

# Large Language Models Threaten Language’s Epistemic and Communicative Foundations

Anonymous ACL submission

## Abstract

Today, Large language models (LLMs) are reshaping the norms of human communication, sometimes decoupling words from genuine human thought. This transformation is deep, and undermines the trust and interpretive norms that were historically tied to authorship. We draw from linguistic philosophy and AI ethics to detail how large-scale text generation can induce semantic drift, erode accountability, and obfuscate intent and authorship. Our work here introduces conceptual frameworks including hybrid authorship graphs (modeling humans, LLMs, and texts in a provenance network), epistemic doppelgängers (LLM-generated texts that are indistinguishable from human-authored texts), and authorship entropy. We explore mechanisms such as “proof-of-interaction” authorship verification and educational reforms to restore confidence in language. While LLMs’ benefits are undeniable (broader access, increased fluency, automation, etc.), the upheavals they introduce to the linguistic landscape demand reckoning. This paper provides a conceptual lens to chart these changes.

## 1 Introduction

*“Last year’s words belong to last year’s language  
And next year’s words await another voice”*

Language has been the keystone of our communication and thought for thousands of years. From ancient cuneiform tablets to modern digital platforms, people have relied on written language as a store of facts, beliefs, and ideas. Underlying this tradition is a widespread assumption: that any text reflects a human mind, shaped by cognitive processes and linked to specific authors. Thus, language has been an expression of human intention, demanding both attention from the reader, and accountability from the author (Winograd, 1972; Bender and Koller, 2020). Over our long evolutionary trajectory, these assumptions have steadily held true, and are now woven into the very fabric of how we understand and interpret language.

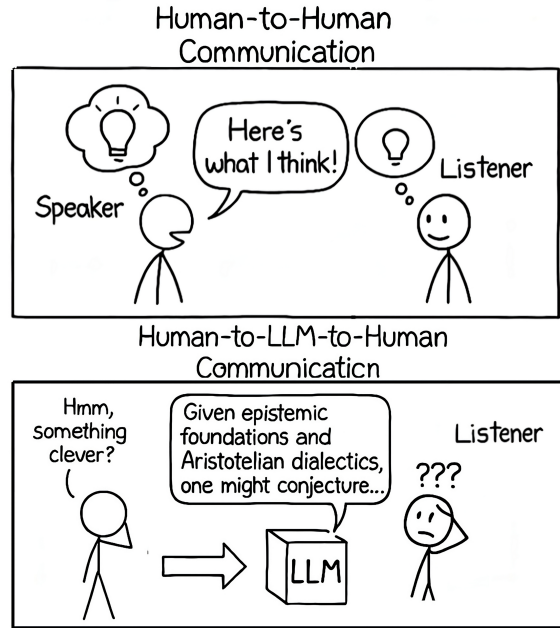


Figure 1: LLMs introduce a shift in communicative dynamics. Traditionally, human-to-human communication directly conveys intentional thought from speaker to listener (top). But when mediated by LLMs, language can lose direct intentional grounding, resulting in messages disconnected from the speaker’s original intent and confusing the listener (bottom).

But the swift ascent of large language models (LLMs) over the past five years has begun to fundamentally reconfigure this relationship. Trained on internet-scale corpora and with representational flexibility from billions of parameters (e.g., GPT (Brown et al., 2020), PaLM (Chowdhery et al., 2022), LLaMA (Touvron et al., 2023)) or DeepSeek (DeepSeek-AI, 2024), these can generate well-polished and coherent text with minimal guidance. They can replicate stylistic nuances, rhetoric, and emotional tones (Schick et al., 2021; Solaiman et al., 2019), which were attributed solely to human creativity till very recently.

On the one hand, these capabilities are surprising

and extraordinary, suggesting potential for a new cognitive revolution in AI. On the other, this transformation weakens the link tying text to a human mind. Teachers worry that student essays reflect an LLM’s fluency rather than the messy traces of a student’s thoughts (Cotton et al., 2023; Zhou et al., 2023). Researchers question whether a paper reflects a scholar’s insights or a model’s reassembly of existing content (Zellers et al., 2019). Even everyday exchanges: emails, tweets, and blogs might stem from digital processes rather than human voices (Duarte et al., 2022; Weidinger et al., 2021). In this sense, it would be ironic if LLMs, instead of illuminating and edifying human language communication, lead to its devolution.

The broader NLP community, as creators of LLM technologies, bears a direct responsibility for their societal implications (Gabriel, 2020; Bender et al., 2021). By ignoring the epistemic and ethical consequences of these systems, we risk having text lose its role as an indicator of human intent and thought (Floridi and Chiriatti, 2020; Raji et al., 2022). This can fundamentally degrade how we conduct discourse, value expertise, and maintain trust (O’Neil, 2016; Zuboff, 2019). This paper grapples with this tension between positive applications of LLMs (Team, 2022; Hutchinson et al., 2023; Miller, 2019) and the challenge LLMs pose to language’s role in human thought. Section 2 explores language’s philosophical aspects like intentionality and authorship. Section 3 identifies two key issues stemming from widespread LLM text, namely semantic drift and erosion in trust. In Section 4, we describe possible approaches to re-establish accountability and measure semantic drift using methods from NLP, cryptography, and HCI. Section 5 looks at social implications and possible responses in education, scholarship, etc. Section 6 argues that approaches like watermarking do not address core issues. Section 7 explores rethinking ideas about language and authorship with ideas on human-AI collaboration, educational changes and human-only publishing spaces. We briefly summarize alternative perspectives and arguments in Section 8. We conclude with a reflection on the need to maintain language’s cognitive and epistemic roles in the future.

## 2 Language, Authorship, and Meaning

Language’s function has been debated for centuries, from Plato’s dialogues on rhetoric to modern ana-

lytic philosophy (Plato, 1997; Wittgenstein, 1953). While language is commonly viewed as an information channel for transmitting information, linguists argue that language is a *communal sense-making act*. It has been closely linked to intention, context, and the ability to hold speakers accountable (Searle, 1969; Austin, 1975; Floridi, 2013).

**Language as an Intentional Act:** Searle’s speech-act theory (Searle, 1969) and Austin’s work on performativity (Austin, 1975) argue that language is not just a conduit for information transfer, but also enacts intentions. To *say* something is often to *do* something: to promise, question, declare, for example. The force of an utterance depends on the speaker’s agency and recognition of those intentions by a listener (Grice, 1975). This has been fundamental to authorship, particularly in academic and legal discourse, where a text is an intellectual act tied to its creator’s identity and responsibility (Dworkin, 1996). Even when ghostwriters were traditionally involved, the text ultimately reflected a coherent cognitive source (Foucault, 1984; Chartier, 1994; Sperber and Wilson, 1986).

The rise of LLM-generated text disrupts these frameworks (Floridi and Chiriatti, 2020). Do AI-generated documents bear the same weight without deliberate intent? This uncertainty can raise questions about authorship, intellectual property, and trust, especially for scientific or legal text (Raji et al., 2022; Huang and Rust, 2021; van Dis et al., 2023). Authorship has historically entailed a social contract: a published text can be challenged or critiqued, holding its human creator(s) responsible for factual or ethical shortcomings (Woodmansee, 1994; Chartier, 1994). But with LLM-authored text, accountability becomes diffused: does it lie with the prompter, the model trainer, or the dataset creator? This diffusion of responsibility strains traditional legal and academic norms (Hacker et al., 2023; Mittelstadt and Floridi, 2016; Kosseff, 2019). As the intent and accountability of text becomes murky, its meaningfulness and trustworthiness can become suspect too.

