



List of contents available at [Lingua Technica](https://journal.arjunu.org/index.php/lingtech)

Lingua Technica: Journal of Digital Literary Studies

homepage: <https://journal.arjunu.org/index.php/lingtech>



The style that was learned: large language models, imitation, and the ethics of literary influence

Nomvuselelo N. Nxumalo ^{1*} and Seroja Ainun Nadhifah ²

¹ University of Eswatini, Kingdom of Eswatini

² Universitas Islam Indonesia, Indonesia

* Corresponding author. Email: nomvuselelonxumalo09@gmail.com

ABSTRACT

Article history:

Received May 24, 2026

Revised Jun 26, 2026

Accepted Jun 29, 2026

Published Jul 31, 2026

Keywords:

authorship;
large language models;
literary influence;
literary style imitation;
textual reproduction

Background: Large language models have transformed literary imitation into a scalable practice, complicating distinctions among stylistic influence, textual reproduction, authorship, and ethical appropriation. **Objective:** This study examines how author-conditioned generation produces stylistic convergence, preserves or transforms source relations, and acquires ethical legitimacy under documented research conditions. **Method:** This study applies paired stylometry, intertextual reproduction auditing, and an ethical literary influence matrix to ten public-domain literary excerpts, ten controlled model outputs, five methodological documents, and seven contextual sources. **Results:** Stylometric comparison reveals graded convergence, with recognizable authorial cues coexisting with substantial divergence in sentence architecture and other distributed features. Reproduction auditing identifies low lexical overlap, no shared four-grams, no source-character or source-event transfer, and consistently high transformation across all pairs. Ethical coding classifies risk as low because sources are public domain, authors are deceased, prompts and exemplars are disclosed, and human control is documented. **Implication:** These findings indicate that literary resemblance should not be treated automatically as authorship, memorization, plagiarism, or infringement, but evaluated through layered textual and contextual evidence. **Novelty:** This study advances a multidimensional framework that theorizes learned style as mediated literary influence shaped by convergence, transformation, provenance, and accountable use across computational, literary, legal, and contemporary ethical domains.

DOI: <https://doi.org/10.64595/80qxyz22>

*) Copyright © 2026 Nomvuselelo N. Nxumalo & Seroja Ainun Nadhifah

This work is licensed under a [Creative Commons Attribution-ShareAlike 4.0 International License](https://creativecommons.org/licenses/by-sa/4.0/)  which permits use, distribution, reproduction, and adaptation in any medium, provided that the original work is properly cited and any adapted material is distributed under the same license.

INTRODUCTION

Generative language systems have moved literary imitation from a marginal experiment to a scalable cultural practice in which authorial voices can be requested, reproduced, circulated, and monetized within seconds. This shift matters because style is not merely decorative form; it embodies accumulated choices concerning syntax, rhythm, viewpoint, imagery, genre memory, and cultural position. CopyBench found that, when Llama 3 models increased from 8B to 70B parameters, literal copying rose from 0.2% to 10.5%, while non-literal copying increased from 2.3% to 5.9%; event and character copying appeared even at 7B parameters (Chen et al., 2024). Meanwhile, scalable watermarking has emerged precisely because synthetic prose can circulate without reliable provenance (Dathathri et al., 2024). Literary influence operates through infrastructures that can obscure whether resemblance reflects learned convention, prompted emulation, memorized expression, or post hoc editing. This ambiguity creates a concrete problem for authorship, attribution, copyright governance, and cultural trust in literary production.

Existing scholarship has clarified components of this problem without integrating them into a literary account of influence. Research on linguistic novelty shows that language models can recombine training patterns rather than reproduce passages verbatim (McCoy et al., 2023), whereas style vectors and controllable generation demonstrate that stylistic features can be steered (Konen et al., 2024; Toshevska & Gievska, 2025). Stylometric studies further show convergence between model outputs and literary targets, although imitation remains uneven across features, genres, and prompting conditions (Mikros, 2025; Wang et al., 2025). Other work identifies synthetic stylistic fingerprints, editorial traces, and demographic homogenization in machine-mediated authorship (Alvero et al., 2024; DeHaan et al., 2025). Digital-literary scholarship has examined platformed textuality and algorithmic poetics (Abdillah, 2026; Fawaid, 2025). Yet these strands rarely distinguish stylistic convergence, narrative borrowing, textual reproduction, and ethically problematic impersonation within one comparative framework. Consequently, literary influence remains undertheorized when imitation is computationally produced.

This study investigates how large language models transform literary influence when prompted to write through authorial styles. It asks three questions. First, which stylometric features converge between public-domain source passages and controlled model outputs, and which features resist imitation? Second, does stylistic proximity arise from lexical overlap, structural echo, narrative transfer, or genre conventions rather than direct reproduction? Third, under what conditions does computational imitation remain a transparent experiment, and when does it approach ethically problematic appropriation, misattribution, or identity substitution? To answer these questions, this study compares paired literary source excerpts and model-generated passages through a stylometric convergence matrix, an intertextual reproduction and novelty audit, and an ethical literary influence evaluation matrix. This design avoids treating statistical similarity as proof of memorization or infringement. Instead, it examines how forms of resemblance become legible across linguistic, narrative, and normative levels, thereby connecting computational measurement with literary interpretation and ethical judgment.

This study argues that model-generated style should be understood neither as autonomous authorship nor as copying, but as mediated influence whose ethical status depends on how resemblance is produced, disclosed, and used. Stylometric convergence may indicate feature alignment without demonstrating possession of an authorial identity; lexical novelty may coexist with structural or narrative appropriation; and public-domain legality does not resolve questions of attribution, deception, or cultural substitution. Accordingly, imitation becomes more defensible when source status is transparent, prompts are documented, outputs are disclosed as synthetic, and analysis distinguishes influence from impersonation. Conversely, ethical risk intensifies when deployment suppresses provenance, invokes living authors without consent, or presents generated voice as continuity. This proposition treats literary style as relational rather than proprietary in every respect, while recognizing that computational scale changes the consequences of imitation. Such a framework can refine debates on authorship, copyright, platform accountability, and literary creativity across literary cultures.

LITERATURE REVIEW

Literary style imitation

Literary style imitation refers to reproducing recurrent linguistic and rhetorical patterns associated with an author, genre, or register while generating new propositional content. Computational research treats style either as a controllable attribute separable from meaning or as an emergent configuration inseparable from authorial practice. Khan et al. (2023) and Konen et al. (2024) operationalize style through learned representations and activation steering, whereas Toshevskaya and Gievska (2025) frame imitation as text style transfer. Mikros (2025), however, approaches literary imitation stylometrically, emphasizing multidimensional resemblance rather than categorical equivalence. Accordingly, imitation denotes stable graded convergence, not possession or reconstruction of an authorial identity.

