



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>



Prompting as poetic practice: human–AI collaboration and the making of algorithmic literary voice

Muhammad Khairul Umam ^{1*} 

¹ International Islamic University, Malaysia

* Corresponding author. Email: umam.mk@live.iium.edu.my

ABSTRACT

Article history:

Received May 22, 2026

Revised Jun 24, 2026

Accepted Jun 30, 2026

Published Jul 31, 2026

Keywords:

algorithmic literary voice;

digital poetry;

human–AI collaboration;

poetic practice; prompting

Background: Generative language models have transformed poetry writing into an iterative practice in which prompts, outputs, revisions, interfaces, and platforms jointly shape literary production. **Objective:** This study examines how prompting functions poetically, how creative agency is negotiated during human–AI collaboration, and how algorithmic literary voice emerges across documented digital writing practices. **Method:** A qualitative multiple-case corpus design analyzed 34 verified public documents, 41 text units, and 41 source-linked coding units through a Prompt–Rhetoric Coding Matrix, Human–AI Negotiation Sequence Analysis, and Algorithmic Literary Voice Profile. **Results:** Dialogic and evaluative prompting formed the largest prompting configuration, indicating that poetic composition depended on reformulation, assessment, and constraint rather than single-turn instruction. Human-directed control remained prominent across interaction sequences, although automation, platform mediation, and distributed participation produced several asymmetrical forms of collaboration. Hybrid and human-curated voice profiles exceeded purely model-conditioned patterns, demonstrating that literary voice developed through combined traces of prompting, editing, stylistic recurrence, attribution, code, and circulation. **Implication:** These findings reposition authorship as a traceable distribution of compositional decisions rather than a binary division between human and machine production. **Novelty:** This study integrates prompt rhetoric, interactional agency, and platform-mediated stylistics within one auditable framework for analyzing AI-assisted poetry in contemporary digital literary culture.

DOI: <https://doi.org/10.64595/kz2knf55>

*) Copyright © 2026 Muhammad Khairul Umam

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

Artificial intelligence has transformed literary production from experimentation to repeatable, instrumented collaboration. CoAuthor records 1,445 English writing sessions involving 63 writers and four GPT-3 configurations, demonstrating that interaction with language models can now be studied at the level of prompts, selections, rejections, and revisions rather than through isolated outputs (Lee et al., 2022). In a complementary professional setting, thirteen published writers used an AI-assisted editor for brainstorming, world-building, research, and textual generation, while reporting difficulty preserving stylistic control (Ippolito et al., 2022). Reception has also become unstable: across 16,340 poem judgments, participants identified AI-generated poems as human-authored more often than poems written by canonical poets, and frequently rated machine poems more favorably on accessible aesthetic qualities (Porter & Machery, 2024). These developments make prompting a consequential literary practice because it organizes creative intention, delegates linguistic decisions, and obscures conventional boundaries among instruction, composition, editing, and authorship within digital culture.

Existing research has established several foundations for examining AI-assisted literary production. Chakrabarty, Padmakumar, and He (2022) developed CoPoet to support attribute-controlled and iterative poetry writing, while Dhillon et al. (2024) showed that different levels of linguistic scaffolding reshape perceived contribution and ownership in co-writing. Henrickson and Meroño-Peñuela (2023) conceptualized prompting as hermeneutic meaning negotiation, and Guo et al. (2025) found that creative writers regulate AI use through commitments to authenticity and craftsmanship. Studies of algorithmic style have examined imitation, recurrent model signatures, and reader attribution (Mikros, 2025; Porter & Machery, 2024), whereas Indonesian digital-literature scholarship has foregrounded platform poetics, procedural repetition, and distributed authorship (Abdillah, 2026). However, these strands remain insufficiently integrated. Prompting is usually treated as technical optimization, interface behaviour, or preparatory instruction, not as a poetic practice whose rhetorical structure, interactional sequencing, and textual consequences jointly participate in producing an identifiable algorithmic literary voice across different media environments.

This study investigates prompting as a situated compositional practice through a publicly verifiable corpus of prompt templates, interaction records, repositories, interface documentation, and scholarship on AI-mediated writing. It asks three connected questions. First, how do prompts translate thematic intentions, formal constraints, stylistic models, and multimodal cues into poetic operations? Second, how is creative agency negotiated across sequences of generation, selection, rejection, redirection, and revision? Third, which recurrent lexical, formal, interactional, and platform traces justify describing an emergent algorithmic literary voice without attributing autonomous authorship to a language model? To answer these questions, this study combines a Prompt–Rhetoric Coding Matrix, Human–AI Negotiation Sequence Analysis, and an Algorithmic Literary Voice Profile. This design treats prompts, responses, revisions, and interfaces as analytically distinct but relational units. It consequently carefully separates observable textual intervention from reported perceptions of collaboration and distinguishes prompt compliance, stylistic regularity, literary evaluation, and authorship attribution as different evidentiary problems.