**Language as a Cognitive Interface:** Beyond communication, language shapes cognition and our capacity to abstract and solve problems (Clark and Chalmers, 1998; Vygotsky, 1978; Whorf, 1956). It is often considered an “interface” to thought. Research in child cognitive development suggests that engagement with language enables reasoning, cognitive flexibility, and problem-solving (Tomasello,

2003; Bruner, 1983; Lakoff and Johnson, 1980). While some argue that LLMs function as cognitive enhancers (Clark and Chalmers, 1998; Warwick, 2003), others caution that reliance on LLM-driven generation can lead to reduced cognitive engagement (Nichols, 2021; Carr, 2010). In particular, LLM-driven writing and summarization has raised concerns about cognitive deskilling (Carr, 2011; Lai and Viering, 2022). Studies show that composition itself is integral to thinking, forcing individuals to clarify ambiguity, structure arguments, and synthesize knowledge (Kellogg, 2008; Galbraith, 1999). Further, LLM-generated summarization risks eroding cognitive *effort*, like digital offloading has been shown to reduce critical engagement (Sparrow et al., 2011; Nichols, 2021).

### Chain-of-Thought and Cognitive Parallels:

The emergence of chain-of-thought prompting (CoT) (Wei et al., 2022) represents a major shift in LLM problem-solving. By externalizing logical steps, CoT compensates for the depth limitations of transformer architectures (Vaswani et al., 2017; Yao et al., 2023). This mirrors how humans articulate thoughts through language, diagrams, or writing to enhance problem-solving (Clark and Chalmers, 1998; Menary, 2010). Beyond computational efficiency, CoT also bears parallels to how externalizing reasoning through symbols has been linked to the expansion of human intelligence (Deacon, 1997; Dor, 2015). If language enabled humans to extend cognition beyond individual memory, CoT might mark a similar milestone in LLM development. Whether CoT augments human intelligence or leads intellectual complacency will depend on how societies integrate LLM-based thought and reasoning in education, work, and decision-making.

## 3 The Crisis of Language

Given these philosophical foundations, the increasing role of LLMs ruptures the linguistic landscape through two forces: (1) *Semantic Drift & Model Collapse*, the idea that the influx of AI-generated text can shift the distribution and meaning of language, and lead to compounding errors; and (2) *Eroding Epistemic Trust* in text, epitomized by what we term epistemic doppelgängers (LLM outputs that are indistinguishable from human outputs). We also suggest a metric, *authorship entropy*, to represent the uncertainty about the origin of a text.

### 3.1 Semantic Drift and Model Collapse

Semantic drift refers to changes in language usage and meaning over time, reflecting cultural and social evolution. However, large-scale LLM-generated content can accelerate or redirect semantic change. For example, they might reinforce common phrases while underrepresenting less frequent expressions (Raji et al., 2022). LLMs trained on text that partially includes their own synthetic outputs can experience compounding errors. Repeated assimilation of AI-generated text leads to a shift away from organic language distributions (Shumailov et al., 2023; Carlini et al., 2023). Over time, certain stylistic artifacts become over-represented, contributing to *model collapse* (Menick et al., 2022), where the system’s expressive range narrows towards the mean of what LLMs produce. While prior work has studied distribution shift in active learning (Blitzer et al., 2007), LLM self-ingestion is a novel feedback loop.

Semantic drift and model collapse are intertwined in a hybrid loop of human and LLM language production (Bommasani et al., 2021). Human writing feeds the training of LLMs, and LLM outputs in turn influence human writing and future training data. Without intervention, this loop may have an unintended equilibrium: language evolution not driven by human innovation, but by statistical characteristics of LLM generated text. The outcome can be a loss of semantic clarity (as meanings shift in unpredictable ways) and reduced linguistic innovation (Weidinger et al., 2021; Bender et al., 2021). If significant portions of text that we read comes to be machine-generated, we should also ask if this can deteriorate our cognitive diversity at a societal level.

### 3.2 Eroding Trust and Accountability

Texts have long been vehicles for accountability. Historically, authors were usually identifiable, and could be praised, critiqued, or legally challenged based on their claims (Woodmansee, 1994; Chartier, 1994). LLM-generated content fragments this chain of responsibility. This has significant implications for defamation suits and retraction practices for erroneous statements (Kosseff, 2019). Even more pressingly, online disinformation campaigns leveraging AI threaten political discourse, as citizenry loses clarity on who authors the narratives shaping public opinion (Weidinger et al., 2021; Chesney and Citron, 2019).

Empirical evidence is growing that synthetic text can fuel coordinated misinformation. Recent experiments demonstrate that AI-generated content is capable of crafting coherent yet deceptive social media campaigns, blurring the line between authentic and automated discourse (Zellers et al., 2019). Furthermore, large-scale language models have been observed to inadvertently plagiarize, amplify biases, and perpetuate stereotypes from their training data (Bender et al., 2021; Hovy and Spruit, 2016; Carlini et al., 2023). Without robust authorship signals or provenance tracking, verifying source credibility becomes increasingly challenging, raising concerns about accountability in digital information ecosystems (Gehrmann et al., 2019; Kirchenbauer et al., 2023).

### 3.3 Epistemic Doppelgängers and Authorship Entropy

LLMs can produce text nearly indistinguishable from human writing. We refer to such outputs as epistemic doppelgängers: texts that impersonate human authorship so convincingly they can fool not only casual readers, but editors, teachers, and even domain experts. As with a human doppelgänger, the deception isn’t necessarily malicious, but its uncanniness can be destabilizing. GPT-generated news articles are often rated as more trustworthy than authentic ones (Zellers et al., 2019), and even the best AI detectors rarely surpass 70% accuracy. Worse, detection systems are often only effective when closely matched to the model they’re trying to catch, making them vulnerable to fine-tuning, or strategic prompting. In short, epistemic doppelgängers erode the assumption that a well-formed sentence signals a human mind.

This epistemic ambiguity leads us to what we call *authorship entropy*, a measure of uncertainty of text authorship. In a world where all documents are confidently human-written, authorship entropy is low: the provenance of text is legible, even if anonymous. But in an AI-saturated ecosystem, the space of plausible authors expands. By modeling this uncertainty as a probability distribution and applying Shannon entropy, we can quantify how “foggy” the authorship landscape is. Rising authorship entropy destabilizes trust: people may become suspicious of legitimate texts, or indifferent to provenance altogether. It weakens accountability: if we don’t know who wrote something, we can’t assign responsibility. Also, AI authors, by definition, evade moral blame.

## 4 Technical Foundations: Authenticating Authorship & Quantifying Drift

Our discussion thus far has been primarily conceptual. In the next section, we explore technical interventions aimed at reclaiming human accountability and reducing authorship entropy.

### 4.1 Author Graphs & Proof-of-Interaction

A possible direction is embedding provenance and requirement of human interaction in the text generation process itself. For example, a hybrid authorship graph can represent relationships between human users, LLMs and texts that they generate or indirectly influence. To explain, a document’s node might have edges from an LLM node (if an AI drafted it) and a human node (who guided or edited it). If an AI’s training data included that document, an edge from the document back to the AI node (“trains”) can be included, forming a cyclic network of influence. Figure 2 shows an example of such a graph. Such explicit representations of provenance and sources can provide grounding to enforce downstream accountability.

A practical implementation of this can be through *Proof-of-Interaction* (PoI) mechanisms that ensure that a human was substantially involved in creating a text. For instance, an editor can sign off on an AI-generated passage after verifying it, or a platform can require that any AI assistance be logged and attested. Some have proposed protocols where documents carry embedded metadata or hashes that link to records of the human-AI collaboration that produced them. If a document cannot present such proof-of-interaction, it might not be trusted for certain uses.

Building on blockchain-inspired ideas (Narayanan et al., 2016) and prior research on authorship verification (Stamatatos, 2009; Layton and Watters, 2020), we propose a *process-based* approach to demonstrate genuine human involvement. Let  $\mathcal{T}$  denote a text document, and let  $\mathcal{C} = \{C_1, C_2, \dots, C_n\}$  denote the set of discrete human contributions (e.g., edits, approvals, interventions). Then we define its proof-of-interaction score for the text as:

$$\text{PoI}(\mathcal{T}) = f(\mathcal{C}, \text{HashChain}).$$

Here, HashChain refers to a cryptographically secure chain of hashes that records the provenance of these edits.  $f(\cdot)$  outputs a scalar or structured PoI score or certificate. This function can be designed

to validate sufficiency (e.g., a minimum amount of human interaction occurred), or provide a verifiable attestation (e.g., a digitally signed summary).