Its principal indicators span lexical, syntactic, rhetorical, and distributional levels. Lexical richness, function-word patterns, sentence length, punctuation, dialogue, and narrative viewpoint capture observable regularities, while embeddings model less explicit stylistic relations. Prompted imitation may reproduce salient surface markers yet miss implicit habits, as Chen and Moscholios (2024) and Wang et al. (2025) suggest. Authorship verification therefore complements feature comparison by testing whether generated passages remain distinguishable from target writing (Weerasinghe et al., 2025). Mikros (2025) similarly demonstrates that literary fidelity varies across feature families. Robust assessment consequently requires converging indicators rather than a single similarity score or intuitive reader judgment.

Intertextual reproduction

Intertextual reproduction concerns the extent to which generated writing reuses protected or source-specific expression, structures, events, or characters. It differs from linguistic novelty, which asks whether sequences are unprecedented, and from plagiarism, which additionally concerns attribution and representational conduct. RAVEN conceptualizes copying through graded novelty across units (McCoy et al., 2023), whereas CopyBench separates literal from non-literal reproduction (Chen et al., 2024). Tripto et al. (2024) show that paraphrasing can transform surface wording while preserving underlying relations. Wahle et al. (2022) situate such transformation within machine-paraphrase plagiarism. Thus, textual difference cannot automatically establish independence from a recognizable specific literary source.

Reproduction is best evaluated through complementary indicators: exact phrase overlap, longest shared sequences, fuzzy similarity, event retention, character retention, semantic correspondence, and structural transformation. Verbatim extraction provides conservative evidence of memorization (Karamolegkou et al., 2023), but non-literal resemblance may remain relevant when protected organization survives lexical alteration (Chen et al., 2024). Copyright-notice compliance and refusal behavior add model-level evidence, although responses vary by prompt and system (Xu et al., 2024; Müller et al., 2024). SHIELD further distinguishes inappropriate reproduction from excessive refusal (Liu et al., 2024). Consequently, reproduction auditing must separate measurable textual dependence from legal conclusions requiring contextual interpretation.

Ethical literary influence

Ethical literary influence describes how generated writing relates to prior authors through legitimacy, disclosure, attribution, consent, and cultural consequence. Unlike copyright analysis, which concerns legally protected interests, ethical assessment addresses deception, identity substitution, asymmetrical appropriation, and representational power. Choksi and Goedicke (2023) foreground ownership and data-governance tensions, while Alvero et al. (2024) show that synthetic authorship can reproduce hegemonic demographic patterns. Sourati et al. (2025) extend this concern to expressive homogenization. Copyright scholarship, meanwhile, emphasizes licensing and risk management rather than treating every resemblance as infringement (Sag, 2023; Smith, 2025). Ethical influence remains relational, context-dependent and broader than textual similarity.

Its indicators include source copyright status, author living status, consent, exemplar provision, prompt specificity, disclosure of synthetic production, attribution, editorial control, market substitution, and identity-confusion risk. Provenance mechanisms such as watermarking may support disclosure but cannot resolve

legitimacy alone (Dathathri et al., 2024). Personalized generation also creates tensions among mimicking, privacy, and verification (Nguyen et al., 2025), while responsible-use frameworks emphasize human accountability and transparent assistance (Mann et al., 2023; Mijatović et al., 2026). This study therefore theorizes learned style as mediated influence: ethically acceptable when provenance and transformation are legible, but problematic when opacity converts resemblance into undeclared impersonation.

METHOD

Units of analysis comprised thirty-two unique textual documents organized into ten paired comparisons and two supporting strata. Twenty primary items consisted of ten public-domain literary excerpts and ten controlled model outputs generated from an identical narrative brief under author-specific style conditions. Five secondary documents supplied validated procedures for stylometry, novelty assessment, and reproduction auditing, while seven contextual documents established institutional guidance on copyright, attribution, disclosure, and responsible use. Each document received a unique identifier, pair identifier, source classification, bibliographic record, access basis, and deduplication key. This layered corpus enabled comparison without treating policy documents or methodological literature as literary evidence.

This study employed a comparative mixed-method corpus design integrating paired stylometry, intertextual auditing, and normative interpretation. Comparison was within pairs rather than across unrelated texts: every generated passage responded to the same narrative brief but adopted one public-domain authorial condition. This control reduced content variation while preserving stylistic differences. Quantitative measures were descriptive and comparative, not inferential, because ten pairs cannot justify generalization or causal claims. Qualitative close reading examined how proximity materialized through syntax, focalization, dialogue, imagery, and narrative organization. This design follows multidimensional approaches to literary imitation and style steering rather than equating one similarity score with authorship.

Primary literary texts were retrieved from Project Gutenberg editions identified as public domain in the United States, with title, author, publication year, edition URL, access date, and excerpt location recorded. Generated texts were produced in one documented ChatGPT session using a constant narrative prompt and author-specific style condition; model label, generation date, prompt, and output were preserved verbatim. Secondary information came from open scholarly articles and repositories describing style vectors, stylometric imitation, linguistic novelty, and literal or non-literal reproduction. Contextual information came from official copyright, authorship, and responsible-use guidance. Only publicly accessible, attributable, and verifiable materials entered this research corpus.

Data collection followed a preregistered-style protocol encoded in the workbook. Searches used combinations of author name, title, literary style, imitation, stylometry, text generation, memorization, reproduction, attribution, copyright, and disclosure. Candidate records were screened for public accessibility, source authority, textual relevance, stable identification, and sufficient methodological detail. Exclusion removed inaccessible items, unverifiable copies, duplicate editions, living-author imitation targets, nontextual materials without transcripts, and records lacking provenance. Metadata and full analytical excerpts were stored separately, linked through document and pair identifiers. Exact-title, URL, author-year, and content-hash checks removed duplication, while retained decisions and exclusion reasons remained auditable. Collection ended after saturation criteria.

Analysis proceeded in three sequential but interpretively integrated stages. First, paired stylometry calculated sentence length, lexical diversity, lexical density, word length, punctuation, first-person reference, dialogue, function-word similarity, and an exploratory convergence index; close reading then explained feature-level patterns (Konen et al., 2024; Mikros, 2025). Second, reproduction auditing measured word-set Jaccard similarity, character-trigram cosine, shared four-grams, longest common sequences, and narrative-event or character overlap, separating literal recurrence from transformation (McCoy et al., 2023; Chen et al., 2024). Third, an ethical matrix assessed rights status, consent, disclosure, attribution, human control, substitution risk, and identity confusion without converting normative indicators into legal determinations.