This study advances the provisional argument that algorithmic literary voice does not reside exclusively in model output but emerges from a distributed arrangement of human instruction, computational generation, editorial intervention, and platform mediation. Prompting therefore functions poetically when it organizes imagery, rhythm, perspective, indirection, and revision as compositional choices rather than merely optimizing task completion. This position extends Gunkel's (2025) distinction between being prompted by a human and generated by ChatGPT, while retaining Elam's (2023) warning that literature cannot be reduced to fluency, efficiency, or instruction compliance. It also aligns with Nadifah's (2025) account of language, code, and platform as jointly constitutive of electronic textuality. Empirically, this proposition will be examined by comparing prompt functions, negotiation sequences, and recurrent stylistic traces across documented cases. Its broader implication is that literary agency should be assessed through traceable contributions and transformations, not through simplistic binary classifications of human versus machine authorship alone.

LITERATURE REVIEW

Prompting in literary production

Prompting in literary production can be understood either as instruction design or as an interpretive act that shapes meaning before generation begins. Engineering-oriented accounts emphasize precision, constraint, and task optimization (Bozkurt, 2024; Park & Choo, 2024), whereas hermeneutic and rhetorical approaches treat prompts as negotiated forms of intention, framing, and revision (Henrickson & Meroño-Peñuela, 2023; Ranade, Saravia, & Johri, 2024). In poetry, this distinction is decisive because prompts may regulate imagery, rhythm, voice, and ambiguity rather than merely improve correctness. Accordingly, prompting is best conceptualized as a compositional practice mediating between human aesthetic judgment and model-generated language within creative processes.

Its principal indicators include semantic direction, formal constraint, stylistic conditioning, iterative reformulation, and multimodal specification. Topic prompts establish conceptual fields; formal prompts impose rhyme, meter, stanza, or lexical endings; stylistic prompts invoke genre, authorial manner, or register; iterative prompts redirect unsatisfactory outputs; and multimodal prompts translate poetic meaning across image, sound, or screen design (Chakrabarty, Padmakumar, & He, 2022; Dang et al., 2022; Jamil et al., 2025). Yet specificity does not guarantee literary value, since excessive constraint may produce mechanical compliance. Prompt analysis must therefore distinguish instruction complexity, rhetorical function, poetic affordance, and openness to productive deviation within each episode.

Human–AI collaboration

Human–AI collaboration is variously defined as assistance, co-creation, augmentation, or distributed authorship. Assistance models preserve human primacy by treating AI as a tool for ideation or revision (Ippolito et al., 2022; Reza et al., 2025), while co-creation models emphasize reciprocal influence and shared contribution (McGuire, De Cremer, & Van De Cruys, 2024). Distributed approaches extend agency across writers, models, interfaces, and platforms (Gunkel, 2025). These definitions remain contested because contribution, control, ownership, and responsibility rarely coincide. Collaboration should therefore be examined empirically through observable interactional sequences and documented textual interventions rather than declared through broad, celebratory labels alone in practice.

Analytically, collaboration can be traced through initiation, generation, selection, rejection, redirection, editing, and integration. CoAuthor and related co-writing research demonstrate that agency appears in micro-decisions about when to request suggestions, which outputs to retain, and how extensively to revise them (Lee, Liang, & Yang, 2022; Dhillon et al., 2024). Team-based prompting also distributes planning and evaluation among multiple human participants (Han et al., 2024), while creative writers regulate AI involvement according to authenticity and craftsmanship (Guo et al., 2024). Relevant indicators therefore include contribution scale, revision depth, acceptance patterns, provenance visibility, and whether human judgment remains decisive across documented episodes.

Algorithmic literary voice

Algorithmic literary voice refers neither to autonomous machine personality nor to simple stylistic imitation. One perspective identifies recurring lexical, syntactic, rhythmic, and structural regularities across model outputs (Mikros, 2025). Another emphasizes hybrid voice produced through prompt conditioning, human editing, and model completion (Huang et al., 2025; Prabowo & Asmarani, 2025). Platform-oriented digital poetics further locates voice in code, interface, circulation, and reader interaction (Nadifah, 2025; Abdillah, 2026). These approaches differ over where voice resides, but collectively reject purely individualist authorship. Algorithmic voice is therefore relational, emerging from specific patterned interactions among linguistic systems, human choices, computational procedures, and technological environments.