Newer works in *human-in-the-loop* text generation have proposed storing metadata that logs partial revisions (Cecchi and Babkin, 2024; Kang et al., 2024). Merging these ideas with zero-knowledge proofs can preserve user privacy. While this does not solve all issues (e.g., adversaries can simply simulate keystrokes or partial edits), such a system raises the cost of deception and provides an auditable trail of interaction.

This idea also aligns with Chain-of-thought (Wei et al., 2022; Kojima et al., 2022) oversight, where an LLM seeks human verification for intermediate reasoning. Some developers propose user-audited chains-of-thought, letting humans see exactly which steps an LLM took (Wang et al., 2022). Future research can unify chain-of-thought logs with proof-of-work, offering a secure record of how text was generated. This can clarify the roles of LLMs and humans in authoring text, although this approach may present challenges in preserving user privacy (Khowaja et al., 2023; Abadi et al., 2016; Glymour et al., 2023). These changes will necessarily introduce friction, and may be tedious for users. However, a proof-of-interaction system can ensure that every text is connected to at least one human via a “verified” edge. This can maintain the principle that for any published text, one can point to a human accountable for it.

## 4.2 Metrics for Semantic Drift

Verifying authorship addresses who wrote the text. But we should also ask what is being written. We propose tracking language changes by defining metrics for semantic drift and linguistic diversity, comparing human-authored and LLM-generated text periodically. By measuring shifts in word frequencies, syntax, or topics, we can identify drift if metaphoric language or dialectal terms decrease while AI-generated phrases increase.

It is also worth monitoring model-internal drift: how successive generations of LLMs differ when trained on data that includes prior LLMs’ outputs. If  $P_{human}$  and  $P_{LLM_\theta}$  indicate the probability distributions of language utterances at a discrete time step  $t$ , and if  $\alpha$  denotes the proportion of LLM generated data, then the distribution of training data that will be used to train the next iteration of the

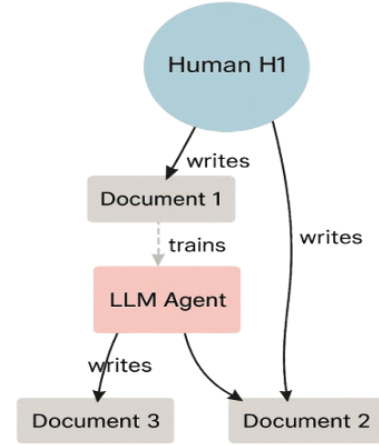


Figure 2: An illustrative hybrid authorship graph, representing provenance and interactions between human agents, an LLM agent, and texts. In this example, Human H1 writes Document 1, which is later used in training the LLM. Document 2 is co-authored by H1 and the LLM (perhaps H1 edited text generated by the LLM). Document 3 is authored solely by the LLM.

LLMs,  $P_{LLM}^{t+1}$ , is given by:

$$P_{mix}^{(t)}(\mathbf{x}) = (1 - \alpha)P_{human}(\mathbf{x}) + \alpha P_{LLM_\theta}^{(t)}(\mathbf{x})$$

as LLM outputs re-enter training data (Shumailov et al., 2023). If  $P_{mix}^{(t)}$  increasingly diverges from  $P_{human}$ , the model parameters  $\theta$  risk converging to a subspace that fails to capture the richness of actual human language patterns. Further computational or theoretical insights might be found in work on catastrophic forgetting (Kirkpatrick et al., 2017) and domain shift (Ganin et al., 2016). While the notion is not new (prior studies on machine-in-the-loop domain adaptation raise similar concerns (Ruder, 2019)), our contribution is to highlight how large volumes of synthetic text can nudge language distributions away from natural usage. We propose coupling distributional metrics (e.g., KL divergence) with textual diversity indices to monitor linguistic homogenization.

Empirically testing these metrics on real corpora that blend human and AI-generated text remains a priority for future research. Experiments with smaller LLMs (Carlini et al., 2023; Menick et al., 2022) suggest that repeated synthetic ingestion amplifies shallow lexical patterns. In Figure 3, we plot the JS divergence between unigram distributions in a base human-authored corpus, and a GPT2-base model that is iteratively re-trained on its own sampled data. We note a clear and increasing semantic drift with an increasing number of steps.

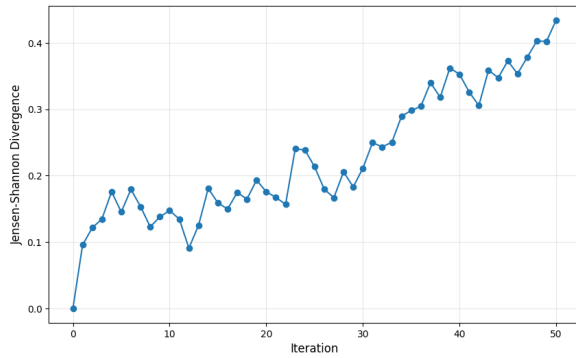


Figure 3: Semantic Drift in Synthetic Text: The plot shows the Jensen-Shannon divergence between a base human-authored text distribution and iteratively drifted synthetic text distributions. As synthetic text is repeatedly generated and reintroduced into training data, the divergence increases, illustrating the risk of semantic drift and potential loss of linguistic diversity over time. At each iteration, the GPT2 model is fine-tuned on text sampled from the GPT2 model in the previous step

In summary, technical tools for authorship authentication (like detection, watermarking, and proof-of-interaction) and for linguistic monitoring (measuring drift and diversity) will be crucial parts of the solution. However, they are not panaceas. Detection can be evaded; verification can be cumbersome; drift metrics can tell us there’s a problem but not fix it. We now turn to the human and institutional side of the response: how education, policy, and cultural norms should adapt in order to safeguard the epistemic foundations of language.

## 5 Societal Implications

While LLMs can improve productivity and streamline workflows, their widespread adoption raises concerns about academic integrity, scholarly publishing, professional hiring, and public discourse (Coglianese and Lehr, 2017; Cotton et al., 2023; van Dis et al., 2023).

**Education** LLMs are increasingly used as study aids, enhancing access for learners with different language skills (Khan et al., 2023; Chaudhuri et al., 2021; Xu et al., 2022; Luckin et al., 2023). However, relying on LLMs for tasks like programming or essay writing can weaken essential skills: algorithmic thinking, structured argumentation, and creativity (Cotton et al., 2023; Perkins and Salomon, 1989). To address this, some schools use real-time or proctored writing tasks or oral exams to ensure understanding (Lund and Wang, 2023). Our chain-of-thought synergy (Sec. 7) encourages student par-

ticipation in the generative process but is difficult to scale. Education should focus on teaching skills that AI cannot easily replace.

**Academic Scholarship** The academic ecosystem assumes text is a reflection of an author’s intellect. Automated text generation challenges this, raising concerns over AI authorship and scholarly contributions (Willis and Williams, 2023). In the short term, LLMs risk hallucinations (Ji et al., 2023) and plagiarism (van Dis et al., 2023). More seriously, they can flood peer reviews and obscure genuine innovation. Ideally though, LLMs should boost research productivity and accelerate scientific progress.

**Professional Settings** In many industries, cover letters, writing samples, and portfolio websites are used to gauge candidates’ communication skills and expertise (Sternberg and Williams, 1997). LLM tools now make it easy to create polished but shallow applications, complicating hiring managers’ ability to assess true abilities. Some organizations are turning to live assessments like real-time writing tests or structured panel interviews (Koch et al., 2015; Levashina et al., 2014). But these can be difficult for introverts, non-native speakers, or candidates who do better with written communication (van Tubergen and Kalmijn, 2014; Hu et al., 2020). Managing this requires a delicate dance between fairness and authenticity.

**The Public Sphere** LLMs are reshaping public discourse via AI-generated content, sparking concerns about amplification and distortion (Zellers et al., 2019; Ferrara, 2020). Disinformation campaigns exploit AI’s capability to produce misleading content, drowning out authentic voices and confusing public understanding. Although detection methods advance, the adversarial landscape perpetuates a constant arms race (see Section 6). Conversely, LLMs present opportunities to democratize communication by reducing barriers for individuals with limited writing skills, disabilities, or those who are non-native speakers (Paritosh et al., 2022; Xu et al., 2022).