Table 1. Composition of dataset

Code	Data type	Source	Location	Period	Number	Principal characteristics
PAIR-01	Literary source–LLM pair	<i>Pride and Prejudice</i> and controlled output	Project Gutenberg; digital generation environment	1813; 2026	2	Austen-conditioned prose; constant narrative brief
PAIR-02	Literary source–LLM pair	<i>A Tale of Two Cities</i> and controlled output	Project Gutenberg; digital generation environment	1859; 2026	2	Dickens-conditioned prose; matched narrative content
PAIR-03	Literary source–LLM pair	<i>Moby-Dick; or, The Whale</i> and controlled output	Project Gutenberg; digital generation environment	1851; 2026	2	Melville-conditioned prose; paired feature comparison
PAIR-04	Literary source–LLM pair	<i>Frankenstein; or, The Modern Prometheus</i> and controlled output	Project Gutenberg; digital generation environment	1818; 2026	2	Shelley-conditioned prose; public-domain target
PAIR-05	Literary source–LLM pair	<i>The Picture of Dorian Gray</i> and controlled output	Project Gutenberg; digital generation environment	1890; 2026	2	Wilde-conditioned prose; identical event structure
PAIR-06	Literary source–LLM pair	<i>Heart of Darkness</i> and controlled output	Project Gutenberg; digital generation environment	1899; 2026	2	Conrad-conditioned prose; narrative and stylistic comparison
PAIR-07	Literary source–LLM pair	<i>Adventures of Huckleberry Finn</i> and controlled output	Project Gutenberg; digital generation environment	1884; 2026	2	Twain-conditioned prose; lexical and narrative comparison
PAIR-08	Literary source–LLM pair	<i>The Adventures of Sherlock Holmes</i> and controlled output	Project Gutenberg; digital generation environment	1892; 2026	2	Doyle-conditioned prose; matched prompting condition
PAIR-09	Literary source–LLM pair	<i>Dracula</i> and controlled output	Project Gutenberg; digital generation environment	1897; 2026	2	Stoker-conditioned prose; epistolary and narrative indicators
PAIR-10	Literary source–LLM pair	<i>Dubliners</i> and controlled output	Project Gutenberg; digital generation environment	1914; 2026	2	Joyce-conditioned prose; paired close reading
SEC-001–005	Secondary methodological documents	ACL Anthology, TACL, scholarly repositories	International digital scholarship	2023–2026	5	Stylometry, style steering, novelty, reproduction, and imitation methods
CTX-001–007	Contextual institutional documents	Project Gutenberg, U.S. Copyright Office, WIPO, Authors Guild, Google	United States and international	1886–2026	7	Rights status, authorship, registration, disclosure, and platform governance
Total	Integrated corpus	Primary, secondary, and contextual sources	Digital and institutional	1813–2026	32	Ten matched pairs and twelve supporting documents

RESULT

Stylometric convergence without authorial equivalence

Pairwise comparison reveals a graded rather than binary pattern of stylistic convergence. Table 2 organizes ten author-conditioned outputs by exploratory similarity band and juxtaposes aggregate similarity with sentence length, lexical density, and feature-level alignment. Two pairs, Shelley and Wilde, occupy the very-high band; Twain falls within the high band; four occupy the moderate band; and three remain limited. These bands summarize relative proximity within this corpus and do not identify authorship or establish memorization. Because every output followed the same rain-coast narrative brief, differences are read against a controlled semantic frame. This arrangement permits stylistic adaptation to be examined as selective feature alignment rather than wholesale replication, consistent with multidimensional approaches to style representation and literary imitation (Konen et al., 2024; Mikros, 2025).

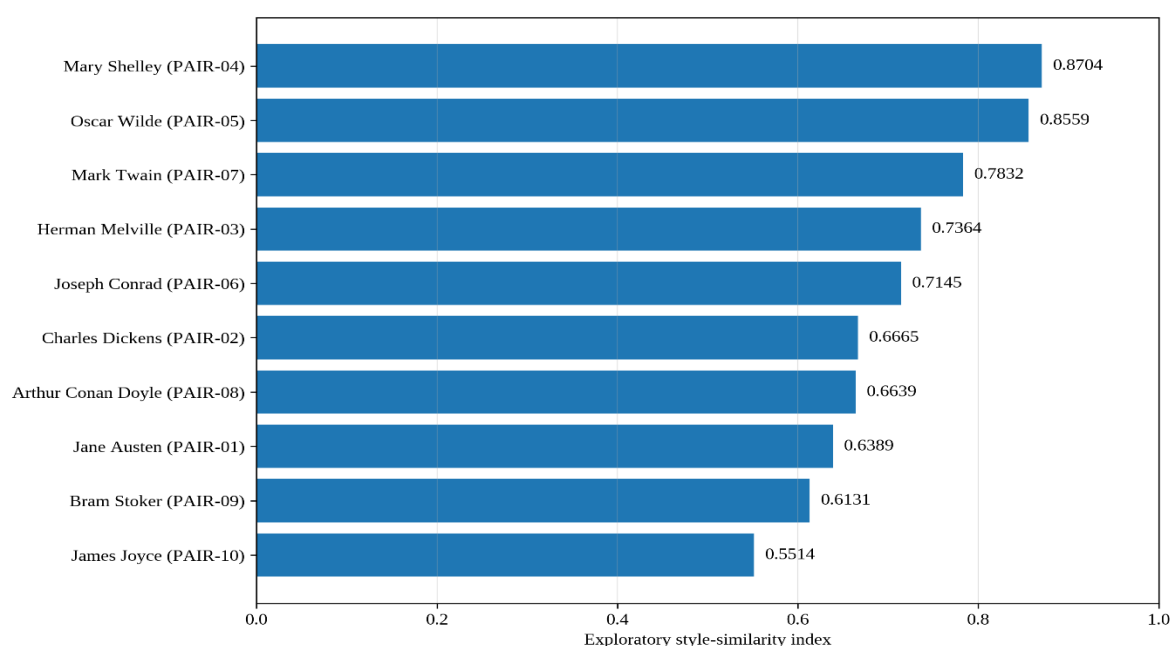


Figure 3. Pairwise stylometric convergence across authorial conditions

Selective convergence constitutes the dominant pattern across all ten pairings. Shelley records the highest index (0.8704), followed by Wilde (0.8559), although Wilde's mean sentence length contracts sharply from 65.333 to 15.333 words. Austen shows almost exact sentence-length alignment, at 15.750 and 16.000 words, but only limited overall convergence (0.6389), confirming that one matched feature cannot determine a composite profile. Generated passages repeatedly foreground recognizable cues: Melville's condition begins, "Call the town Salt Mercy"; Stoker's adopts the dated entry "11 October. Reached Varen Head after sunset"; and Twain's uses colloquial narration about rain "saving up all summer." Joyce produces the lowest index (0.5514), alongside increased lexical density and comma use. Surface recognizability therefore coexists with measurable structural divergence.

These distributions support a view of literary style as a bundle of partially independent features rather than a transferable authorial essence. High composite proximity can coexist with substantial divergence in sentence architecture, while lower composite scores may contain striking local resemblance. Such unevenness accords with research showing that style steering captures salient lexical and rhetorical markers more reliably than implicit author-specific habits (Chen & Moscholios, 2024; Wang et al., 2025). It reinforces Mikros's (2025) argument that literary imitation must be assessed across feature families rather than through impressionistic resemblance. Within this corpus, model outputs operate as selective reconstructions: they activate recognizable signals, redistribute less visible regularities, and preserve independent wording. Stylometric convergence consequently indicates mediated influence, not authorial equivalence or complete stylistic reproduction.

