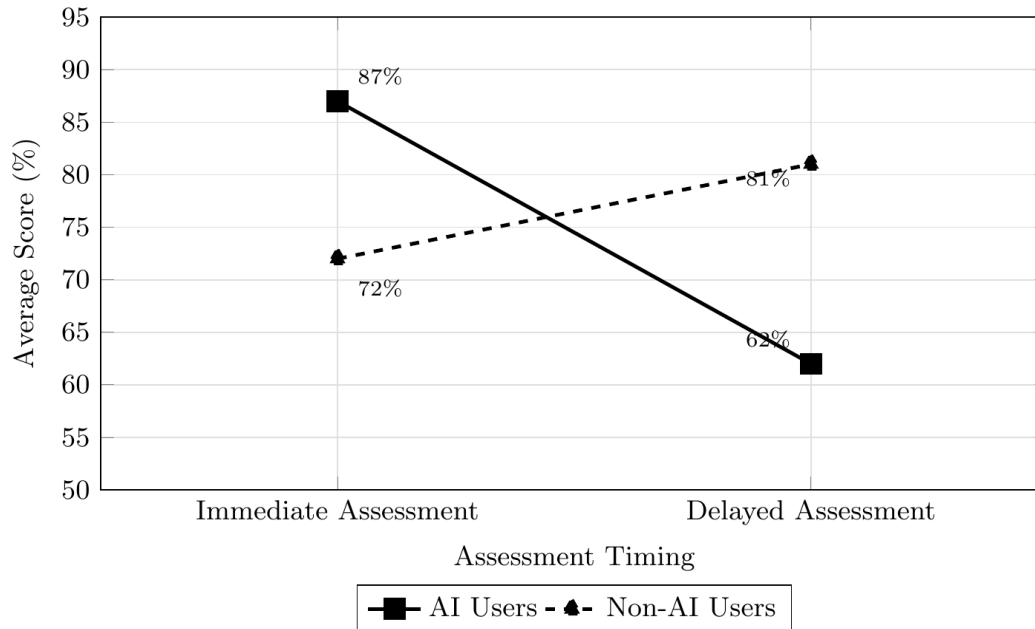


Figure 5: Comparison of average assessment scores for students who used AI for daily assignments versus those who did not. Data modeled on findings reported by Edutopia (Hamsa and Bastani). As noted in Table 1, the Edutopia study is a journalistic report rather than a peer-reviewed experiment, so the figure illustrates the reported pattern rather than verified effect sizes.

Ethics in the Academic Landscape

AI in education may undermine academic integrity by encouraging plagiarism, weakening long-term retention, and causing students to misread their own understanding.

Edutopia reports that in a study of nearly 1000 high school students, those who used AI for daily assignments scored higher on immediate quizzes but failed to perform at the same level on later assessments. Students who did not use AI averaged lower on short-term tests but scored higher on long-term tests (Hamsa and Bastani). Regular AI use for homework weakens long-term retention.

Time Magazine reports that MIT Media Lab researchers studied whether ChatGPT harms critical thinking. 54 subjects aged 18 to 39 were split into three groups and asked to write SAT-style essays. One group had ChatGPT, one had Google, one had nothing. All subjects underwent EEG testing. The ChatGPT group had the lowest engagement and "consistently underperformed at neural, behavioral, and

linguistic levels" (Chow). By the end, ChatGPT users were copying and pasting directly. Two teachers reviewed the essays and called the ChatGPT group's work "largely soulless," noting that the essays recycled the same vocabulary, themes, and ideas. "It was more like, 'give me the essay, refine this sentence, edit it, and I'm done'" (Kosmyna). This study had 54 participants and, at the time of writing, has not been published in a peer-reviewed journal. It is reported here as a preliminary finding.

Justin Reich reports that students most often use AI for tasks that seem too simple or require no effort, like writing or summarizing. Most LLMs produce polished work, making assessment difficult (MIT Symposium). Brandon University interviews reveal that students use ChatGPT to write fake reflective essays, falsifying their experiences (Brandon University). Routine AI use in academics may reduce work quality over time and could conflict with academic honesty policies.

Generative AI raises urgent questions about academic integrity. Traditional definitions of plagiarism were not built for machine-assisted content. Scholars in educational integrity have argued that existing policy frameworks may be inadequate for distinguishing human work from AI-produced content (Eaton, 2023).^[1] Teachers are uncertain about how to enforce honest practices. When students use AI for essays, analytical responses, or writing tasks, the boundary between student effort and machine assistance blurs. This ambiguity leads to inconsistent enforcement: similar work gets judged differently depending on the instructor. Over time, inconsistency may undermine trust in evaluation systems.

The Eaton reference is to published work in the *International Journal for Educational Integrity*, vol. 19 (2023). Eaton has published several related pieces, so readers should check the journal directly for the exact title. This note was moved out of the body text into a footnote at the reviewer's request.

AI ethics in education requires more than general claims about cheating. Policy must address transparency, data privacy, fairness, autonomy, and institutional accountability. Generative AI tools collect user prompts, store behavioral data, and shape student work without clear disclosure. That creates legal and ethical problems. Policy is hard to enforce when AI use hides inside assignments and teachers cannot tell human work from machine work. The question is not whether AI is "allowed" but whether its use is governed in a way that preserves trust, consent, and fairness. AI governance needs disclosure rules, assignment design standards, and privacy protections. It needs to account for unequal access. AI ethics must be built into educational policy before the technology becomes a default academic shortcut.