In all of these domains just discussed, a common thread is that trust is threatened by automatic text generation from LLMs. As trust erodes, institutions will react by imposing stricter verification, leading to friction, surveillance, or cynicism. The challenge is developing norms that preserve the value of human contribution and ensure transparency.

## 6 Labels & Classifiers wont save us

Proposals such as watermarking and policy bans, while helpful in the short term, offer only superficial remedies (Gehrmann et al., 2019; Zellers et al., 2019; Papernot et al., 2016).

**Watermarking and Detection Arms Races** Watermarking remains fragile against adversarial attacks like paraphrasing (Kirchenbauer et al., 2023). Detection classifiers also struggle with robustness as LLMs adapt and human post-editing obfuscates machine origins (Holtzman et al., 2020). This creates a resource-intensive cat-and-mouse dynamic without stable solutions (Gallagher et al., 2023). More critically, these methods do not resolve attribution, leaving ethical and legal questions unanswered (Authors, 2023; Devinney, 2023).

**Policy Bans and their Limitations** Bans on AI-assisted writing are unenforceable due to weak detection and strong incentives for LLM use, effectively becoming honor systems (Devinney, 2023). Lagging legislation creates a fragmented regulatory landscape (Hacker et al., 2023), and the global nature of digital communication allows easy circumvention of local policies (Katyal and Epps, 2022), failing to address authorship and accountability.

### Neglecting the Deeper Interpretive Question

Fundamentally, current measures do not restore a discernible *human* presence. If language’s epistemic function relies on text as an intentional artifact, superficial labeling fails to reattach text to a mind (Bender and Koller, 2020). It still leaves us with a fundamentally ambiguous communicative landscape, and risks putting us in an era of permanent ambiguity in textual interpretation.

## 7 Rethinking Language & Authorship

We propose recalibrating the role of LLMs in language by leveraging their benefits while preserving human traits like intentionality, accountability, and diversity of thought. This requires an interdisciplinary approach involving NLP, cognitive science, ethics, law, and education. In this section, we refine previous suggestions and introduce ideas on human-AI collaboration frameworks, governance, and cultural appreciation of human-only work (Mittelstadt and Floridi, 2016; Floridi, 2019).

### 7.1 Chain-of-Thought with Human Oversight

As previously mentioned in Section 4, we advocate for embedding AI within a structured chain-of-

thought framework that requires human oversight at key decision points (Wei et al., 2022). In this paradigm, LLMs can generate partial outlines, intermediate arguments, or suggested revisions, but finalization has to be authenticated by a human user after consideration. By logging human-AI interactions through an auditable chain (with appropriate privacy safeguards), this method establishes a transparent record that delineates AI-generated content from human refinement, addressing concerns about accountability and intellectual ownership (Wang et al., 2022; Christiano, 2022).

### 7.2 Governance & Collaborative Policy

Governance for LLM-usage has to be a negotiation (between policymakers, educators, and user communities, etc.) for it to work, rather than a prescription (Floridi and Chiriatti, 2020; Hacker et al., 2023). Several directions seem promising. First, AI contribution statements, similar to conflict-of-interest disclosures, can prompt authors to declare the extent and nature of LLM involvement (Devinney, 2023). Second, labeling protocols for governmental or legal texts can introduce metadata or disclaimers to flag LLM-generated contents (Union, 2023). Also, ethical AI certification programs, modeled on data protection seals, can help LLM developers conform with regulations such as the EU AI Act (Union, 2023).

### 7.3 Educational Reforms and Cultural Shifts

To prevent cognitive deskilling, education has to pivot and adapt. Assignments will have to adapt to the inevitability of the use of LLMs for drafting, but can require students to justify revisions and incentivize peer engagement (Lai and Viering, 2022). Assessments like live problem-solving and debates will have to focus on substance over polish (Paul and Elder, 2007; Lipman, 2003; Chi and Wylie, 2014; Freeman et al., 2014). Finally, students should be encouraged to play with LLMs, and taught to interrogate them. AI literacy should be a form of critical literacy, where students learn not simply how to use LLMs, but when and when not to (Bowman and Reeves, 2015).

### 7.4 Human-Only Publishing Spaces

A potential direction is establishing “human-only” publishing spaces: media outlets, or creative communities that employ verification measures (such as mechanisms like proof-of-work logs) to ensure

that any content reflects considerable human intellectual effort and creativity. These spaces would offer a parallel track for those who value direct human expression. Some journals already forbid undisclosed AI collaboration for final submissions (Board, 2023). A fiction community might pride itself on entirely human-crafted stories. Like organic labels in food, these spaces can serve audiences that value authentic human expression—akin to “slow food” movements in a fast-food world (Petrini, 2001). Over time, such enclaves can serve as a ‘control group’, preserving the standards and norms associated with human authorship.

## 8 Alternate Views

We have emphasized how LLMs may erode traditional assumptions about language and authorship. But many scholars have a more sanguine outlook, and do not frame LLM-driven automation as a threat to language. For balance, we summarize key arguments in these perspectives.

**Democratization and Accessibility** LLMs can enable non-native speakers and individuals with disabilities to participate in public discourse (Norvig and Thrun, 2009; Ogawa et al., 2022). By automating surface-level writing concerns, these tools allow users to focus on substantive ideas. For example, spell-checkers, were once controversial too (Felton, 2023; Christiansen, 2021). Additionally, LLMs increase accessibility for users with impairments (Wagner et al., 2020), reframing ‘linguistic inclusivity’ as a positive evolution.

**Accelerated Knowledge Dissemination** Summarization tools help researchers digest literature efficiently (Fabbri et al., 2022; Sharma et al., 2022), and multilingual translation expands access to specialized knowledge (Fan et al., 2021; Artetxe and Schwenk, 2019). With editorial oversight, these outputs can enhance comprehension without compromising reliability (Szegedy et al., 2022). Advocates argue that with transparency, LLMs can strengthen epistemic ecosystems rather than harm them (Diakopoulos, 2016).

**Evolving Norms of Collaboration** In many domains, collaborative authorship is standard (technical manuals, corporate reports, etc.), which rarely reflect a single voice (Darics, 2020; Leonard and Noonan, 2020). In this context, LLMs are seen as additional collaborators (Krause et al., 2022; Dinan et al., 2022). Rather than undermining authorship,

they may shift workflows, with new roles emerging for human editors and fact-checkers (Eisenstein and McNamara, 2023; Roose and Sullivan, 2023).

## Empirical Evidence of Positive Outcomes

Some studies suggest that, when used responsibly, LLMs can enhance writing without weakening critical thinking. They support non-native and novice writers in building fluency (Lee et al., 2022; Laubrock et al., 2022). In collaborative environments, there is evidence that AI systems help clarity, and can identify redundancy (Yosinski et al., 2023; Rahimi et al., 2021).

Broadly, these perspectives argue that LLMs are not existential threats to the integrity of language, and that a ‘crisis’ of authorship is neither new nor uniquely AI-induced. Rather, this is a natural evolution in how we produce and share ideas. Possibly, when questions about the origin and intent of a text fade, newer and better-suited norms can emerge in the linguistic landscape to replace them.

## 9 Conclusion & Reflection

LLMs are here to stay. Yet, they challenge the epistemic and communicative foundations of language by decoupling text from human intent. While automated or impersonal writing is not new (for example, memos or legal boilerplates), LLMs amplify this disconnection at an unseen scale. This shift raises several prickly questions about intent, accountability and trust. Any solutions here must preserve the role of language in human agency and accountability. The key will be to assert human presence in language: whether through cryptographic attestations, proof-of-interaction, or new cultural norms.