Table 2. Paired Stylometric convergence across author-conditioned outputs

Convergence band	Authorial condition / pair	Band frequency	Similarity index	Mean sentence length: source → generated	Lexical density: source → generated	Closest aligned feature: source → generated	Largest divergence: source → generated	Illustrative generated phrase	Data code
Very high	Mary Shelley / PAIR-04	2 pairs (20%)	0.8704	21.333 → 18.500	0.5000 → 0.5459	Mean word length: 4.552 → 4.632	Comma rate: 6.250 → 7.568	“daylight, exhausted by the storm, withdrew from every object”	PRI-GEN-SHE-001
	Oscar Wilde / PAIR-05	2 pairs (20%)	0.8559	65.333 → 15.333	0.5612 → 0.5924	Hapax ratio: 0.8971 → 0.8647	Sentence length: 65.333 → 15.333	“nature’s least imaginative form of criticism”	PRI-GEN-WIL-001
High	Mark Twain / PAIR-07	1 pair (10%)	0.7832	16.167 → 16.455	0.5052 → 0.5304	Mean word length: 3.804 → 3.862	Comma rate: 7.216 → 5.525	“saving up all summer just to spend itself on me”	PRI-GEN-TWA-001
Moderate	Herman Melville / PAIR-03	4 pairs (40%)	0.7364	24.125 → 15.818	0.5440 → 0.6149	Mean word length: 4.399 → 4.437	Comma rate: 5.181 → 9.195	“Call the town Salt Mercy”	PRI-GEN-MEL-001
	Joseph Conrad / PAIR-06	4 pairs (40%)	0.7145	21.222 → 16.455	0.5183 → 0.5359	Mean word length: 4.319 → 4.309	Sentence length: 21.222 → 16.455	“a dark smudge at the edge of an immense and indifferent water”	PRI-GEN-CON-001
	Charles Dickens / PAIR-02	4 pairs (40%)	0.6665	47.750 → 19.000	0.4188 → 0.5965	Comma rate: 10.471 → 9.357	Sentence length: 47.750 → 19.000	“Rain upon the roofs, rain in the gutters”	PRI-GEN-DIC-001

	Arthur Conan Doyle / PAIR-08	4 pairs (40%)	0.6639	17.455 → 13.214	0.5104 → 0.5351	Hapax ratio: 0.7840 → 0.7812	Sentence length: 17.455 → 13.214	"I regarded it almost as an observer"	PRI-GEN-DOY-001
Limited	Jane Austen / PAIR-01	3 pairs (30%)	0.6389	15.750 → 16.000	0.4815 → 0.5682	Sentence length: 15.750 → 16.000	Hapax ratio: 0.6696 → 0.8160	"grey sea, grey road"	PRI-GEN-AUS-001
	Bram Stoker / PAIR-09	3 pairs (30%)	0.6131	18.900 → 10.438	0.4974 → 0.5868	Hapax ratio: 0.7787 → 0.8661	Sentence length: 18.900 → 10.438	"11 October.— Reached Varen Head after sunset"	PRI-GEN-STO-001
	James Joyce / PAIR-10	3 pairs (30%)	0.5514	19.000 → 11.000	0.4579 → 0.5615	Mean word length: 3.826 → 4.278	Comma rate: 4.211 → 7.487	"Rain made the coast town intimate and mean"	PRI-GEN-JOY-001

Reproduction boundaries: lexical novelty, structural echo, and narrative transformation

Reproduction was assessed across the same ten source–output pairs used for stylistometric comparison, preserving document codes, narrative controls, and paired units of analysis. Table 3 combines word-set overlap, character-trigram similarity, function-word similarity, longest shared sequence, four-gram recurrence, narrative transfer, and transformation degree. These indicators separate lexical reuse from broader structural resemblance because a model may approximate rhythm or register without repeating distinctive expressions. Following RAVEN and CopyBench, literal and non-literal dependence were therefore coded independently (McCoy et al., 2023; Chen et al., 2024). Values describe relations between approximately 190-word source excerpts and generated scenes; they do not identify training-data membership, plagiarism, or copyright infringement. This design locates reproduction boundaries through converging textual evidence rather than a single threshold.

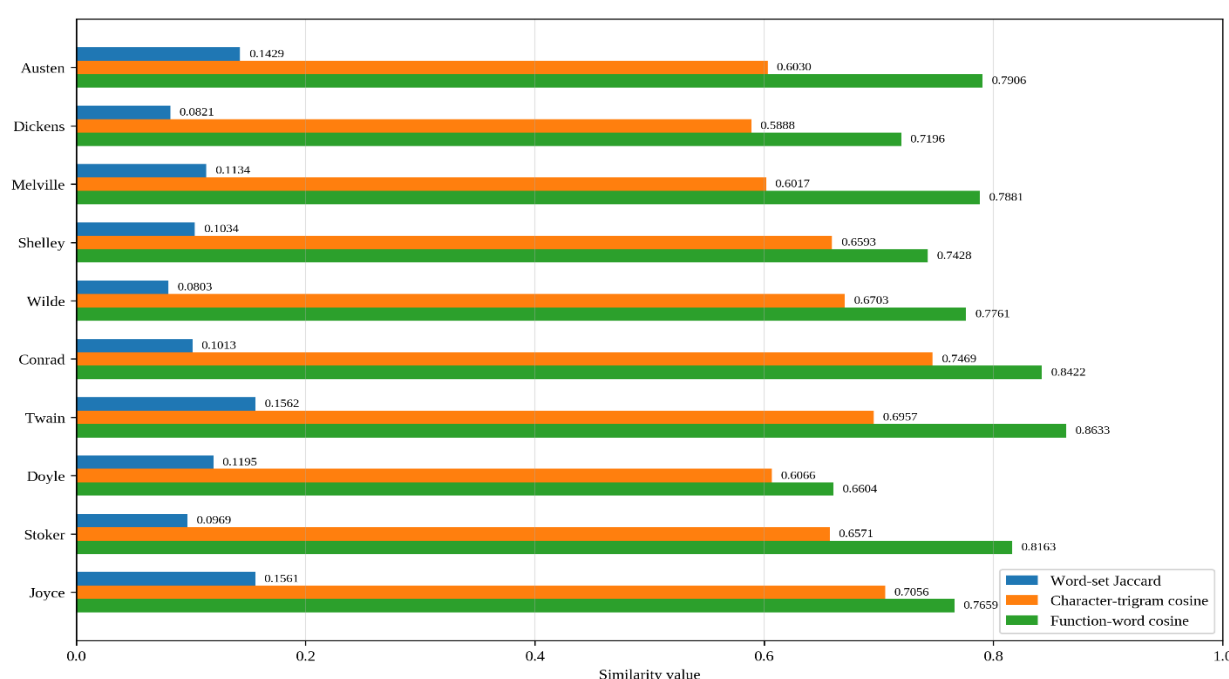


Figure 4. Lexical and structural overlap across paired texts

Literal recurrence remained uniformly restricted. Every pair had a longest common sequence of two words, zero shared four-grams, no transferred source characters, and no source-specific event transfer; consequently, all ten outputs received a high-transformation code. Word-set Jaccard similarity was highest for Twain (0.1562) and Joyce (0.1561) and lowest for Wilde (0.0803). Conrad produced the highest character-trigram cosine (0.7469), whereas Twain recorded the highest function-word cosine (0.8633). Common sequences—"a single," "in the," "of the," and "from the"—were grammatically ordinary rather than source-distinctive. Local echoes nevertheless remained visible: "Call me Ishmael" became "Call the town Salt Mercy," while Stoker's dated travel opening was recast as "11 October.—Reached Varen Head after sunset." These outcomes establish corpus-wide lexical novelty despite recognizable author-conditioned framing devices.