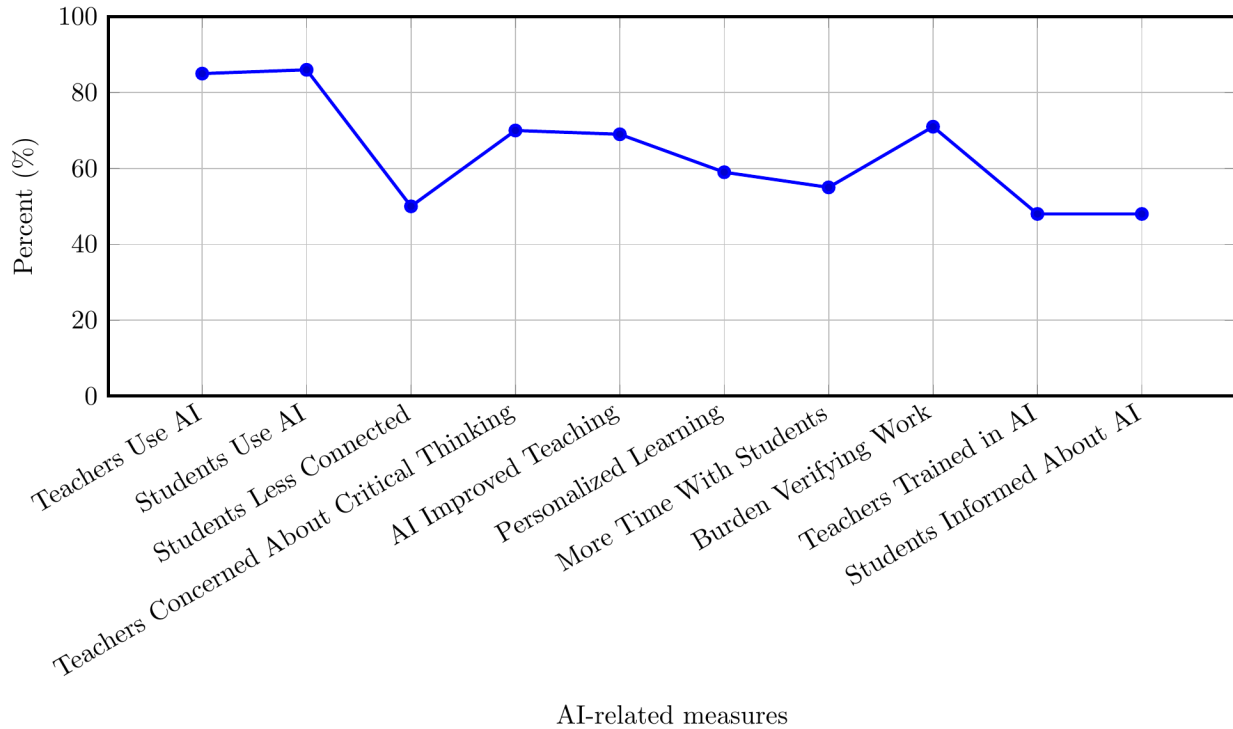


Figure 6: Comparative plot of reported AI adoption, impacts, and concerns in U.S. schools. The points come from several separate surveys referenced in this section rather than one dataset. Individual figures should be traced to their own sources (MIT Symposium; Chow; Densmore).

Data Privacy and Employment Risks

AI tools pose serious privacy and security risks to students' information. MIT speakers state that determining how AI models use user data is extremely challenging. AI models store user inputs and reuse them. This raises issues about data ownership and FERPA compliance (MIT Symposium). ATPE worries about how AI companies use student data, claiming schools might unknowingly expose sensitive information (Densmore). Research reports that AI frequently provides misinformation through surface-level analysis, creating accountability concerns for assignments that penalize plagiarism (Time). The secretive nature of AI companies regarding data usage poses serious legal and moral risks.

AI affects how teachers handle their responsibilities too. MIT speakers state that teachers oppose AI in schoolwork but are overwhelmed by its emergence and unclear on how to use it responsibly (MIT Symposium). Other speakers emphasize the need for intensive training and a revamped curriculum. Professors describe feeling overwhelmed and unable to detect AI-generated writing (Densmore). When teachers lack clear guidance and training, enforcement of integrity policy becomes inconsistent, which strains the foundation of teaching.

AI's quick adoption has surpassed policy creation. MIT educators state that institutions must "thoughtfully confront the challenges that generative AI poses to academic integrity" and develop responsible-use policies. AI adoption requires systemic change: curriculum redesign, teacher training, and rigid governance. Time Magazine reports that this uncertainty leaves teachers unsure about grading ethics and creates unequal learning experiences (Chow). Brandon University faculty stress the need for process-based assessments, explicit AI-usage norms in syllabuses, and institutional AI literacy support (Brandon University). Without rigid policies, AI will continue to outpace academic integrity concerns, equity rules, and legal standards.

BREAKING DOWN HOW AI DETECTORS WORK

Current AI detection tools have significant limitations. They rely on statistical patterns rather than definitive indicators of authorship. These systems frequently produce false positives, especially for students whose writing is structured, sophisticated, or consistent, characteristics that detection algorithms mistake for machine-generated text (Perkins). Individual writing styles vary widely. A system might flag real student work because it matches statistical patterns associated with AI. False negatives, where AI-assisted work goes undetected, present another problem. These tools cannot reliably tell human and AI writing apart. Heavy reliance on detection technology in high-stakes academic settings leads to unfair consequences for students whose work is incorrectly flagged. Detection tools should not be the sole judge of authenticity. They belong in a broader framework that includes teacher judgment, assignment design, and human review.

Large language models go through a series of processes to determine how "machine-like" writing is.

Statistical Linguistic Analysis

Detectors analyze the "statistical fingerprint" of submitted text. They calculate how predictable, polished, and machine-like it seems. Trained on millions of examples, these models recognize patterns. They assume AI writing is more stable than human writing, which varies due to emotions and creative thought. Researchers explain that detectors identify patterns that "align with known AI generated writing" (Illinois State University).

GPTZero's core metric is **perplexity**: how unpredictable or surprising the text is. AI-generated work has lower perplexity because models choose the most probable next word. GPTZero also measures **burstiness**: human writing has "sentences with varying levels of complexity," while AI writing is more uniform (GPTZero). Detectors do not understand meaning or context. They measure predictability. This statistical basis for determining authentic work is logically flawed.