The future of language doesn’t hinge as much on building better models, but on what we choose to protect. If we can collaborate to implement verifiable authorship and enforceable audit mechanisms, we may harness LLMs’ advantages without surrendering the uniquely human dimensions of language. But a failure to act can lead to communication devoid of color, reeking with hollow expressions, and diluted cognitive depth in people. The path to hell is famously paved with good intentions. Still, through ingenuity and foresight, LLM innovations could be steered toward enhancing human creativity, rather than eroding the intellectual bedrock which is its basis.

## 698 Limitations

699 By design, this paper is more diagnostic than pre-  
700 scriptive. We introduce conceptual tools like epis-  
701 temic doppelgängers and authorship entropy to  
702 make sense of the shifting linguistic terrain. How-  
703 ever, many of these constructs remain speculative  
704 without empirical grounding. Nor do we pretend  
705 that quantitative metrics alone can capture the con-  
706 sequences of LLM saturation. What we offer is a  
707 framework to think with, not a solution to deploy.

708 Second, some of our proposals (such proof-of-  
709 interaction logs, and human-only publishing en-  
710 claves) are challenging to implement and reify.  
711 They require infrastructure, extensive cooperation,  
712 and cultural shifts that may not be welcome. We  
713 should also acknowledge a significant tension: the  
714 paper champions the pre-eminence of human inten-  
715 tion in language, but we do not wish to gatekeep  
716 expression or discourage the increasingly creative  
717 and original uses of LLMs. The challenge is to  
718 protect the epistemic integrity of language with-  
719 out devolving into ‘purity tests’. To truly solve this  
720 challenge will requires contending with lived social  
721 realities of people, not just technical design.

## 722 AI Use Acknowledgment

723 In this work, we acknowledge the use of AI assis-  
724 tance in the following cases in accordance with the  
725 ACL Policy on AI Writing Assistance: assistance  
726 with literature search and review, proof-reading  
727 and refining the language of the paper, and ana-  
728 lytical code. We utilized AI tools for searching  
729 for relevant literature, help with writing code for  
730 analysis related to Figure 3, and code for diagram  
731 generation for Figure 2.

## 732 References

- 733 Martin Abadi, Andy Chu, Ian Goodfellow, H. Bren-  
734 dan McMahan, Ilya Mironov, Kunal Talwar, and  
735 Li Zhang. 2016. [Deep learning with differential pri-  
736 vacy](#). In *Proceedings of the 2016 ACM SIGSAC Con-  
737 ference on Computer and Communications Security  
738 (CCS)*, pages 308–318.
- 739 Mikel Artetxe and Holger Schwenk. 2019. [Mas-  
740 sively multilingual sentence embeddings for zero-  
741 shot cross-lingual transfer and beyond](#). *Transactions  
742 of the Association for Computational Linguistics*,  
743 7:597–610.
- 744 J. L. Austin. 1975. *How to Do Things with Words*,  
745 second edition. Harvard University Press.

- Anonymous Authors. 2023. [Gpt and the plagiarism  
problem: Assessing ai’s impact on academic integrity](#).  
*Journal of Ethics in AI*.
- Emily M. Bender, Timnit Gebru, Angelina McMillan-  
Major, and Shmargaret Shmitchell. 2021. [On the  
dangers of stochastic parrots: Can language models  
be too big?](#) In *Proceedings of the 2021 ACM Confer-  
ence on Fairness, Accountability, and Transparency  
(FAccT)*, pages 610–623.
- Emily M. Bender and Alexander Koller. 2020. [Climbing  
towards NLU: On meaning, form, and understanding  
in the age of data](#). *Proceedings of the 58th Annual  
Meeting of the Association for Computational Lin-  
guistics (ACL)*, pages 5185–5198.
- John Blitzer, Mark Dredze, and Fernando Pereira. 2007. [Biographies, bollywood, boom-boxes, and blenders:  
Domain adaptation for sentiment classification](#). In  
*Proceedings of the 45th Annual Meeting of the Asso-  
ciation for Computational Linguistics (ACL)*, pages  
440–447.
- Nature Editorial Board. 2023. [Ai and authorship: Na-  
ture’s policy on undisclosed ai collaboration](#).
- Rishi Bommasani, Drew A Hudson, Ehsan Adeli, Russ  
Altman, Simran Arora, Sydney von Arx, Michael S  
Bernstein, Jeannette Bohg, Antoine Bosselut, Emma  
Brunskill, and 1 others. 2021. On the opportunities  
and risks of foundation models. *arXiv e-prints*, pages  
arXiv–2108.
- Nicholas A. Bowman and Anne E. Reeves. 2015. [Re-  
thinking media literacy in the age of algorithmic cu-  
ration](#). *Journal of Media Education*, 6(3):14–24.
- Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie  
Subbiah, Jared D. Kaplan, Prafulla Dhariwal, Arvind  
Neelakantan, Pranav Shyam, Girish Sastry, Amanda  
Askell, Sandhini Agarwal, Ariel Herbert-Voss,  
Gretchen Krueger, Tom Henighan, Rewon Child,  
Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu,  
Clemens Winter, and 12 others. 2020. [Language  
models are few-shot learners](#). *Advances in Neural  
Information Processing Systems (NeurIPS)*, 33:1877–  
1901.
- Jerome Bruner. 1983. *Child’s Talk: Learning to Use  
Language*. W. W. Norton Company.
- Nicholas Carlini, Matthew Jagielski, Chiyuan Zhang,  
Katherine Lee, Christopher A. Choquette-Choo, Ja-  
cob Imber, Andreas Terzis, Nicholas Frosst, Ilya  
Mironov, Vasisht Duddu, and 1 others. 2023. [Poison-  
ing web-scale datasets is practical](#). *arXiv preprint  
arXiv:2305.00956*.
- Nicholas Carr. 2010. *The Shallows: What the Internet  
Is Doing to Our Brains*. W.W. Norton Company.
- Nicholas Carr. 2011. *The Shallows: What the Internet  
is Doing to Our Brains*. W. W. Norton & Company.