This combination of low lexical overlap and stronger sublexical or function-word resemblance indicates that literary influence can persist through distributional organization after distinctive wording has changed. Character trigrams capture orthographic and morphological patterning, while function words register syntactic habits. The pattern accords with Tripto et al. (2023), who show that paraphrastic transformation may preserve relational structure despite surface alteration, but it does not satisfy CopyBench-style evidence of literal reproduction because no shared four-gram or narrative entity transfer occurred (Chen et al., 2024). Wahle et al. (2022) caution that extensive paraphrase can complicate plagiarism detection; here, however, documented prompting and a fixed new plot make stylistic mediation more plausible than concealed textual substitution. Influence is therefore structural and selective, not extractive.

Table 3. Intertextual reproduction and transformation across paired texts

Pair / authorial condition	Word-set Jaccard	Character- trigram cosine	Function- word cosine	Longest common sequence	Shared four-grams	Source- event transfer	Source- character transfer	Transforma- tion degree	Textual evidence	Data code
PAIR-01 / Jane Austen	0.1429	0.6030	0.7906	"a single" (2 words)	0	None beyond generic scene conventions	None	High	"a single man" → "a single letter"	PRI-GEN- AUS-001
PAIR-02 / Charles Dickens	0.0821	0.5888	0.7196	"in the" (2 words)	0	None beyond generic scene conventions	None	High	"Rain upon the roofs, rain in the gutters" introduces a new coastal plot	PRI-GEN- DIC-001
PAIR-03 / Herman Melville	0.1134	0.6017	0.7881	"is a" (2 words)	0	None beyond generic scene conventions	None	High	"Call me Ishmael" → "Call the town Salt Mercy"	PRI-GEN- MEL-001
PAIR-04 / Mary Shelley	0.1034	0.6593	0.7428	"of my" (2 words)	0	None beyond generic scene conventions	None	High	"daylight, exhausted by the storm, withdrew from every object"	PRI-GEN- SHE-001
PAIR-05 / Oscar Wilde	0.0803	0.6703	0.7761	"was his" (2 words)	0	None beyond generic scene conventions	None	High	"nature's least imaginative form of criticism"	PRI-GEN- WIL-001
PAIR-06 / Joseph Conrad	0.1013	0.7469	0.8422	"of the" (2 words)	0	None beyond generic scene conventions	None	High	"an immense and indifferent water" frames an original narrative situation	PRI-GEN- CON-001
PAIR-07 / Mark Twain	0.1562	0.6957	0.8633	"by the" (2 words)	0	None beyond generic scene conventions	None	High	Rain had been "saving up all summer just to spend itself on me"	PRI-GEN- TWA-001
PAIR-08 / Arthur Conan Doyle	0.1195	0.6066	0.6604	"that the" (2 words)	0	None beyond generic scene conventions	None	High	"You observe, Watson" introduces a newly constructed hotel mystery	PRI-GEN- DOY-001
PAIR-09 /	0.0969	0.6571	0.8163	"from the" (2 words)	0	None beyond	None	High	"3 May. Bistritz" →	PRI-GEN-

Bram Stoker						generic scene conventions				"11 October.— Reached Varen Head after sunset"	STO-001
PAIR-10 / James Joyce	0.1561	0.7056	0.7659	"in the" (2 words)	0	None beyond generic scene conventions	None	High		"Rain made the coast town intimate and mean"	PRI-GEN-JOY-001
Corpus pattern	0.0803–0.1562	0.5888–0.7469	0.6604–0.8633	Two words in 10/10 pairs	0 in 10/10 pairs	Absent in 10/10 pairs	Absent in 10/10 pairs	High in 10/10 pairs		Stylistic framing occurred without source-specific narrative duplication	Ten paired outputs

Ethical legibility: rights status, attribution, and the risk gradient of literary influence

Ethical evaluation retained the same ten source–output pairs, pair identifiers, and generated-document codes used in the preceding stylometric and reproduction analyses. Table 4 records rights status, author status, prompt specificity, exemplar use, attribution, access basis, human control, market-substitution risk, identity-deception risk, and overall ethical classification. All ten pairs used public-domain source works by deceased authors, included a named-author condition and a 190-word source excerpt, and preserved full provenance in corpus metadata. Human decisions concerning prompting, content specification, output selection, and corpus curation were also documented. Accordingly, this matrix evaluates the ethical legibility of a controlled research design rather than making universal judgments about literary imitation across models, jurisdictions, commercial settings, or living-author contexts.