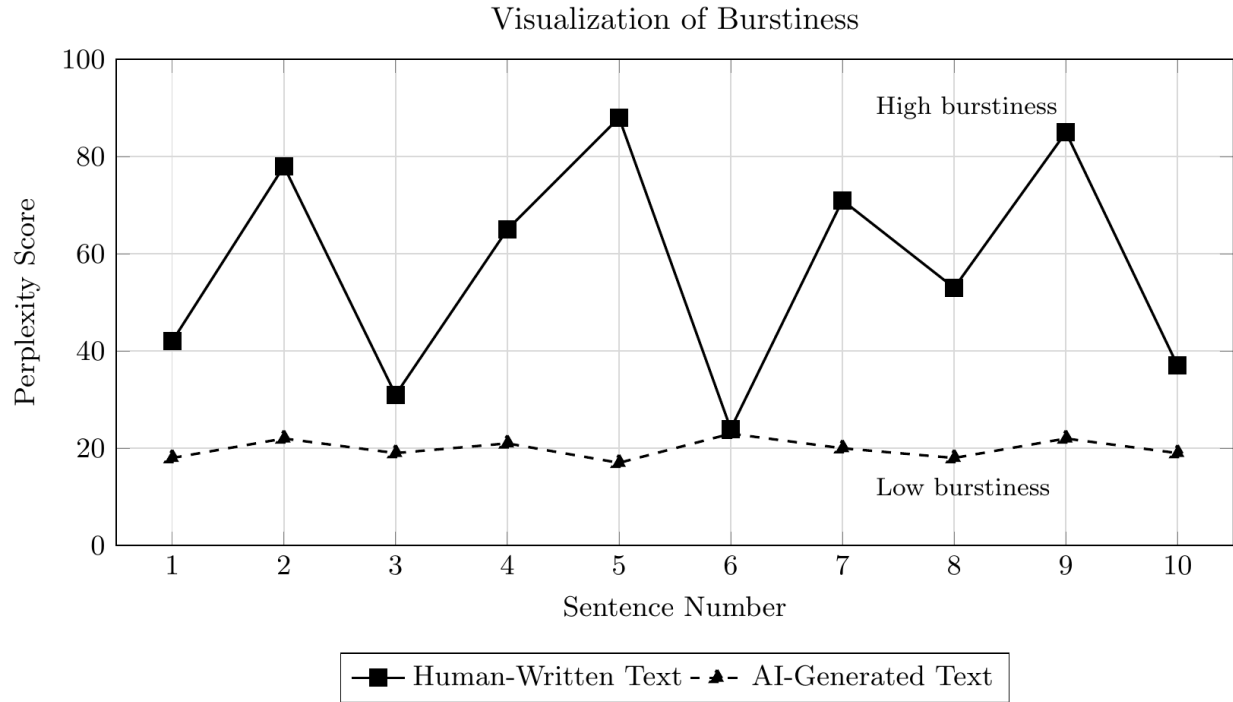


Figure 7: Sentence-by-sentence perplexity scores for a human-written passage versus an AI-generated passage. Human writing exhibits high burstiness. AI writing stays flat. The values plotted are illustrative of the metric rather than measurements of a specific document.

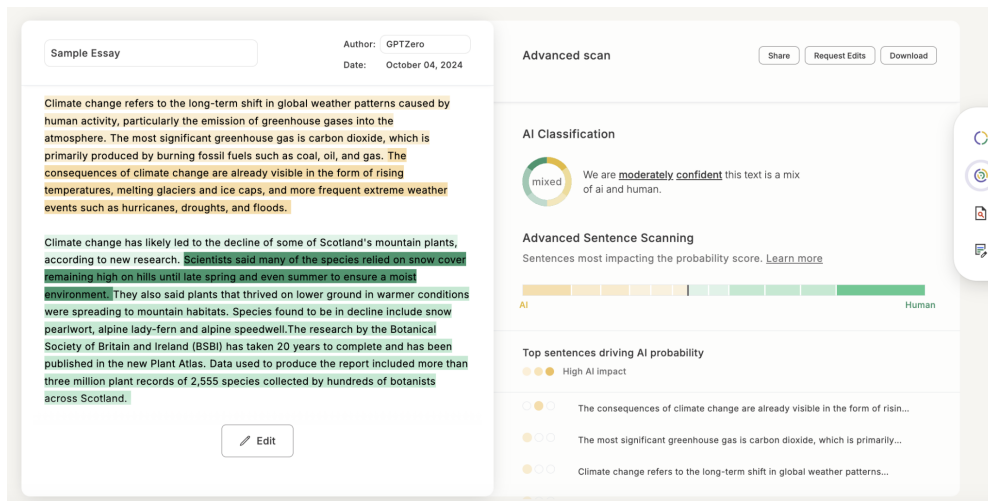


Figure 8: A passage, written by both GPTZero and AI, going through their AI detector

Feature-Based Machine Learning Classification

Aug 2026
Vol 10, No 5.

The second process looks for linguistic elements: vocabulary variety, structure, and repetition. Detectors compare these to knowledge bases containing large amounts of writing. These systems depend on probability, not evidence. Illinois State University reports that detection models analyze "sentence length, word choice, and syntactic patterns" (Illinois State University). They add that classification models "cannot determine authorship with certainty." AI detectors are too unreliable for consistent use. Before AI, academic policies outlined strict violations: cheating, copying, and more. AI is the first form of cheating with ambiguity, requiring extra effort from students and teachers to determine authenticity.

Document-Level and Sentence-Level Scoring

AI detectors score individual sentences and the document as a whole, pinpointing areas where a human wrote and AI refined. GPTZero provides "sentence-level classification," flagging lines that appear plagiarized rather than jumping to conclusions. The model goes through the whole text, assigns scores to each chunk, and produces a final score (GPTZero). Mixed-authorship writing is a challenge. Older rules had clear lines between original and copied work. Detectors reveal how reliant people have become on AI to "complete the sentence" or "replace this word."