Its indicators should include recurrent diction, stanzaic regularity, perspective, figurative patterning, stylistic imitation, revision traces, attribution signals, and platform effects. Reception studies show that readers may misidentify AI poetry or respond differently when authorship is labelled (Porter & Machery, 2024; Stańko-Kaczmarek, Dera, & Koscielska, 2024), demonstrating that perceived voice and production provenance are distinct. Digital literary scholarship likewise shows that multimodality and platform affordances influence

textual meaning (Assyabani, 2025; Atikurrahman, 2025). This study consequently theorizes algorithmic literary voice as a layered profile integrating prompt rhetoric, negotiated agency, stylistic recurrence, and platform mediation across publicly documented cases and materially consequential contexts.

METHOD

This study examines publicly accessible artifacts documenting poetic interaction between human writers and generative systems. Its material units comprise prompt templates, prompting scripts, interface records, model outputs, revision traces, repositories, metadata tables, and published corpus descriptions. The finalized corpus contains 34 verified documents: 12 primary records, 10 secondary empirical studies, and 12 contextual sources. These documents yield 41 text units and 41 preliminary coding units, each connected to a unique identifier, source locator, canonical link, and data category. Table 1 summarizes corpus composition, while individual prompts, interaction episodes, and comparable output sets constitute the principal analytical units for subsequent interpretation.

A qualitative multiple-case corpus design was adopted because this study compares heterogeneous but traceable instances of poetic prompting across platforms, interfaces, and publication contexts. Rather than treating generated poems as isolated products, this design reconstructs relations among instructions, outputs, selections, rejections, edits, and technological affordances. Case comparison enables patterns to be identified without assuming that all systems, genres, or users operate identically. This approach draws on human–AI writing research that treats collaboration as processual and interface-dependent (Lee, Liang, & Yang, 2022; Dhillon et al., 2024). Accordingly, claims remain interpretive, comparative, and bounded by documented public evidence rather than causal alone.

Information was obtained from official scholarly publishers, university repositories, journal platforms, project websites, and author-maintained code repositories. Primary sources included ACL Anthology, Stanford CoAuthor, CEUR Workshop Proceedings, DergiPark, GitHub repositories linked from publications, and Lingua Technica records. Secondary evidence came from empirical studies of creative writing, scaffolding, professional authorship, collaboration, and reception. Contextual materials addressed hermeneutic prompting, algorithmic authorship, stylometry, platform textuality, and Indonesian digital poetics (Henrickson & Meroño-Peñuela, 2023; Gunkel, 2025; Mikros, 2025; Nadifah, 2025; Abdillah, 2026). Official DOI pages and canonical repositories were systematically prioritized over derivative summaries, commercial descriptions, or reposted content circulated elsewhere online through sources.

Data collection proceeded through predefined keywords covering poetry prompting, collaborative writing, prompt engineering, algorithmic voice, style imitation, platform poetics, and human–AI revision. Records published between 2022 and 2026 were prioritized, while earlier materials were retained only when theoretically indispensable. Each item was screened for public accessibility, identifiable provenance, relevance, and sufficient metadata. Titles, creators, dates, institutions, document types, languages, geographic scope, prompts, short quotations, contextual notes, visual descriptions, and links were then entered into structured Excel sheets. DOI and canonical URL normalization supported deduplication, while papers, datasets, interfaces, and repositories from one project remained separate when analytically distinct and verifiable.

Analysis followed three connected stages, illustrated in Figure 2. First, a Prompt–Rhetoric Coding Matrix classified semantic direction, formal restriction, stylistic conditioning, negative constraint, multimodal specification, and evaluation criteria. Second, Human–AI Negotiation Sequence Analysis ordered initiating, generating, selecting, rejecting, redirecting, editing, and integrating moves within documented episodes. Third, an Algorithmic Literary Voice Profile examined recurrent diction, stanzaic organization, perspective, imitation, revision traces, attribution, and platform mediation. Coding distinguished observable evidence from interpretation and participant perception. Source-to-code links preserved auditability, while provisional categories were reviewed against operational definitions. This layered procedure supports triangulation without equating compliance, regularity, reception, or authorship as equivalent.