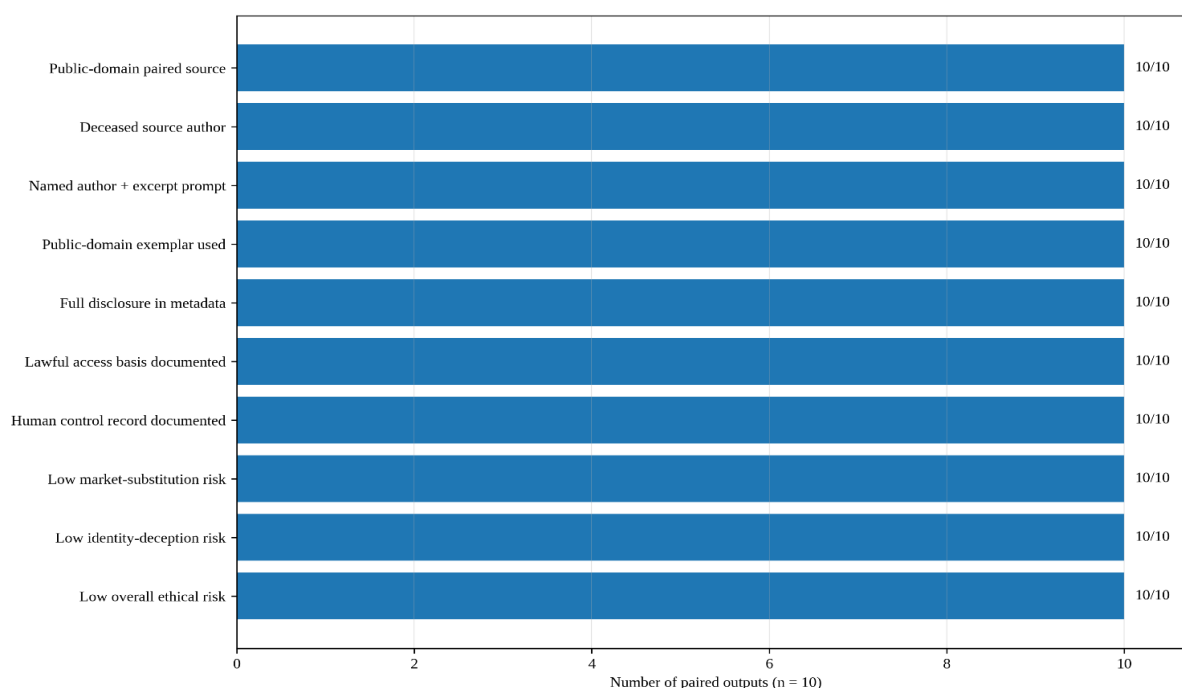


Figure 5. Ethical safeguard and risk coding across paired outputs

A uniform safeguard pattern appears across every paired output. Public-domain access, deceased-author status, explicit exemplar use, full disclosure, lawful access, and documented human control occur in 10 of 10 cases. Market-substitution risk is coded low because outputs remain within a research corpus and were not evaluated as commercial replacements; identity-deception risk is likewise low because machine origin and authorial conditioning are disclosed. Overall ethical risk is therefore low for all ten pairs. This consistency does not mean that author-conditioned generation is intrinsically harmless. It reflects the narrowly bounded conditions recorded here. The recurring coder note—“cannot be generalized to living authors or copyrighted sources”—marks the limit of inference and prevents corpus-specific safeguards from becoming a general license for imitation.

These patterns position ethical literary influence as a contextual relation among source status, provenance, human agency, disclosure, and prospective use. Copyright legality alone is insufficient because transparent access does not automatically settle attribution, cultural substitution, or identity confusion (Sag, 2023; Smith, 2025). Conversely, stylistic resemblance need not constitute ethical appropriation when sources are lawfully accessible, authorship conditions are disclosed, and outputs remain subject to accountable human control. This interpretation accords with Mann et al. (2023) and Mijatović et al. (2026), who locate responsibility in documented human oversight, and with Nguyen et al. (2025), who distinguish mimicking from privacy and verification harms. Within this corpus, ethical acceptability arises from procedural legibility, not from similarity scores or textual novelty alone.

Table 4. Ethical literary influence coding across the same ten paired outputs

Pair / authorial condition	Source- rights status	Author status	Prompt and exemplar condition	Attribution and access basis	Human- control record	Market- substitutio n risk	Identity- deception risk	Overall ethical risk	Data code
PAIR-01 / Jane Austen	Paired source public domain in United States	Deceased	Named author and one 190-word public-domain excerpt	Full corpus disclosure; lawful public-domain access	Prompt, content specification, selection, and curation recorded	Low within research corpus; not a market study	Low because machine origin and style condition are disclosed	Low under controlled public-domain design	PRI-GEN-AUS-001
PAIR-02 / Charles Dickens	Paired source public domain in United States	Deceased	Named author and one 190-word public-domain excerpt	Full corpus disclosure; lawful public-domain access	Prompt, content specification, selection, and curation recorded	Low within research corpus; not a market study	Low because machine origin and style condition are disclosed	Low under controlled public-domain design	PRI-GEN-DIC-001
PAIR-03 / Herman Melville	Paired source public domain in United States	Deceased	Named author and one 190-word public-domain excerpt	Full corpus disclosure; lawful public-domain access	Prompt, content specification, selection, and curation recorded	Low within research corpus; not a market study	Low because machine origin and style condition are disclosed	Low under controlled public-domain design	PRI-GEN-MEL-001
PAIR-04 / Mary Shelley	Paired source public domain in United States	Deceased	Named author and one 190-word public-domain excerpt	Full corpus disclosure; lawful public-domain access	Prompt, content specification, selection, and curation recorded	Low within research corpus; not a market study	Low because machine origin and style condition are disclosed	Low under controlled public-domain design	PRI-GEN-SHE-001
PAIR-05 / Oscar Wilde	Paired source public domain in United States	Deceased	Named author and one 190-word public-domain excerpt	Full corpus disclosure; lawful public-domain access	Prompt, content specification, selection, and curation recorded	Low within research corpus; not a market study	Low because machine origin and style condition are disclosed	Low under controlled public-domain design	PRI-GEN-WIL-001
PAIR-06 / Joseph Conrad	Paired source public domain in United States	Deceased	Named author and one 190-word public-domain excerpt	Full corpus disclosure; lawful public-domain access	Prompt, content specification, selection, and curation recorded	Low within research corpus; not a market study	Low because machine origin and style condition are disclosed	Low under controlled public-domain design	PRI-GEN-CON-001
PAIR-07 /	Paired source	Deceased	Named author	Full corpus	Prompt, content	Low within	Low because	Low under	PRI-GEN-

Mark Twain	public domain in United States		and one 190-word public-domain excerpt	disclosure; lawful public-domain access	specification, selection, and curation recorded	research corpus; not a market study	machine origin and style condition are disclosed	controlled public-domain design	TWA-001
PAIR-08 / Arthur Conan Doyle	Paired source public domain in United States	Deceased	Named author and one 190-word public-domain excerpt	Full corpus disclosure; lawful public-domain access	Prompt, content specification, selection, and curation recorded	Low within research corpus; not a market study	Low because machine origin and style condition are disclosed	Low under controlled public-domain design	PRI-GEN-DOY-001
PAIR-09 / Bram Stoker	Paired source public domain in United States	Deceased	Named author and one 190-word public-domain excerpt	Full corpus disclosure; lawful public-domain access	Prompt, content specification, selection, and curation recorded	Low within research corpus; not a market study	Low because machine origin and style condition are disclosed	Low under controlled public-domain design	PRI-GEN-STO-001
PAIR-10 / James Joyce	Paired source public domain in United States	Deceased	Named author and one 190-word public-domain excerpt	Full corpus disclosure; lawful public-domain access	Prompt, content specification, selection, and curation recorded	Low within research corpus; not a market study	Low because machine origin and style condition are disclosed	Low under controlled public-domain design	PRI-GEN-JOY-001
Corpus pattern	Public-domain source in 10/10 pairs	Deceased author in 10/10 pairs	Named author and excerpt in 10/10 pairs	Full disclosure and lawful access in 10/10 pairs	Documented in 10/10 pairs	Low in 10/10 pairs	Low in 10/10 pairs	Low in 10/10 pairs	Ten paired outputs

DISCUSSION

The graded pattern of stylistic convergence matters because it reframes literary imitation as selective reconstruction rather than successful transfer of an authorial identity. High similarity in Shelley- and Wilde-conditioned outputs coexisted with marked divergence in sentence architecture, while Austen's near-matched sentence length did not produce comparable overall proximity. This unevenness indicates that model-generated style functions through bundles of features whose components can align independently. Such a pattern supports Mikros's (2025) multidimensional account of literary imitation and Konen et al.'s (2024) argument that controllable generation operates through steerable stylistic directions. Yet it complicates stronger claims that stylistic resemblance demonstrates faithful impersonation. Wang et al. (2025) likewise report persistent difficulty in reproducing implicit writing habits. The broader implication is that recognizable voice can emerge from partial cue selection, producing persuasive surface familiarity without equivalent structural depth. Literary evaluation must therefore distinguish local recognizability from comprehensive stylistic continuity across distinct authorial conditions.