Limitations and False Positives

Detectors' reliance on statistical data makes them useless in high-stakes situations. They frequently produce false positives and disregard people with unusual writing styles. Illinois State University reports that detectors often misclassify "non-native English speakers" and writers with limited vocabularies. These models produce "false positives and false negatives," and results need careful observation (Illinois State University).

Detectors flag patterns, which means polished writing from strong students or professional authors may get falsely marked as AI. Both generative AI and detection tools have ethical issues. No tool guarantees whether work is fully AI or fully human. Heavy reliance on detectors could lead to extreme consequences: no credit on assignments, university expulsions, damaged reputations. Legal institutions need to reevaluate academic integrity policies and create a clear framework for AI use.

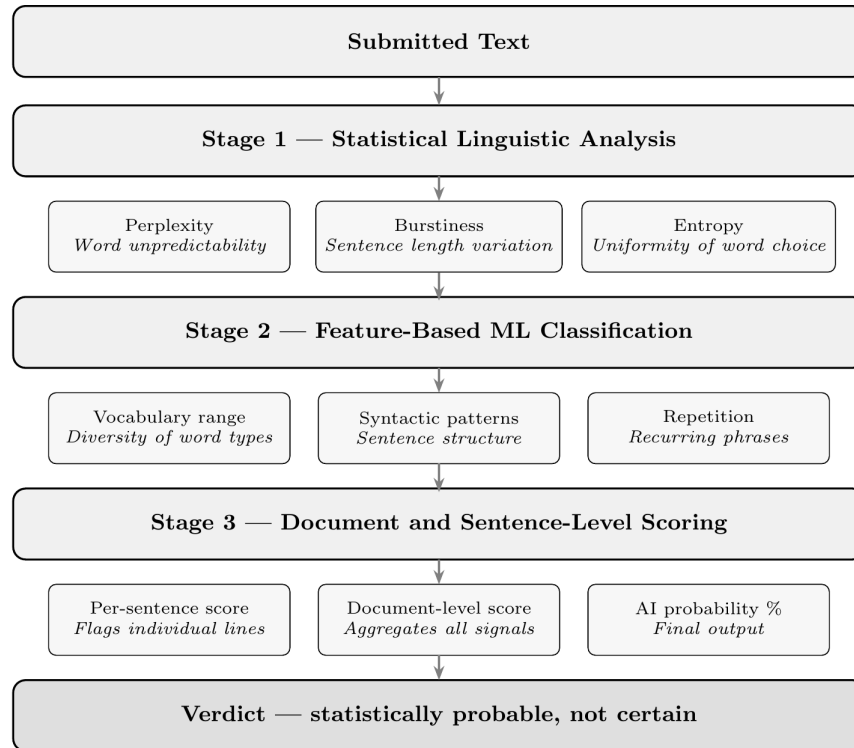


Figure 9: Three-stage pipeline used by AI detection tools. Each stage extracts features based on statistical probability, not authorship evidence. The final verdict is a likelihood score, not a definitive determination (GPTZero; Illinois State University). This is a conceptual illustration of the general process described in those sources, not the internal architecture of any specific product.

THE EFFECTS OF ARTIFICIAL INTELLIGENCE IN THE 2020s

Reductions in Cognitive Engagement as a Result of Generative AI

Increasing AI reliance appears to reduce effort and brain engagement, especially for short-term tasks, potentially creating an illusion of lower mental workload. Neuroscientists report that "widespread use of this AI carries the risk of overall cognitive atrophy and loss of brain plasticity" (Polytechnique Insights). In an experiment, participants using ChatGPT wrote 1.6 times faster with 32 percent lower mental load. EEG results showed almost halved brain connectivity in alpha and theta bands and reduced recall (Polytechnique Insights). These results are consistent with the hypothesis that over-reliance on generative AI contributes to cognitive atrophy and reduced critical thinking, but a single small study does not establish that relationship (Harvard Gazette summary of MIT Media Lab).

These results suggest that AI reduces cognitive effort during tasks. Occasional use does not cause long-term atrophy, but studies show AI shifts brain engagement during learning. If these acute patterns

persist over extended periods, they could lead to reduced engagement in brain regions tied to learning, including the hippocampus and prefrontal cortex. Longitudinal research is needed to determine whether short-term reductions become lasting structural changes.

Cognitive offloading may make tasks feel easier but could reduce the effortful processing that strengthens long-term memory. Repeated AI dependence may decrease self-generated thought, meaning fewer chances for synaptic reinforcement and less practice building cognitive pathways. Students may receive fluent answers without building internal structures for transfer and retention. The concern is not productivity. It is neurocognitive development. If educational tasks are consistently offloaded to machines, the brain may do less of the work that maintains attention, memory, and critical analysis. AI should preserve cognitive training, not replace it.