Table 1. Dataset composition

Code	Data Type	Source	Location	Period	Records	Principal Characteristics
DS-01	Primary	CoPoet paper and repository	Global—ACL Anthology and GitHub	2022	2	Prompt examples, interface workflow, instruction data, and executable code
DS-02	Primary	ChatGPT poetry corpus paper and prompting script	United States—CEUR and GitHub	2024	2	Parameterized topic-form prompts, annotated outputs, and reproducible script
DS-03	Primary	Stanford CoAuthor portal and metadata	United States—Stanford University	2021–2022	2	Session replays, prompts, model suggestions, acceptance, dismissal, and editing actions
DS-04	Primary	Poetry-translation prompting study	Türkiye—DergiPark	2025	1	Seven English–Turkish poetry-translation prompts and evaluation dimensions
DS-05	Primary	Poem-to-image prompt-tuning study	Global—ACL Anthology	2025	1	Emotion, visual-element, thematic, and cross-modal prompt components
DS-06	Primary	Poetry-generation and style repositories	Global—GitHub and arXiv	2023–2025	2	Author-conditioned datasets, prompt patterns, generation code, and evaluation resources
DS-07	Primary	Indonesian Twitter/X poetry-bot study	Indonesia—Lingua Technica	2026	1	Published description of 240 bot-generated poems and platform-based poetic production
DS-08	Primary	Human–AI poetic-dialogue study	East Asia—MDPI	2025	1	Iterative poetic dialogue developed through environmental storytelling
DS-09	Secondary	Empirical human–AI writing studies	Global—ACM, arXiv, Nature, MDPI, SAGE, and Springer	2022–2025	10	Creative writers, scaffolding, team prompting, co-creation, reception, and workflow
DS-10	Context	Authorship, hermeneutics, reception, and stylometry	Global—scholarly publishers	2023–2025	6	Prompt meaning, attribution, originality, stylistic imitation, and literary value
DS-11	Context	Indonesian digital-literature studies	Indonesia—Lingua Technica	2025–2026	5	Digital textuality, multimodality, platform mediation, and generative literature
DS-12	Context	AI-poetry reception study	Global—Wiley	2024	1	Authorship attribution, reader perception, and poetic evaluation

RESULTS

Poetic functions of prompting

Prompting emerged as a structured literary operation rather than a single technical request. Thirty-two eligible coding units were assigned exclusively to six main categories and thirteen subcategories through the Prompt–Rhetoric Coding Matrix. Dialogic and evaluative negotiation formed the largest category, followed by formal, stylistic, and rhetorical constraint. Semantic direction, multimodal and procedural prompting, technical parameterization, and authorship framing supplied complementary layers of practice. This distribution indicates that public artifacts document prompting at several levels: lexical instruction, formal composition, interactional revision, executable configuration, and metadiscursive attribution. Table 2 presents both category totals and subcategory evidence so that numerical patterns remain connected to identifiable prompts, interfaces, repositories, and conceptual records. Figure 3 visualizes only the six category totals derived from the same 32 units.