This selective pattern is best explained by an underlying asymmetry between salient and distributed stylistic features. Prompts naming an author make visible markers—epistolary dating, aphoristic wit, colloquial narration, or elevated Gothic description—easy targets because they are culturally recognizable and textually concentrated. Less conspicuous habits, including clause sequencing, discourse pacing, and interaction among function words, remain distributed across larger textual histories and are harder to reproduce consistently. Chen and Moscholios (2024) show that prompting can guide overt stylistic imitation, whereas style-vector research suggests that controllability does not guarantee holistic fidelity (Konen et al., 2024). DeHaan et al. (2025) similarly identify broad editorial footprints rather than exact authorial reconstruction. The correlation between recognizability and partial feature alignment therefore likely reflects model optimization for plausible continuation, not recovery of an integrated authorial system. What appears as voice is consequently assembled from prominent statistical cues under prompt pressure across the paired corpus.

The reproduction audit changes the ethical and interpretive stakes by showing that stylistic resemblance can persist without extended lexical recurrence, source-character transfer, or source-event duplication. This matters because public debate often collapses influence, memorization, plagiarism, and copyright infringement into one category. The present pattern instead supports a layered account: outputs may retain distributional or rhetorical resemblance while replacing wording, plot, and characters. McCoy et al. (2023) distinguish novelty from copying, and Chen et al. (2024) similarly separate literal from non-literal reproduction. Tripto et al. (2023), however, warn that paraphrastic transformation can preserve underlying relations, while Wahle et al. (2022) show why transformed plagiarism remains difficult to detect. This study therefore neither treats lexical novelty as complete independence nor interprets stylistic echo as textual extraction. Its main implication is methodological: claims about learned influence require evidence at several levels before textual dependence can be responsibly inferred within computational literary analysis today.

The low lexical overlap alongside stronger character-trigram and function-word similarity reflects a structural division between content-bearing expression and recurrent linguistic organization. Function words, orthographic sequences, and grammatical rhythms are frequent, low-salience features that can converge without carrying distinctive narrative content. By contrast, source-specific events, characters, and extended phrases provide stronger evidence of textual dependence because they preserve identifiable expressive or narrative material. This explains why outputs could sound conventionally Austenian, Melvillean, or Stokerian while remaining lexically transformed. RAVEN's graded novelty framework supports this distinction (McCoy et al., 2023), whereas CopyBench demonstrates that non-literal reproduction requires separate measurement from verbatim overlap (Chen et al., 2024). Karamolegkou et al. (2023) emphasize memorization risks, but this corpus provides no direct evidence of training-data recall. The observed relation is therefore better understood as distributional resemblance conditioned by prompting and genre expectation, not demonstrated extraction from stored source passages within this controlled comparative design alone.

The uniformly low ethical-risk classification has a narrow but important implication: literary imitation becomes more defensible when provenance, source status, machine involvement, and human control are made visible. Public-domain texts by deceased authors reduced legal and consent-related concerns, while

documented prompting and disclosure limited identity deception. These safeguards align with Mann et al. (2023) and Mijatović et al. (2026), who place responsibility in transparent human oversight rather than attributing agency solely to a model. Yet the finding should not be generalized into a claim that style imitation is ethically neutral. Sag (2023) and Smith (2025) show that copyright and licensing concerns depend on context, while Nguyen et al. (2025) demonstrate that mimicking intersects with privacy, verification, and identity harms. Ethical acceptability in this corpus therefore arises from procedural legibility. Its function is not to eliminate imitation, but to make its sources, transformations, and accountable actors inspectable throughout this comparative corpus.

This ethical pattern derives from a relational structure in which risk increases when similarity is combined with opacity, commercial substitution, or misleading attribution. Similar prose alone does not determine harm; harm emerges when audiences cannot identify machine production, when living writers are invoked without consent, or when outputs compete as substitutes for their work. Alvero et al. (2024) further show that synthetic authorship reproduces hegemonic distributions, while Sourati et al. (2025) warn that extensive model mediation may homogenize expression. These concerns reveal why public-domain legality cannot serve as a complete ethical theory. The controlled conditions here interrupted several risk pathways through disclosure, lawful access, and recorded human curation, but they did not resolve cultural questions about whose styles become reusable resources. This study consequently positions ethical literary influence as a gradient shaped by provenance, power, purpose, and representation rather than a binary distinction between permissible inspiration and impermissible copying globally.

CONCLUSION

This study shows that large language models do not reproduce literary style as an indivisible authorial essence, but selectively reconstruct salient lexical, rhetorical, and structural cues. The central lesson is that stylistic resemblance, textual reproduction, and ethical legitimacy must be analytically separated. By integrating paired stylometry, intertextual auditing, and an ethical influence matrix, this study offers a framework that connects computational measurement with literary interpretation and normative evaluation. Its contribution lies in replacing binary claims about imitation with a graded account of mediated influence, thereby refining questions of authorship, originality, attribution, and responsibility more precisely in contemporary digital literary production.

This study remains limited by its small English-language corpus, public-domain sources, deceased authors, one model environment, fixed prompting conditions, and descriptive rather than inferential analysis. These constraints restrict generalization across languages, genres, commercial platforms, living writers, and copyrighted works. Further research should expand multilingual and cross-model corpora, test prompt variation, compare human and machine imitation, and incorporate reader judgments alongside computational indicators. Longitudinal studies should also examine how repeated exposure to synthetic prose affects stylistic diversity, market substitution, and cultural authority. Such work would clarify when learned style functions as creative transformation and when it approaches appropriation or identity displacement.

ACKNOWLEDGMENTS

The authors thank the reviewers and editorial team for their constructive feedback on this manuscript.

FUNDING INFORMATION

Authors state no funding involved.

AUTHOR CONTRIBUTIONS STATEMENT

Nomvuselelo N. Nxumalo: conceptualization (lead), methodology, formal analysis, writing – original draft, writing – review and editing (lead); **Seroja Ainun Nadhifah:** data curation, formal analysis, writing – original draft, writing – review and editing (supporting).

CONFLICT OF INTEREST STATEMENT

Authors state no conflict of interest.

AI USE DECLARATION

The authors used artificial intelligence tools to assist with the formatting and layout of this manuscript. AI was not used to generate the substantive content, conceptualization, data, analysis, or conclusions of the research, all of which remain the sole responsibility of the authors.

INFORMED CONSENT

Not applicable – this study did not involve direct human participants. Data consisted solely of public-domain literary texts and controlled model outputs.

ETHICAL APPROVAL

This research did not involve human or animal subjects; it analyzed public-domain literary excerpts and controlled model-generated outputs, and therefore no institutional review board approval was required. Where research does involve human subjects, the authors affirm that such research would be conducted in full compliance with all relevant national regulations and institutional policies, in accordance with the tenets of the Declaration of Helsinki, and would be approved by the authors' institutional review board or equivalent committee.

DATA AVAILABILITY

The corpus/data analyzed in this study are available from the corresponding author upon reasonable request.