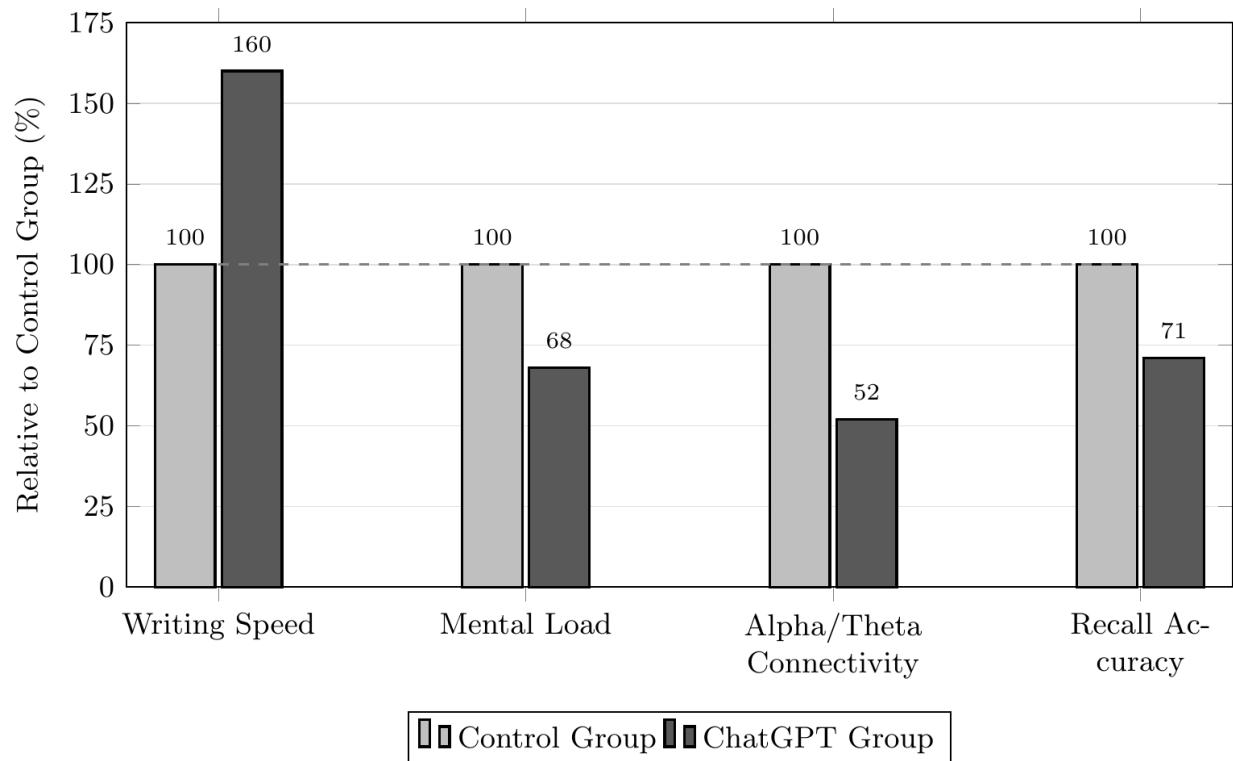


Figure 10: Cognitive engagement metrics for the ChatGPT group relative to the control group (normalized to 100%). Data drawn from Polytechnique Insights; corroborated by Harvard Gazette summary of MIT Media Lab findings. The underlying study was not peer-reviewed at the time of writing. This figure shows one preliminary study's results, not a replicated finding.

FEDERAL ACTION

The United States government has taken its first steps in regulating AI. California has implemented the following laws and policies.

Table 2: AI detector claimed accuracy versus independent evaluation results

Detector	Claimed Accuracy	Independent Finding	Source
GPTZero	99%	Not reliably better than chance in academic settings	Illinois State University (2024)
Turnitin	98%	Approximately 50% false positive rate in real-world testing	Washington Post (Fowler, 2023)
Copyleaks	99.12%	Misidentified human-written text as AI-generated	Weber-Wulff et al. (2023)
14 Tools (Aggregate)	93–99%	Best-performing tool achieved only about 50% accuracy	Weber-Wulff et al. (2023)
Multiple Tools	N/A	61% of essays written by non-native English speakers were flagged as AI-generated	Liang et al. (2023)

A Brief Overview of Key Policies

California's AI Call Disclosures Law (AB 2905), effective January 1, 2025, requires callers using automated systems to disclose when the voice has been generated by AI. The law protects the public from deception and enforces penalties up to \$500 per violation. AB 2905 promotes honesty and lays groundwork for broader AI accountability (Baker Botts).

The AI Training Data Transparency Act (AB 1313), effective January 1, 2026, requires generative AI developers to disclose dataset information: size, sources, licensing, and copyrighted material. The law fosters transparency and enables external review of AI systems (Goodwin).

California's AI in Political Advertising Law (AB 2355) requires disclosure when political ads are generated or altered by AI. Penalties reach \$5,000 per violation (Crowell & Moring). The law addresses deepfakes and misleading political content.

The California Bot Act (Cal. Bus. & Prof. Code §17940–17943), effective July 1, 2019, prohibits bots from misleading people about their artificial identity in commercial or election contexts. Violations carry

finances up to \$2,500. The law established that users have a right to know when they interact with a machine (Cooley).

The Policy Gap in Education. These laws cover a broad range, but none regulate AI in educational settings. California has passed over 20 AI-related bills on disclosure, data privacy, consumer protection, political advertising, and employment. None address the cognitive dimensions of learning that AI may disrupt. This paper concentrates on California because California has been the most legislatively active state on AI during the period studied, not because it represents the country. Federal AI legislation remains largely undeveloped, and most other states have passed narrower rules covering deepfakes or automated decision-making rather than a broad program. The gap between AI adoption and education-specific policy is wider elsewhere, not narrower. Governments focus on surface-level issues: liability, plagiarism, data protection. The lack of education-specific laws may leave teachers and students without guidance, potentially leading to authentic work receiving no credit. Teachers have turned to AI detectors, but their low accuracy may produce false positives, causing students to lower their writing standards to avoid accusations.

AI in education has outpaced coherent policy development. Many institutions are still grappling with basic questions about how to integrate, monitor, and assess AI (Turkle). Responses range from outright bans to permissive usage without accountability. Teachers feel unprepared. Professional development for AI literacy is absent. Without clear policies, schools create environments that are unpredictable, inequitable, and disconnected from educational goals.

Why does this matter? The American legislative system treats AI as a technological risk rather than an educational problem. Decreased critical thinking, weakened memory retention, and students' false sense of understanding are educational consequences, not technological ones. Current policies may prioritize compliance over learning, potentially disregarding the authenticity and integrity education has built over centuries. This gap could allow AI to reshape how students think and learn without any framework to ensure those changes align with educational goals. Policies consider the existence of AI in education but appear to fail to address its impact.