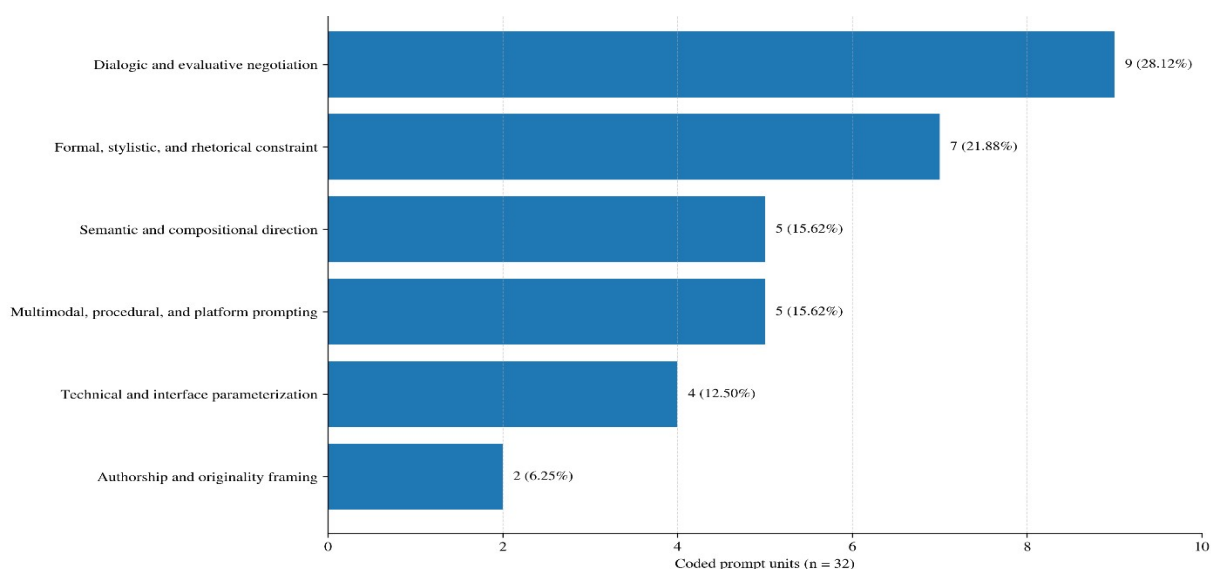


Figure 3. Distribution of poetic prompting practices

Dialogic and evaluative practices accounted for 9 units (28.125%), including iterative reformulation, alternative selection, explicit assessment, and stage-sensitive calibration. Their prominence is visible in formulations such as “Prompting is an iterative hermeneutic negotiation of meaning” (TXT-033) and in interface sequences where writers request, accept, dismiss, or edit suggestions (TXT-003; TXT-008). Formal, stylistic, and rhetorical constraints comprised 7 units (21.875%). Direct examples include “Write a sentence ending in ‘fly’” (TXT-002) and “Do not use the actual word(s) X or Y” (TXT-005). Semantic and compositional direction and multimodal, procedural, and platform prompting each contributed 5 units (15.625%). Technical and interface parameterization supplied 4 units (12.500%), whereas explicit authorship and originality framing appeared in 2 units (6.250%).

These patterns position poetic prompting between constraint-based composition and negotiated interpretation. Formal instructions resemble attribute controls described by Chakrabarty, Padmakumar, and He (2022), but their operation cannot be reduced to prompt optimization because dialogic redirection and value-based selection remain central. Henrickson and Meroño-Peñuela’s (2023) hermeneutic account clarifies why a prompt acquires meaning through response and reformulation, while Gunkel (2025) helps distinguish initiation from generation in authorship claims. Technical templates and platform procedures also support Nadifah’s (2025) argument that language, code, and interface jointly organize electronic textuality. In an Indonesian bot context, repetition and fragmentation become platform-mediated aesthetic features rather than merely compositional defects (Abdillah, 2026). Prompting therefore operates as a layered poetic practice whose agency remains distributed but traceable.

Table 2. Distribution of prompting practices across coded corpus units

Main Category (Category Total)	Subcategory	Operational Indicators	Frequency	Percentage	Data Variation	Evidence or Textual Marker	Data Codes
Semantic and compositional direction (5; 15.625%)	Topical and ideational direction	Explicit topic specification; brainstorming; imagery and world-building prompts	2	6.250%	Primary; Secondary	"Write a sentence about 'love'; brainstorming and world-building assistance	TXT-001; TXT-020
	Topic-form, translation, and scaffold framing	Combination of subject and form; translation strategy; sentence- or paragraph-level assistance	3	9.375%	Primary; Secondary	"Write a poem about the subject of X in the following form or style: Y"	TXT-004; TXT-011; TXT-022
Formal, stylistic, and rhetorical constraint (7; 21.875%)	Formal and lexical control	Terminal-word requirement; named poetic form; lexical prohibition	3	9.375%	Primary	"Write a sentence ending in 'fly'; 'Do not use the actual word(s) X or Y"	TXT-002; TXT-005; TXT-006
	Stylistic and rhetorical transformation	Author conditioning; adaptation; rewriting; stylistic imitation	3	9.375%	Primary; Context	Author-organized datasets; prompt patterns involving "adaptations and rewrites"	TXT-014; TXT-017; TXT-035
	Literary-value resistance	Resistance to instrumental optimization; preservation of ambiguity and literary excess	1	3.125%	Context	"Poetry resists reduction to optimization"	TXT-030
Dialogic and evaluative negotiation (9; 28.125%)	Metapoetic and dialogic steering	Selection among alternatives; suggestion requests; redirection; iterative interpretation	4	12.500%	Primary; Context	"Prompting is an iterative hermeneutic negotiation of meaning"	TXT-003; TXT-008; TXT-018; TXT-033
	Evaluative and stage-sensitive calibration	Creativity criteria; responsiveness; authenticity; stage-aligned AI involvement	3	9.375%	Primary; Secondary	"Creativity, prompt responsiveness, and literary value"	TXT-012; TXT-021; TXT-023
	Distributed prompt design and methodological integration	Collaborative planning; testing; refinement; combined close and computational reading	2	6.250%	Secondary; Context	Team-based prompting and integrated analytical procedures	TXT-024; TXT-037