|     |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
|-----|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 799 | Lucas Cecchi and Petr Babkin. 2024. <a href="#">Reportgpt: Human-in-the-loop verifiable table-to-text generation</a> . In <i>Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing: Industry Track</i> , pages 529–537.                                                                                                                                                                  | 852 |
| 800 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 853 |
| 801 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 802 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 803 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 804 | Roger Chartier. 1994. <i>The Order of Books: Readers, Authors, and Libraries in Europe Between the Fourteenth and Eighteenth Centuries</i> . Stanford University Press.                                                                                                                                                                                                                                                   | 856 |
| 805 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 857 |
| 806 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 858 |
| 807 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 808 | Soham Chaudhuri, Varun Kumar, and Marti Hearst. 2021. <a href="#">Ai-powered writing assistants: Enhancing learning and creativity</a> . In <i>Proceedings of the Conference on Educational Data Mining (EDM)</i> , pages 344–355.                                                                                                                                                                                        | 859 |
| 809 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 860 |
| 810 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 861 |
| 811 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 862 |
| 812 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 813 | Bobby Chesney and Danielle Keats Citron. 2019. <a href="#">Deepfakes and the new disinformation war: The coming age of post-truth geopolitics</a> . <i>Foreign Affairs</i> , 98(1):147–155.                                                                                                                                                                                                                               | 866 |
| 814 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 867 |
| 815 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 868 |
| 816 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 869 |
| 817 | Micheline T. H. Chi and Rachel Wylie. 2014. <a href="#">The icap framework: Linking active learning to cognitive engagement</a> . <i>Educational Psychologist</i> , 49(4):219–243.                                                                                                                                                                                                                                        | 870 |
| 818 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 819 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 820 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 821 | Aakanksha Chowdhery, Sharan Narang, Jacob Devlin, Maarten Bosma, Gaurav Mishra, Adam Roberts, Paul Barham, Hyung Won Chung, Charles Sutton, Sebastian Gehrmann, Parker Schuh, Kensen Shi, Sasha Tsvyashchenko, Joshua Maynez, Abhishek Rao, Anselm Levskaya, Tyler Wang, Nan Du, Yinhan Liu, and 6 others. 2022. <a href="#">Palm: Scaling language modeling with pathways</a> . <i>arXiv preprint arXiv:2204.02311</i> . | 871 |
| 822 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 872 |
| 823 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 873 |
| 824 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 825 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 826 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 827 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 828 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 829 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 830 | Paul Christiano. 2022. <a href="#">Ai alignment and the role of human oversight in generative models</a> . <i>Alignment Research Journal</i> , 4:101–128.                                                                                                                                                                                                                                                                 | 874 |
| 831 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 875 |
| 832 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 876 |
| 833 | Meredith Christiansen. 2021. <a href="#">Algorithmic writing and digital literacy: How ai is reshaping composition</a> . <i>Digital Studies</i> , 12:55–73.                                                                                                                                                                                                                                                               | 877 |
| 834 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 835 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 836 | Andy Clark and David Chalmers. 1998. <a href="#">The extended mind</a> . <i>Analysis</i> , 58(1):7–19.                                                                                                                                                                                                                                                                                                                    | 878 |
| 837 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 879 |
| 838 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 890 |
| 839 | Cary Coglianese and David Lehr. 2017. <a href="#">Regulating by robot: Administrative decision making in the machine-learning era</a> . <i>Georgetown Law Journal</i> , 105:1147–1223.                                                                                                                                                                                                                                    | 891 |
| 840 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 892 |
| 841 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 893 |
| 842 | Debbie Cotton, Peter Cotton, and Sarah Shipway. 2023. <a href="#">Chatgpt, ai and the impact on academic integrity: Ai-assisted student writing</a> . <i>International Journal for Educational Integrity</i> , 19(1):1–16.                                                                                                                                                                                                | 894 |
| 843 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 895 |
| 844 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 845 |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
| 846 | Erika Darics. 2020. <a href="#">E-voice: A multimodal perspective on institutional writing in digital environments</a> . <i>Discourse, Context Media</i> , 35:100391.                                                                                                                                                                                                                                                     | 896 |
| 847 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 897 |
| 848 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 898 |
| 849 | Terrence W. Deacon. 1997. <i>The Symbolic Species: The Co-evolution of Language and the Brain</i> . W.W. Norton Company.                                                                                                                                                                                                                                                                                                  | 899 |
| 850 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 900 |
| 851 |                                                                                                                                                                                                                                                                                                                                                                                                                           | 901 |
|     | DeepSeek-AI. 2024. <a href="#">Deepseek-v3 technical report</a> . <i>arXiv preprint arXiv:2412.19437</i> .                                                                                                                                                                                                                                                                                                                | 902 |
|     |                                                                                                                                                                                                                                                                                                                                                                                                                           | 903 |
|     | Timothy Devinney. 2023. <a href="#">Plagiarism in the age of ai: Who owns the output?</a> <i>Journal of Business Ethics</i> .                                                                                                                                                                                                                                                                                             | 904 |
|     |                                                                                                                                                                                                                                                                                                                                                                                                                           |     |
|     | Nicholas Diakopoulos. 2016. <a href="#">Accountability in algorithmic decision making</a> . <i>Communications of the ACM</i> , 59(2):56–62.                                                                                                                                                                                                                                                                               |     |
|     | Emily Dinan, Laura Perez, and Jason Weston. 2022. <a href="#">Collaborative ai: Integrating language models into professional writing teams</a> . In <i>Proceedings of NeurIPS 2022</i> , pages 5123–5136.                                                                                                                                                                                                                |     |
|     | Daniel Dor. 2015. <i>The Instruction of Imagination: Language as a Social Communication Technology</i> . Oxford University Press.                                                                                                                                                                                                                                                                                         |     |
|     | Fernando Duarte, Maarten Sap, and Yejin Choi. 2022. <a href="#">Ai-generated speech and the decline of authentic online discourse</a> . <i>Proceedings of the 2022 Conference on Fairness, Accountability, and Transparency (FACCT)</i> .                                                                                                                                                                                 |     |
|     | Ronald Dworkin. 1996. <i>Freedom’s Law: The Moral Reading of the American Constitution</i> . Harvard University Press.                                                                                                                                                                                                                                                                                                    |     |
|     | Jacob Eisenstein and Danielle McNamara. 2023. <a href="#">Human-ai collaboration in writing: Rethinking authorship and editorial oversight</a> . <i>AI Society</i> , 38(2):177–192.                                                                                                                                                                                                                                       |     |
|     | Alexander Fabbri, Irene Li, Tianyi Tang, Caiming Xiong, and Dragomir Radev. 2022. <a href="#">Qmsum: A new benchmark for query-based multi-document summarization</a> . In <i>Proceedings of NAACL 2022</i> , pages 6145–6162.                                                                                                                                                                                            |     |
|     | Angela Fan, Shruti Bhosale, Holger Schwenk, Alexander Baevski, Guillaume Lample, and Michael Auli. 2021. <a href="#">Beyond english-centric multilingual machine translation</a> . In <i>Proceedings of ACL 2021</i> , pages 4403–4419.                                                                                                                                                                                   |     |
|     | James Felton. 2023. <a href="#">From spell-checkers to ai writing assistants: The evolution of digital literacy tools</a> . <i>Journal of Digital Communication</i> , 17(2):89–107.                                                                                                                                                                                                                                       |     |
|     | Emilio Ferrara. 2020. <a href="#">Characterizing social media manipulation in the 2020 u.s. presidential election</a> . <i>First Monday</i> , 25(11).                                                                                                                                                                                                                                                                     |     |
|     | Luciano Floridi. 2013. <i>The Ethics of Information</i> . Oxford University Press.                                                                                                                                                                                                                                                                                                                                        |     |
|     | Luciano Floridi. 2019. <a href="#">Establishing the rules for building trustworthy ai</a> . <i>Nature Machine Intelligence</i> , 1(6):261–262.                                                                                                                                                                                                                                                                            |     |
|     | Luciano Floridi and Marcello Chiriatti. 2020. <a href="#">Gpt-3: Its nature, scope, limits, and consequences</a> . <i>Minds and Machines</i> , 30:681–694.                                                                                                                                                                                                                                                                |     |
|     | Michel Foucault. 1984. What is an author? In Paul Rabinow, editor, <i>The Foucault Reader</i> , pages 101–120. Pantheon Books.                                                                                                                                                                                                                                                                                            |     |