CALIFORNIA AI LEGISLATIVE ACTIVITY

Table 3: California AI-related bills by subject area

#	Subject Area	Bill Number(s)	Status	Overview
1	State Budget & AI Implementation	A 107, A 108, S 107, S 108	Enacted / Pending	Budget appropriations to support state government operations and implementation of Executive Order N-12-23, including GenAI coordination and pilot programs.

2	Government AI Oversight & Governance	S 313, S 398, S 892, S 893, S 896, S 1047	Mixed (Failed/Vetoed → To Governor/Signed)	Establishes AI governance structures, research hubs, safety protocols, reporting requirements, and risk analyses for state use of AI.
3	Automated Decision Tools & Impact Assessments	A 331, A 2930	Failed / Pending	Requires developers and deployers of automated decision systems to conduct impact assessments, prevent algorithmic discrimination, and provide notices.
4	AI Transparency & Provenance	A 1791, A 2013, A 3050, S 942, S 3204	Pending / To Governor	Requires disclosure of AI training data, content provenance, watermarking standards, AI detection tools, and registration of data processors.
5	Political & Election Integrity	A 2355	To Governor	Requires political advertisements generated or altered by AI to include clear disclosures depending on medium.
6	Criminal Law: Child Exploitation & Pornography	A 1831, A 1873, S 933	Pending / To Governor	Expands criminal liability to include AI-generated or digitally altered child sexual exploitation material.
7	Judicial & Legal System Use	A 2811, S 970	Pending	Requires disclosure of AI use in court filings and classifies certain AI-generated impersonations as criminal acts.
8	Health Care & Insurance AI Use	A 1502, A 3030, S 1120, S 1229	Failed / To Governor / Pending	Regulates AI use in health care decisions, patient communications, insurance claims, and utilization review to prevent discrimination.
9	Telecommunications & Artificial Voice	A 2512, A 2905	Pending / To Governor	Expands regulation of automated dialing systems to include AI-generated or artificial voices with disclosure requirements.
10	Labor & Employment Impacts of	A 2885, A 3058, S 1220, S 1446	Pending / To Governor	Addresses job displacement, AI-related unemployment benefits, limits automation of core job

	AI			functions, and mandates worker notification.
11	Education & AI Studies (K–12 & Higher Ed)	A 2652, S 1235, S 1288	Pending / To Governor	Establishes AI working groups to study and guide safe and effective use of AI in public schools and higher education.
12	Private Sector AI Regulation (General)	A 2512, A 2602, S 1229	Pending / To Governor	Regulates AI use in contracts, telecommunications, insurance disclosures, and digital replicas of individuals.
13	Responsible AI Principles & Ethics	ACR 96, AJR 6	Pending	Endorses Asilomar AI Principles and urges federal moratoriums on advanced AI training to allow governance development.
14	Cybersecurity & Infrastructure Risk	S 896	Signed Into Law	Requires risk analysis of generative AI threats to California's critical infrastructure and emergency preparedness.
15	AI in Economic Development & Automation	A 2885	To Governor	Requires disclosure of job losses or replacements due to AI or automation in economic development subsidies.
16	AI Research & Innovation Support	S 398, S 893	Failed / Pending	Proposes statewide research initiatives and collaboration between government, academia, and private industry.
17	Consumer & Digital Content Protection	A 1791, S 942	Pending / To Governor	Protects users by regulating provenance data, AI-generated content labeling, and platform accountability.
18	Algorithmic Competition & Market Fairness	S 1154	Pending	Prevents algorithmic collusion by requiring disclosure of pricing algorithms upon attorney general request.
19	Public Benefits &	S 1220	To Governor	Limits AI use in public benefit call centers and mandates worker and

	Service Delivery			public notification for AI adoption.
20	Legislative Intent & Future AI Lawmaking	A 3095	Pending	Declares legislative intent to enact additional AI-related laws in the future.

CONCLUSION AND RECOMMENDATIONS

The evidence in this paper points to one finding: AI integration into education may have outpaced the cognitive, ethical, and legal frameworks needed to protect learning. But the more specific finding, the one that connects every section of this paper, is that governance failure distributes harm unevenly. Students with fewer resources face the widest access gaps. Students with non-standard writing styles face the highest false positive rates from detectors. Students who use AI to manage unequal workloads may be the ones who lose the most cognitively over time. And students without institutional support have the fewest tools to challenge how their data is used. The Edutopia retention study showed AI users scored higher on immediate quizzes but lower on delayed assessments, matching decades of research on desirable difficulties and the testing effect (Hamsa and Bastani; Bjork; Roediger and Karpicke). The MIT Media Lab EEG experiment showed ChatGPT users had 32 percent lower mental load and nearly halved brain connectivity in alpha and theta bands, providing neurological evidence that AI reduces the effortful processing learning depends on (Polytechnique Insights; Kosmyna). The evaluation of AI detection tools showed that no detector is reliable enough to serve as the sole basis for integrity enforcement, with false positive rates reaching 50 percent and disproportionate misclassification of non-native English speakers (Weber-Wulff et al.; Liang et al.; Fowler). AI may be undermining how students learn, how teachers assess, and how institutions enforce honesty. These findings are hypotheses supported by evidence, not proven causes.

Fixing these problems requires concrete action. First, schools should adopt process-based assessment: evaluating drafts, revisions, and in-class components, not relying on final submissions that AI easily produces. Second, teacher training must include mandatory AI literacy: understanding how generative AI works, when its use is appropriate, and how to design assignments that preserve cognitive effort. Third, federal and state legislatures should develop education-specific AI legislation addressing the cognitive, ethical, and equity dimensions documented here. Fourth, institutions should require students to declare any AI assistance, creating transparency without relying on flawed detection technology.