Technical and interface parameterization (4; 12.500%)	Automated and parameterized construction	Reusable templates; model settings; prompt variables; executable constraints	3	9.375%	Primary; Context	Programmatic assembly of topic, form, constraint, temperature, and model conditions	TXT-007; TXT-010; TXT-038
	Interface-mediated prompting	Suggestion activation; provenance marking; visual distinction of human and model text	1	3.125%	Primary	Human and model contributions are differentiated through interface position and colour	TXT-009
Multimodal, procedural, and platform prompting (5; 15.625%)	Multimodal specification	Emotion, visual element, theme, image, typography, and layout coordination	2	6.250%	Primary; Context	"Prompts encode emotions, visual elements, and themes"	TXT-013; TXT-039
	Procedural and platform generation	Rule-based recombination; bot circulation; repetition; fragmentation; absent conversational prompting	3	9.375%	Primary; Context	"Algorithmic repetition and structural fragmentation function as dominant formal strategies"	TXT-015; TXT-016; TXT-040
Authorship and originality framing (2; 6.250%)	Attribution and originality differentiation	Separation of prompting, generation, prompt originality, and output originality	2	6.250%	Context	"Prompted by me; generated by ChatGPT"	TXT-031; TXT-034

Negotiating agency across prompt-response-revision sequences

Creative agency appeared through observable decisions made before, during, and after model generation. Forty eligible units were classified using Human-AI Negotiation Sequence Analysis, preserving continuity with the source-linked coding units employed for prompt analysis. Human-directed control and evaluation constituted the largest configuration, while automated production, iterative negotiation, distributed collaboration, and instrumented comparison represented distinct interactional arrangements. Such classification avoids treating every use of generative technology as equivalent collaboration. A model could function as a conditional generator, alternative provider, ideation assistant, stylistic imitator, procedural system, or object of evaluation, depending on the documented sequence. Table 3 therefore connects numerical distributions with actor roles, revision status, uptake patterns, and evidence spans. Figure 4 preserves these frequencies while displaying how each main category is composed of mutually exclusive subcategories.

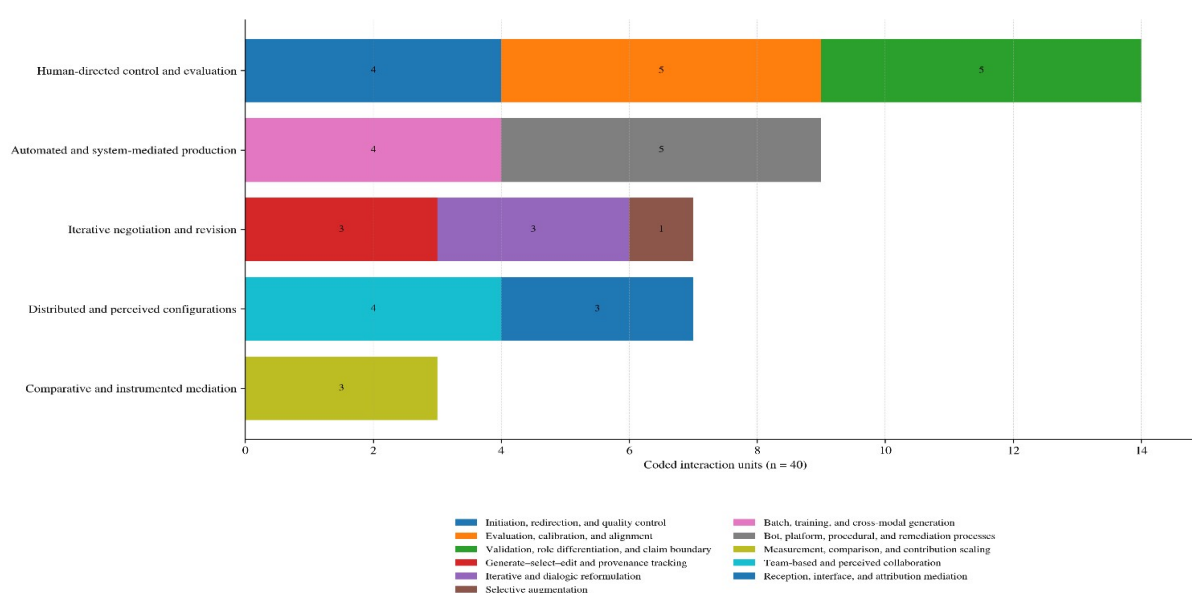


Figure 4. Distribution of human-ai negotiation configurations

Human-directed control and evaluation accounted for 14 units (35.00%), including initiation, redirection, quality assessment, validation, and differentiation of authorship roles. Automated and system-mediated production contributed 9 units (22.50%), particularly through batch generation, fine-tuning, cross-modal transformation, bots, and procedural recombination. Iterative negotiation and distributed configurations each comprised 7 units (17.50%). The former included sequences in which writers “accepted them or dismissed them, and edited accepted text,” whereas the latter covered team prompting, perceived collaboration, interface interaction, and authorship attribution. Comparative and instrumented mediation represented 3 units (7.50%), recording model settings, alternative outputs, and assistance granularity. These distributions show that direct editing constitutes only one form of intervention; evaluation, parameter selection, platform configuration, and refusal also organize collaborative production.