|     |                                                                                |                                                                                 |      |
|-----|--------------------------------------------------------------------------------|---------------------------------------------------------------------------------|------|
| 905 | Scott Freeman, Sarah L. Eddy, Miles McDonough,                                 | Ben Hutchinson, Jasmine Collins, Mark Diaz, and Qian                            | 957  |
| 906 | Michelle K. Smith, Nnadozie Okoroafor, Hannah                                  | Yang. 2023. <a href="#">Ai for accessibility: Enabling inclusive</a>            | 958  |
| 907 | Jordt, and Mary Pat Wenderoth. 2014. <a href="#">Active learn-</a>             | <a href="#">digital communication</a> . <i>CHI Conference on Human</i>          | 959  |
| 908 | <a href="#">ing increases student performance in science, engi-</a>            | <a href="#">Factors in Computing Systems</a> .                                  | 960  |
| 909 | <a href="#">neering, and mathematics</a> . <i>Proceedings of the Na-</i>       |                                                                                 |      |
| 910 | <a href="#">tional Academy of Sciences (PNAS)</a> , 111(23):8410–              | Zhengbao Ji, Zhiting Hu, Patrick Lewis, and Meta AI.                            | 961  |
| 911 | 8415.                                                                          | 2023. <a href="#">Survey of hallucination in large language mod-</a>            | 962  |
|     |                                                                                | <a href="#">els</a> . <i>Transactions of the Association for Computa-</i>       | 963  |
| 912 | Iason Gabriel. 2020. <a href="#">Artificial intelligence, values, and</a>      | <a href="#">tional Linguistics</a> .                                            | 964  |
| 913 | <a href="#">alignment</a> . <i>Minds and Machines</i> , 30(4):411–437.         |                                                                                 |      |
| 914 | David Galbraith. 1999. <a href="#">Writing as a knowledge-</a>                 | Hong Jin Kang, Fabrice Harel-Canada, Muham-                                     | 965  |
| 915 | <a href="#">constituting process</a> . <i>Erkenntnis</i> , 50:357–370.         | mad Ali Gulzar, Violet Peng, and Miryung Kim.                                   | 966  |
|     |                                                                                | 2024. <a href="#">Human-in-the-loop synthetic text data in-</a>                 | 967  |
| 916 | John Gallagher, Jacob Hilton, and Owain Evans. 2023.                           | <a href="#">spection with provenance tracking</a> . <i>arXiv preprint</i>       | 968  |
| 917 | <a href="#">Adversarial robustness in ai-generated text detection:</a>         | <i>arXiv:2404.18881</i> .                                                       | 969  |
| 918 | <a href="#">A losing battle?</a> <i>arXiv preprint arXiv:2305.07692</i> .      |                                                                                 |      |
| 919 | Yaroslav Ganin, Evgeniya Ustinova, Hana Ajakan,                                | Neal Katyal and Daniel Epps. 2022. <a href="#">The criminal regu-</a>           | 970  |
| 920 | Pascal Germain, Hugo Larochelle, François Lavi-                                | <a href="#">lation of artificial intelligence</a> . <i>Harvard Law Review</i> , | 971  |
| 921 | olette, Mario Marchand, and Victor Lempitsky.                                  | 135:412–468.                                                                    | 972  |
| 922 | 2016. <a href="#">Domain-adversarial training of neural net-</a>               |                                                                                 |      |
| 923 | <a href="#">works</a> . <i>Journal of Machine Learning Research</i> ,          | Ronald T. Kellogg. 2008. <a href="#">Training writing skills: A cog-</a>        | 973  |
| 924 | 17(59):1–35.                                                                   | <a href="#">nitive developmental perspective</a> . <i>Journal of Writing</i>    | 974  |
|     |                                                                                | <i>Research</i> , 1(1):1–26.                                                    | 975  |
| 925 | Sebastian Gehrmann, Hendrik Strobelt, and Alexan-                              | Salman Khan, Pranav Rajpurkar, and Daphne Koller.                               | 976  |
| 926 | der M. Rush. 2019. <a href="#">Glr: Statistical detection and</a>              | 2023. <a href="#">Ai tutors: The role of large language</a>                     | 977  |
| 927 | <a href="#">visualization of ai-generated text</a> . <i>Proceedings of the</i> | <a href="#">models in personalized education</a> . <i>arXiv preprint</i>        | 978  |
| 928 | <i>57th Annual Meeting of the Association for Computa-</i>                     | <i>arXiv:2303.11288</i> .                                                       | 979  |
| 929 | <i>tional Linguistics (ACL)</i> .                                              |                                                                                 |      |
| 930 | Clark Glymour, David Danks, and Peter Spirtes. 2023.                           | Sahar Khowaja, Ammar Ahmad, Khaled Salah, Raja                                  | 980  |
| 931 | <a href="#">Ai explainability and its limits: The role of human</a>            | Jayaraman, and Ibrar Yaqoob. 2023. <a href="#">Blockchain</a>                   | 981  |
| 932 | <a href="#">oversight</a> . <i>Artificial Intelligence Review</i> .            | <a href="#">for ai: Review and open research challenges</a> . <i>IEEE</i>       | 982  |
|     |                                                                                | <i>Transactions on Artificial Intelligence</i> , 4(1):1–15.                     | 983  |
| 933 | H.P. Grice. 1975. Logic and conversation. In <i>Syntax</i>                     | Johannes Kirchenbauer, Jonas Geiping, Yuxuan Han,                               | 984  |
| 934 | <i>and Semantics</i> , volume 3, pages 41–58. Academic                         | Elliot Creager, Ari S. Morcos, and Tom Goldstein.                               | 985  |
| 935 | Press.                                                                         | 2023. A watermark for large language models. <i>arXiv</i>                       | 986  |
|     |                                                                                | <i>preprint arXiv:2307.06624</i> .                                              | 987  |
| 936 | Philipp Hacker, Andreas Engel, and Marco Mauer. 2023.                          | James Kirkpatrick, Razvan Pascanu, Neil Rabinowitz,                             | 988  |
| 937 | Regulating chatgpt and other large generative ai mod-                          | Joel Veness, Guillaume Desjardins, Andrei A. Rusu,                              | 989  |
| 938 | els. In <i>Proceedings of the 2023 ACM Conference on</i>                       | Kieran Milan, John Quan, Tiago Ramalho, Ag-                                     | 990  |
| 939 | <i>Fairness, Accountability, and Transparency</i> , pages                      | nieszka Grabska-Barwinska, Demis Hassabis, Clau-                                | 991  |
| 940 | 1112–1123.                                                                     | dia Clopath, Dharshan Kumaran, and Raia Hadsell.                                | 992  |
|     |                                                                                | 2017. <a href="#">Overcoming catastrophic forgetting in neural</a>              | 993  |
| 941 | Ari Holtzman, Jan Buys, Li Du, Maxwell Forbes, and                             | <a href="#">networks</a> . <i>Proceedings of the National Academy of</i>        | 994  |
| 942 | Yejin Choi. 2020. <a href="#">The curious case of neural text</a>              | <i>Sciences</i> , 114(13):3521–3526.                                            | 995  |
| 943 | <a href="#">degeneration</a> . In <i>Proceedings of the International</i>      |                                                                                 |      |
| 944 | <i>Conference on Learning Representations (ICLR)</i> .                         | Andrew J. Koch, Robert D. Gerber, and Sarah C.                                  | 996  |
|     |                                                                                | Roberts. 2015. <a href="#">Structured interviews: Reducing bias</a>             | 997  |
| 945 | Dirk Hovy and Shannon L. Spruit. 2016. <a href="#">The social</a>              | <a href="#">and increasing hiring effectiveness</a> . <i>Journal of Ap-</i>     | 998  |
| 946 | <a href="#">impact of natural language processing</a> . <i>Proceedings</i>     | <i>plied Psychology</i> , 100(3):775–789.                                       | 999  |
| 947 | <i>of the 54th Annual Meeting of the Association for</i>                       |                                                                                 |      |
| 948 | <i>Computational Linguistics (ACL)</i> , pages 591–598.                        | Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yu-                             | 1000 |
|     |                                                                                | taka Matsuo, and Yusuke Iwasawa. 2022. Large lan-                               | 1001 |
| 949 | Yang Hu, Heather R. Younger, and Patricia M. Green-                            | <a href="#">guage models are zero-shot reasoners</a> . <i>Advances in</i>       | 1002 |
| 950 | field. 2020. <a href="#">Language and hiring bias: How intro-</a>              | <i>neural information processing systems</i> , 35:22199–                        | 1003 |
| 951 | <a href="#">verts and neurodiverse candidates face discrimination</a>          | 22213.                                                                          | 1004 |
| 952 | <a href="#">in spontaneous evaluations</a> . <i>Journal of Business and</i>    |                                                                                 |      |
| 953 | <i>Psychology</i> , 35(2):215–230.                                             | Jeff Kosseff. 2019. <a href="#">The Twenty-Six Words That Created</a>           | 1005 |
|     |                                                                                | <a href="#">the Internet</a> . Cornell University Press.                        | 1006 |
| 954 | Ming-Hui Huang and Roland T. Rust. 2021. <a href="#">Artifi-</a>               | Ben Krause, Ethan Wilcox, Max Bittker, and Christo-                             | 1007 |
| 955 | <a href="#">cial intelligence in business: Promise, pitfalls, and</a>          | pher D. Manning. 2022. <a href="#">Co-authoring with ai: How</a>                | 1008 |
| 956 | <a href="#">prospects</a> . <i>Journal of Service Research</i> , 24(1):3–6.    | <a href="#">llms shape collaborative writing practices</a> . <i>Transac-</i>    | 1009 |
|     |                                                                                | <i>tions of the ACL</i> , 10:413–429.                                           | 1010 |