None of this argues for banning AI outright. Cognitive load theory says that offloading extraneous load can free working memory for the germane effort that builds schemas. AI used as a scaffold could preserve that effort rather than replace it. A student who has AI produces a rough outline and then has to critique it,

restructure it, and rewrite it is still doing the retrieval and generation that make learning stick. A student who has AI explains a concept a second way before attempting the problem set alone is still doing the work. What matters is not whether AI is used but who does the germane thinking. That is the calculator argument from Section 6.6, and it holds as long as the conditions it depends on are met: teacher scaffolding, redesigned assignments, and accountability for process rather than product. Most schools do not meet them yet.

This paper has limitations. The MIT EEG study involved only 54 participants. While its findings are consistent with broader research, longitudinal studies are needed to determine whether acute reductions in brain engagement translate into lasting changes over months or years. The neuroplasticity arguments in Section 5.2 are theoretical. No study has shown structural brain changes from AI dependence specifically. This paper focuses on secondary and undergraduate education in the United States. Cognitive effects could differ for younger children with developing neural systems, for graduate students with established schemas, or for students outside the United States. This is also a narrative review rather than a systematic one, and its source selection followed no pre-registered protocol. Future research should test whether structured AI use, with pedagogical scaffolding and process-based accountability, preserves the cognitive benefits of effortful learning while offering the accessibility advantages AI proponents describe. Until that research exists and is translated into policy, educational institutions should default to caution.

REFERENCES

- [1] Bill Text - AB-2013 Generative Artificial Intelligence: Training Data Transparency. [leginfo.legislature.ca.gov/faces/billTextClient.xhtml?bill_id=202320240AB2013](https://leginfo.ca.gov/faces/billTextClient.xhtml?bill_id=202320240AB2013). Accessed 31 Mar. 2026.
- [2] "AI Found to Boost Individual Creativity – at the Expense of Less Varied Content." *ScienceDaily*, 24 July 2024, www.sciencedaily.com/releases/2024/07/240712222127.htm.
- [3] Association of Texas Professional Educators. "ATPE News Summer 2023." *ATPE*, www.atpe.org/News-Media/Magazine/atpe-News-Summer-2023. Accessed 31 Mar. 2026.
- [4] Bjork, Robert A. "Memory and Metamemory Considerations in the Training of Human Beings." *Metacognition: Knowing About Knowing*, edited by J. Metcalfe and A. Shimamura, MIT Press, 1994, pp. 185–205.
- [5] Bjork, Elizabeth L., and Robert A. Bjork. "Making Things Hard on Yourself, but in a Good Way: Creating Desirable Difficulties to Enhance Learning." *Psychology and the Real World: Essays Illustrating Fundamental Contributions to Society*, edited by M. A. Gernsbacher et al., Worth Publishers, 2011, pp. 56–64.
- [6] Bjork, Robert A., and Elizabeth L. Bjork. "Desirable Difficulties in Theory and Practice." *Journal of Applied Research in Memory and Cognition*, vol. 9, no. 4, 2020, pp. 475–479.

- [7] Bjork, Robert A., John Dunlosky, and Nate Kornell. "Self-Regulated Learning: Beliefs, Techniques, and Illusions." *Annual Review of Psychology*, vol. 64, 2013, pp. 417–444.
- [8] "Brain Basics: Know Your Brain." *National Institute of Neurological Disorders and Stroke*, www.ninds.nih.gov/health-information/public-education/brain-basics/brain-basics-know-your-brain. Accessed 31 Mar. 2026.
- [9] BU CARES. "Artificial Intelligence and ChatGPT in Education." *YouTube*, uploaded by BU CARES, 3 May 2023, www.youtube.com/watch?v=Lw3Mcz8TNKg.
- [10] "California's AB 2013: Challenges and Opportunities in Generative AI Compliance." *Baker Botts*, 25 Nov. 2024, www.bakerbotts.com/thought-leadership/publications/2024/november/ca-ab-2013-genai-compliance.
- [11] *California's New AI Laws Focus on Training Data, Content Transparency*. Cooley, 16 Oct. 2024, www.cooley.com/news/insight/2024/2024-10-16-californias-new-ai-laws-focus-on-training-data-content-transparency.
- [12] Carr, Nicholas. *The Shallows: What the Internet Is Doing to Our Brains*. W. W. Norton and Company, 2010.
- [13] Chow, Andrew R. "ChatGPT May Be Eroding Critical Thinking Skills, According to a New MIT Study." *TIME*, 23 June 2025, time.com/7295195/ai-chatgpt-google-learning-school.
- [14] "Cognitive Load Theory." *The Decision Lab*, thedecisionlab.com/reference-guide/psychology/cognitive-load-theory.
- [15] Coursera Staff. "What Is ChatGPT? How It Works, How to Use It, and More." *Coursera*, 26 Jan. 2026, www.coursera.org/articles/chatgpt.
- [16] Danelladebel. "Governor Newsom Signs Bills to Combat Deepfake Election Content." *Governor of California*, 23 Sept. 2024, www.gov.ca.gov/2024/09/17/governor-newsom-signs-bills-to-combat-deepfake-election-content.
- [17] Eaton, Sarah Elaine. Work on artificial intelligence and academic integrity published in *International Journal for Educational Integrity*, vol. 19, 2023.
- [18] Tuomi, Ilkka. "What Counts as Evidence in AI & ED: Towards Science-for-Policy 3.0." *European Journal of Education Policy and Practice*, vol. 1, no. 1, 2025, pp. 1–31. doi.org/10.5117/EJEP2025.1.001.TUOM.
- [19] "Fluency Illusion: A Review on Influence of ChatGPT in Classroom Settings." *Information*, vol. 17, no. 3, 2026, article 299. mdpi.com/2078-2489/17/3/299.