Agency within these sequences is better understood as asymmetrically distributed than equally shared. Lee, Liang, and Yang (2022) demonstrate that requesting, accepting, dismissing, and editing suggestions create distinct moments of human control, while Dhillon et al. (2024) show that assistance granularity changes perceived ownership and contribution. Dialogic reformulation nevertheless complicates a tool-only account because model responses can redirect subsequent human choices, producing the recursive interpretation described by Henrickson and Meroño-Peñuela (2023). Creative writers further regulate this exchange through authenticity and craftsmanship rather than indiscriminate adoption (Guo et al., 2024). Automated and platform-based cases extend agency beyond dyadic interaction, but they do not erase responsibility. Human-AI collaboration therefore emerges as a traceable negotiation of permissions, selections, revisions, infrastructures, and attribution rather than a symmetrical fusion of creative subjects.

Table 3. Distribution of human–ai negotiation configurations

Main Category (Category Total)	Subcategory	Observable Interaction Indicators	Frequency	Percentage	Human Agency	Predominant AI Role	Revision or Uptake Pattern	Textual or Documentary Evidence	Data Codes
Human-directed control and evaluation (14; 35.00%)	Initiation, redirection, and quality control	Initiating instructions; changing constraints; identifying stylistic limitations; repairing unsuitable output	4	10.00%	High	Conditional generator or assistant	Instruction, redirection, rejection, or repair	"Write a sentence about 'love'"; professional writers reported difficulty preserving authorial voice	COD-001; COD-002; COD-005; COD-019
	Evaluation, calibration, and alignment	Assessing form, creativity, responsiveness, style, user needs, and appropriate contribution level	5	12.50%	High or researcher-led	Object of evaluation, selective collaborator, or style imitator	Analytical rejection, comparison, calibration, or alignment	"Creativity, prompt responsiveness, and literary value"; requested limerick realized through repeated quatrains	COD-006; COD-012; COD-021; COD-023; COD-035
	Validation, role differentiation, and claim boundary	Verifying evidence; distinguishing prompting from generation; separating contribution, attribution, and interpretation	5	12.50%	High or researcher-led	Generator, analytic aid, or coding assistant	Human validation and interpretive boundary setting	"Prompted by me; generated by ChatGPT"; source-to-code validation remained human-controlled	COD-028; COD-029; COD-031; COD-034; COD-037
Iterative negotiation and revision (7; 17.50%)	Generate–select–edit and provenance tracking	Requesting alternatives; accepting, dismissing, or editing suggestions; recording textual origin	3	7.50%	High	Alternative or suggestion provider	Mixed acceptance, rejection, and human editing	Writers requested suggestions, accepted or dismissed them, and edited retained text	COD-003; COD-008; COD-009
	Iterative and dialogic reformulation	Comparing prompt strategies; redirecting responses; interpreting and reformulating meaning across turns	3	7.50%	High	Translator, dialogic collaborator, or interpretive interlocutor	Recursive and comparative revision	Seven distinct poetry-translation prompts; poetic expression developed through human–AI dialogue	COD-011; COD-018; COD-033