|      |                                                                                                                                                                                                                                                               |      |
|------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------|
| 1011 | Vivian Lai and Laura Viering. 2022. <a href="#">Ai assistance and over-reliance: Implications for human judgment and decision-making</a> . <i>Proceedings of the ACM Conference on Human Factors in Computing Systems (CHI)</i> .                             | 1064 |
| 1012 |                                                                                                                                                                                                                                                               | 1065 |
| 1013 |                                                                                                                                                                                                                                                               | 1066 |
| 1014 |                                                                                                                                                                                                                                                               |      |
| 1015 | George Lakoff and Mark Johnson. 1980. <i>Metaphors We Live By</i> . University of Chicago Press.                                                                                                                                                              | 1067 |
| 1016 |                                                                                                                                                                                                                                                               | 1068 |
| 1017 | Jochen Laubrock, Clara Martin, and Marc Brysbaert. 2022. <a href="#">Readership enhancement via ai-assisted writing tools: A cognitive science perspective</a> . <i>Cognitive Science</i> , 46(4):e13120.                                                     | 1069 |
| 1018 |                                                                                                                                                                                                                                                               | 1070 |
| 1019 |                                                                                                                                                                                                                                                               |      |
| 1020 |                                                                                                                                                                                                                                                               |      |
| 1021 | Rebekah Layton and Carolyn Watters. 2020. <a href="#">Authorship attribution with deep learning</a> . <i>Digital Scholarship in the Humanities</i> , 35(2):317–331.                                                                                           | 1071 |
| 1022 |                                                                                                                                                                                                                                                               | 1072 |
| 1023 |                                                                                                                                                                                                                                                               | 1073 |
| 1024 | Jisoo Lee, Daniel McNamara, and Tanja Käser. 2022. <a href="#">Co-authoring with ai: How language models support second-language writing development</a> . <i>Computers Education</i> , 186:104536.                                                           | 1074 |
| 1025 |                                                                                                                                                                                                                                                               | 1075 |
| 1026 |                                                                                                                                                                                                                                                               | 1076 |
| 1027 |                                                                                                                                                                                                                                                               | 1077 |
| 1028 | Brian Leonard and Kevin Noonan. 2020. <a href="#">Computational text generation in professional and technical writing</a> . <i>Journal of Business and Technical Communication</i> , 34(4):451–474.                                                           | 1078 |
| 1029 |                                                                                                                                                                                                                                                               | 1079 |
| 1030 |                                                                                                                                                                                                                                                               |      |
| 1031 |                                                                                                                                                                                                                                                               |      |
| 1032 | Julia Levashina, Christian J. Hartwell, Frederick P. Morgeson, and Michael A. Campion. 2014. <a href="#">The structured employment interview: Narrative and quantitative review of the research literature</a> . <i>Personnel Psychology</i> , 67(1):241–293. | 1080 |
| 1033 |                                                                                                                                                                                                                                                               | 1081 |
| 1034 |                                                                                                                                                                                                                                                               | 1082 |
| 1035 |                                                                                                                                                                                                                                                               | 1083 |
| 1036 |                                                                                                                                                                                                                                                               | 1084 |
| 1037 | Matthew Lipman. 2003. <i>Thinking in Education</i> , 2nd edition. Cambridge University Press.                                                                                                                                                                 | 1085 |
| 1038 |                                                                                                                                                                                                                                                               | 1086 |
| 1039 | Rose Luckin, Wayne Holmes, and Joshua Greer. 2023. <a href="#">Ai and education: The future of personalized learning</a> . <i>AI Education Journal</i> , 4(1):112–130.                                                                                        | 1087 |
| 1040 |                                                                                                                                                                                                                                                               |      |
| 1041 |                                                                                                                                                                                                                                                               |      |
| 1042 | Brian Lund and Weilin Wang. 2023. <a href="#">Reinventing assessments in the age of ai: From written exams to oral evaluations</a> . <i>AI Education Journal</i> , 2(1):22–35.                                                                                | 1088 |
| 1043 |                                                                                                                                                                                                                                                               | 1089 |
| 1044 |                                                                                                                                                                                                                                                               | 1090 |
| 1045 | Richard Menary. 2010. <a href="#">Cognitive integration and the extended mind</a> . In Richard Menary, editor, <i>The Extended Mind</i> , pages 227–243. MIT Press.                                                                                           | 1091 |
| 1046 |                                                                                                                                                                                                                                                               | 1092 |
| 1047 |                                                                                                                                                                                                                                                               |      |
| 1048 | Jacob Menick, Jeffrey Shlens, Xiaohua Zhai, Neil Houlsby, Andrea Gesmundo, Avital Oliver, and Karen Simonyan. 2022. <a href="#">Reducing the recurrence of errors in language model outputs</a> . <i>NeurIPS</i> .                                            | 1093 |
| 1049 |                                                                                                                                                                                                                                                               | 1094 |
| 1050 |                                                                                                                                                                                                                                                               |      |
| 1051 |                                                                                                                                                                                                                                                               |      |
| 1052 | Arthur I Miller. 2019. <i>The artist in the machine: The world of AI-powered creativity</i> . MIT Press.                                                                                                                                                      | 1095 |
| 1053 |                                                                                                                                                                                                                                                               | 1096 |
| 1054 | Brent Daniel Mittelstadt and Luciano Floridi. 2016. <a href="#">The ethics of algorithms: Mapping the debate</a> . <i>Big Data Society</i> , 3(2):1–21.                                                                                                       | 1097 |
| 1055 |                                                                                                                                                                                                                                                               | 1098 |
| 1056 |                                                                                                                                                                                                                                                               |      |
| 1057 | Arvind Narayanan, Joseph Bonneau, Edward Felten, Andrew Miller, and Steven Goldfeder. 2016. <i>Bitcoin and Cryptocurrency Technologies: A Comprehensive Introduction</i> . Princeton University Press.                                                        | 1099 |
| 1058 |                                                                                                                                                                                                                                                               | 1100 |
| 1059 |                                                                                                                                                                                                                                                               | 1101 |
| 1060 |                                                                                                                                                                                                                                                               | 1102 |
| 1061 | Tom Nichols. 2021. The death of expertise: The campaign against established knowledge and why it matters. <i>Oxford University Press</i> .                                                                                                                    | 1103 |
| 1062 |                                                                                                                                                                                                                                                               | 1104 |
| 1063 |                                                                                                                                                                                                                                                               | 1105 |
|      |                                                                                                                                                                                                                                                               | 1106 |
|      |                                                                                                                                                                                                                                                               | 1107 |
|      |                                                                                                                                                                                                                                                               | 1108 |
|      |                                                                                                                                                                                                                                                               | 1109 |
|      |                                                                                                                                                                                                                                                               | 1110 |
|      |                                                                                                                                                                                                                                                               | 1111 |
|      |                                                                                                                                                                                                                                                               | 1112 |
|      |                                                                                                                                                                                                                                                               | 1113 |
|      |                                                                                                                                                                                                                                                               | 1114 |
|      |                                                                                                                                                                                                                                                               | 1115 |



Shunyu Yao, Dian Yu, Jeffrey Zhao, Nan Du, Izhak Shafran, Thomas L. Griffiths, Yuan Cao, and Karthik Narasimhan. 2023. [Tree of thoughts: Deliberate problem solving with large language models](#). *arXiv preprint arXiv:2305.10601*.

Jason Yosinski, Dario Amodei, and Alec Radford. 2023. [Co-editing with ai: The role of large language models in real-time group writing](#). In *Proceedings of ACL 2023*, pages 3145–3159.

Rowan Zellers, Ari Holtzman, Hannah Rashkin, Yonatan Bisk, Ali Farhadi, Franziska Roesner, and Yejin Choi. 2019. [Defending against neural fake news](#). In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 32.

Xia Zhou, Yang Xu, and Feng Liu. 2023. [Can ai-generated text be reliably detected? evaluating plagiarism detection tools on gpt-based texts](#). *arXiv preprint arXiv:2304.08979*.

Shoshana Zuboff. 2019. *The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power*. PublicAffairs.

## A Example Appendix

This is an appendix.