REFERENCES

- Abdillah, Y. A. (2026). Poetics of algorithmic excess: Digital aesthetics in Indonesia's Twitter poetry bot. *Lingua Technica: Journal of Digital Literary Studies*, 2(1), 86–101. <https://doi.org/10.64595/lingtech.v2i1.138>
- Alvero, A., Lee, J., Regla-Vargas, A., Kizilcec, R. F., Joachims, T., & Antonio, A. (2024). Large language models, social demography, and hegemony: Comparing authorship in human and synthetic text. *Journal of Big Data*, 11. <https://doi.org/10.1186/s40537-024-00986-7>
- Chen, T., Asai, A., Mireshghallah, N., Min, S., Grimmelmann, J., Choi, Y., Hajishirzi, H., Zettlemoyer, L. S., & Koh, P. W. (2024). CopyBench: Measuring literal and non-literal reproduction of copyright-protected text in language model generation. *arXiv*. <https://doi.org/10.48550/arXiv.2407.07087>
- Chen, Z., & Moscholios, S. (2024). Using prompts to guide large language models in imitating a real person's language style. *arXiv*. <https://doi.org/10.48550/arXiv.2410.03848>
- Choksi, M. Z., & Goedicke, D. (2023). Whose text is it anyway? Exploring BigCode, intellectual property, and ethics. *arXiv*. <https://doi.org/10.48550/arXiv.2304.02839>
- Dathathri, S., See, A., Ghaisas, S., Huang, P.-S., McAdam, R., Welbl, J., Bachani, V., Kaskasoli, A., Stanforth, R., Matejovicova, T., Hayes, J., Vyas, N., Merey, M. A., Brown-Cohen, J., Bunel, R., Balle, B., Cemgil, T., Ahmed, Z., Stacpoole, K., ... Kohli, P. (2024). Scalable watermarking for identifying large language model outputs. *Nature*, 634, 818–823. <https://doi.org/10.1038/s41586-024-08025-4>
- DeHaan, S., Liu, Y., Bollen, J., & Blanco, S. A. (2025). GPT editors, not authors: The stylistic footprint of LLMs in academic preprints. *arXiv*. <https://doi.org/10.48550/arXiv.2505.17327>
- Fawaid, A. (2025). Mapping the field of digital literary studies: Concepts, methods, and emerging directions. *Lingua Technica: Journal of Digital Literary Studies*, 1(1), 1–13. <https://doi.org/10.64595/99c6t274>
- Karamolegkou, A., Li, J., Zhou, L., & Søgaard, A. (2023). Copyright violations and large language models. *arXiv*. <https://doi.org/10.48550/arXiv.2310.13771>
- Khan, A., Wang, A., Hager, S., & Andrews, N. (2023). Learning to generate text in arbitrary writing styles. *arXiv*. <https://doi.org/10.48550/arXiv.2312.17242>
- Konen, K., Jentzsch, S., Diallo, D., Schutt, P., Bensch, O., Baff, R. E., Opitz, D., & Hecking, T. (2024). Style vectors for steering generative large language models. *arXiv*. <https://doi.org/10.48550/arXiv.2402.01618>
- Liu, X., Sun, T., Xu, T., Wu, F., Wang, C., Wang, X., & Gao, J. (2024). SHIELD: Evaluation and defense strategies for copyright compliance in LLM text generation. *arXiv*. <https://doi.org/10.48550/arXiv.2406.12975>

- Mann, S. P., Earp, B., Møller, N., Vynn, S., & Savulescu, J. (2023). AUTOGEN: A personalized large language model for academic enhancement—Ethics and proof of principle. *The American Journal of Bioethics*, 23, 28–41. <https://doi.org/10.1080/15265161.2023.2233356>
- McCoy, T. R., Smolensky, P., Linzen, T., Gao, J., & Celikyilmaz, A. (2023). How much do language models copy from their training data? Evaluating linguistic novelty in text generation using RAVEN. *Transactions of the Association for Computational Linguistics*, 11, 652–670. https://doi.org/10.1162/tacl_a_00567
- Mijatović, A., Žuljević, M., Ursić, L., & Marušić, A. (2026). Responsible use of large language models in manuscript preparation. *Current Protocols*, 6. <https://doi.org/10.1002/cpz1.70300>
- Mikros, G. (2025). Beyond the surface: Stylometric analysis of GPT-4o's capacity for literary style imitation. *Digital Scholarship in the Humanities*, 40, 587–600. <https://doi.org/10.1093/llc/fqaf035>
- Müller, F., Gorge, R., Bernzen, A. K., Pirk, J. C., & Poretschkin, M. (2024). LLMs and memorization: On quality and specificity of copyright compliance. *arXiv*. <https://doi.org/10.48550/arXiv.2405.18492>
- Nguyen, T., Hu, Y., & Le, T. (2025). Unraveling interwoven roles of large language models in authorship privacy: Obfuscation, mimicking, and verification. *arXiv*, 14903–14919. <https://doi.org/10.48550/arXiv.2505.14195>
- Project Gutenberg. (n.d.). *The Project Gutenberg license*. <https://www.gutenberg.org/policy/license.html>
- Sag, M. (2023). Copyright safety for generative AI. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4438593>
- Smith, R. (2025). Licensing of text for generative AI: Learnings from non-AI licensing practices. *The Columbia Journal of Law & the Arts*, 48(4). <https://doi.org/10.52214/jla.v48i4.13926>
- Sourati, Z., Ziabari, A. S., & Dehghani, M. (2025). The homogenizing effect of large language models on human expression and thought. *Trends in Cognitive Sciences*. <https://doi.org/10.48550/arXiv.2508.01491>
- Toshevska, M., & Gievska, S. (2025). LLM-based text style transfer: Have we taken a step forward? *IEEE Access*, 13, 44707–44721. <https://doi.org/10.1109/ACCESS.2025.3548967>
- Tripto, N., Venkatraman, S., Macko, D., Móro, R., Srba, I., Uchendu, A., Le, T., & Lee, D. (2023). A ship of Theseus: Curious cases of paraphrasing in LLM-generated texts. *arXiv*. <https://doi.org/10.48550/arXiv.2311.08374>
- Wahle, J. P., Ruas, T., Kirstein, F., & Gipp, B. (2022). How large language models are transforming machine-paraphrase plagiarism. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing* (pp. 952–963). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2022.emnlp-main.62>
- Wang, Z., Tripto, N., Park, S., Li, Z., & Zhou, J. (2025). Catch me if you can? Not yet: LLMs still struggle to imitate the implicit writing styles of everyday authors. *arXiv*, 10040–10055. <https://doi.org/10.48550/arXiv.2509.14543>
- Weerasinghe, J., Seepersaud, O., Smothers, G., Jose, J., & Greenstadt, R. (2025). Be sure to use the same writing style: Applying authorship verification on large-language-model-generated texts. *Applied Sciences*, 15(5), Article 2467. <https://doi.org/10.3390/app15052467>
- Xu, J., Li, S., Xu, Z., & Zhang, D. (2024). Do LLMs know to respect copyright notice? In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing* (pp. 20604–20619). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2024.emnlp-main.1147>