- [20] Fowler, Geoffrey A. "Detecting AI May Be Impossible. That's a Big Problem for Teachers." *The Washington Post*, 2 June 2023.
- [21] Hollis, Chuck, et al. "California and Artificial Intelligence Watermarking Law." *Data Protection Report*, 20 Sept. 2024,
www.dataprotectionreport.com/2024/09/california-and-artificial-intelligence-watermarking-law.
- [22] "How AI Vaporizes Long-Term Learning." *Edutopia*, 28 Jan. 2025,
www.edutopia.org/video/how-ai-vaporizes-long-term-learning. Accessed 31 Mar. 2026.
- [23] "How Does Our Brain Form Creative and Original Ideas?" *ScienceDaily*, 15 Nov. 2015,
www.sciencedaily.com/releases/2015/11/151119104105.htm.
- [24] "How Will AI Affect the Global Workforce?" *Goldman Sachs*, 13 Aug. 2025,
www.goldmansachs.com/insights/articles/how-will-ai-affect-the-global-workforce.
- [25] Kahneman, Daniel. *Thinking, Fast and Slow*. Farrar, Straus and Giroux, 2011.
- [26] Karpicke, Jeffrey D., and Henry L. Roediger, III. "Repeated Retrieval During Learning Is the Key to Long-Term Retention." *Journal of Memory and Language*, vol. 57, no. 2, 2007, pp. 151–162.
- [27] Koriat, Asher, and Robert A. Bjork. "Illusions of Competence in Monitoring One's Knowledge During Study." *Journal of Experimental Psychology: Learning, Memory, and Cognition*, vol. 31, no. 2, 2005, pp. 187–194.
- [28] Kosmyna, Nataliya. "Your Brain on ChatGPT: Accumulation of Cognitive Debt When Using an AI Assistant for Essay Writing Task." *MIT Media Lab*,
www.media.mit.edu/publications/your-brain-on-chatgpt. Accessed 31 Mar. 2026.
- [29] Liang, Weixin, et al. "GPT Detectors Are Biased Against Non-Native English Writers." *Patterns*, vol. 4, no. 7, 2023.
- [30] Manfredi, Richard. "Regulating the Future: Eight Key Takeaways From California's SB 1047, Vetoed by Governor Newsom." *Gibson Dunn*, 5 Mar. 2025,
www.gibsondunn.com/regulating-the-future-eight-key-takeaways-from-californias-sb-1047-vetoed-by-governor-newsom.
- [31] Massachusetts Institute of Technology. "Generative AI + Education: Will Generative AI Transform Learning and Education." *YouTube*, uploaded by MIT, 18 Dec. 2023,
www.youtube.com/watch?v=yUlt7nLNNKE.
- [32] Mineo, Liz. "Is AI Dulling Our Minds?" *Harvard Gazette*, 13 Nov. 2025,
news.harvard.edu/gazette/story/2025/11/is-ai-dulling-our-minds.

- [33] Perkins, Mike. "Academic Integrity in the Age of Artificial Intelligence: A Study of AI Detection Tools." *Assessment and Evaluation in Higher Education*, 2023.
- [34] *Public Schools: AI Working Group - Professional Learning* (CA Dept of Education). www.cde.ca.gov/ci/pl/aiineducationworkgroup.asp. Accessed 31 Mar. 2026.
- [35] Reuell, Peter. "How The Brain Builds New Thoughts." *Harvard Gazette*, 5 Oct. 2015, news.harvard.edu/gazette/story/2015/10/how-the-brain-builds-new-thoughts.
- [36] Roediger, Henry L., III, and Jeffrey D. Karpicke. "Test-Enhanced Learning: Taking Memory Tests Improves Long-Term Retention." *Psychological Science*, vol. 17, no. 3, 2006, pp. 249–255.
- [37] Roediger, Henry L., III, and Jeffrey D. Karpicke. "The Power of Testing Memory: Basic Research and Implications for Educational Practice." *Perspectives on Psychological Science*, vol. 1, no. 3, 2006, pp. 181–210.
- [38] Slamecka, Norman J., and Peter Graf. "The Generation Effect: Delineation of a Phenomenon." *Journal of Experimental Psychology: Human Learning and Memory*, vol. 4, no. 6, 1978, pp. 592–604.
- [39] Sparrow, Betsy, et al. "Google Effects on Memory: Cognitive Consequences of Having Information at Our Fingertips." *Science*, vol. 333, no. 6043, 2011, pp. 776–778.
- [40] Sweller, John. "Cognitive Load During Problem Solving: Effects on Learning." *Cognitive Science*, vol. 12, no. 2, 1988, pp. 257–285.
- [41] Sweller, John. "Element Interactivity and Intrinsic, Extraneous, and Germane Cognitive Load." *Educational Psychology Review*, vol. 22, no. 2, 2010, pp. 123–138.
- [42] "The AI Paradox: Building Creativity to Protect Against AI." *ScienceDaily*, 24 May 2024, www.sciencedaily.com/releases/2024/05/240530132651.htm.
- [43] Turkle, Sherry. *Reclaiming Conversation: The Power of Talk in a Digital Age*. Penguin Press, 2015.
- [44] Weber-Wulff, Debora, et al. "Testing of Detection Tools for AI-Generated Text." *International Journal for Educational Integrity*, vol. 19, no. 1, 2023.
- [45] Withers, Bethany P. "California's AB 1033 Takes Effect: Navigating AI Training Data Transparency and Trade Secret Risk." *Goodwin*, 9 Mar. 2026, www.goodwinlaw.com/en/insights/publications/2026/01/alerts-other-industries-californias-ab-1033-takes-effect.

[46] Yan, Veronica X., Elizabeth Ligon Bjork, and Robert A. Bjork. "On the Difficulty of Mending Metacognitive Illusions: A Priori Theories, Fluency Effects, and Misattributions of the Interleaving Benefit." *Journal of Experimental Psychology: General*, vol. 145, no. 7, 2016, pp. 918–933.

ABOUT THE AUTHOR

Siddharth Roy is a high school junior at Washington High School in Fremont, California. He takes interest in political science and cognitive science, hoping to pursue further education at law school. He plans to study a STEM-related field in college, focusing on legal ethics and psychology.