	Selective augmentation	Using AI for brainstorming, imagery, story details, research, or world-building while retaining editorial control	1	2.50%	High	Ideation assistant	Selective appropriation and human integration	"Brainstorming, story details, world-building, and research assistance"	COD-020
Automated and system-mediated production (9; 22.50%)	Batch, training, and cross-modal generation	Producing parallel outputs; executing reusable prompt templates; fine-tuning; translating poems into images	4	10.00%	Medium or researcher-directed	Corpus, API, visual, or fine-tuned generator	Unrevised, generated, or parallel output	Prompt templates were assembled from topic, form, and constraint variables	COD-004; COD-007; COD-013; COD-014
	Bot, platform, procedural, and remediation processes	Automated recombination; bot circulation; platform mediation; code-based textual production; cross-media transformation	5	12.50%	Distributed or low direct intervention	Bot, platform actor, computational actor, or procedural generator	Unrevised generation, circulation, or remediation	Twitter/X was positioned as a co-author; repetition and fragmentation marked bot-generated poetry	COD-015; COD-016; COD-036; COD-038; COD-040
Comparative and instrumented mediation (3; 7.50%)	Measurement, comparison, and contribution scaling	Recording model parameters; comparing variants; varying sentence- and paragraph-level assistance	3	7.50%	Researcher-led or medium	Instrumented partner, transformation agent, or scaffold provider	Observed, parallel, and differently scaled contributions	Metadata recorded prompt codes, temperature, queries, selections, and assistance granularity	COD-010; COD-017; COD-022
Distributed and perceived configurations (7; 17.50%)	Team-based and perceived collaboration	Team negotiation; perceived co-creation; writing scenarios; user interpretation of AI roles	4	10.00%	Distributed or variable	Shared resource, collaborative system, or chat collaborator	Mixed or perceived collaboration	Teams distributed planning, testing, and refinement around generative prompting	COD-024; COD-025; COD-026; COD-027
	Reception, interface, and attribution mediation	Reader judgment; screen interaction; authorship labels; provenance and attribution exposure	3	7.50%	Reader- or interface-dependent	Multimodal system or attributed generator	Reception rather than direct textual revision	Authorship labels could alter poetic perception even when production provenance remained unchanged	COD-032; COD-039; COD-041

From model regularity to algorithmic literary voice

Algorithmic literary voice appeared as a layered configuration of textual regularity, human intervention, technological mediation, and attribution. Forty-one source-linked units were examined through the Algorithmic Literary Voice Profile, maintaining the same coding chain used in the preceding analyses. Hybrid and human-curated configurations formed the largest category, followed by platform, procedural, and multimodal mediation. Model-conditioned regularity, reception and attribution, and methodological boundaries supplied additional dimensions. This structure prevents voice from being reduced to recurrent vocabulary or model personality. Table 4 connects stylistic indicators with authorship signals, platform traces, and variation status, while Figure 5 displays the hierarchical composition of all twelve mutually exclusive configurations. Consequently, literary voice becomes observable through relationships among textual features, production histories, editorial actions, model conditions, interfaces, and circulation environments.

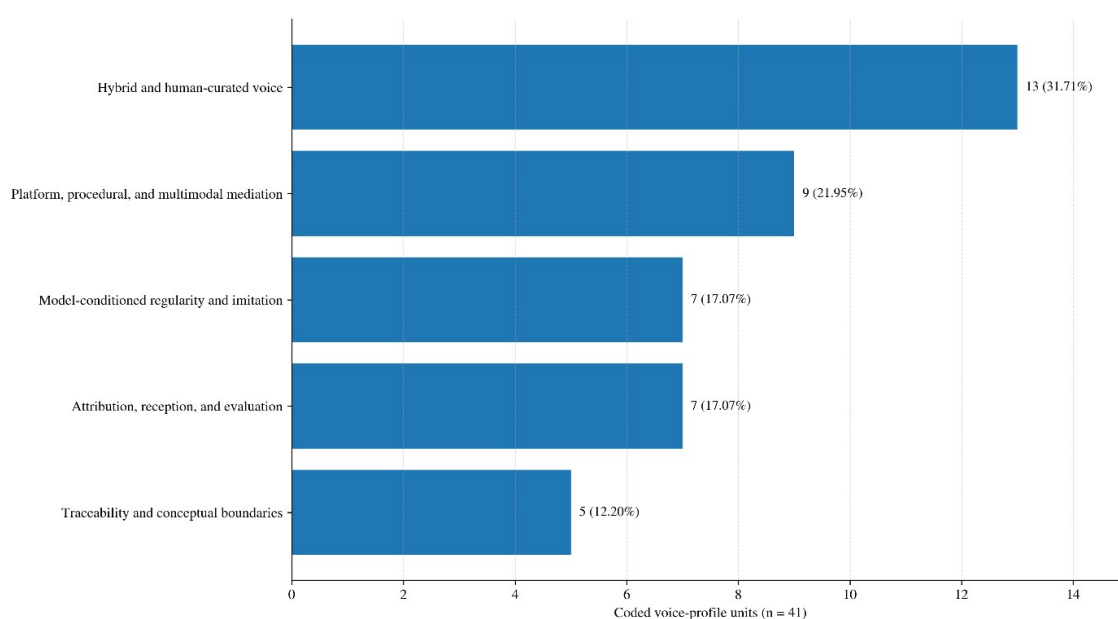


Figure 5. Composition of algorithmic literary voice profiles

Hybrid and human-curated voice accounted for 13 units (31.71%), including prompt conditioning, visible editing, authenticity management, and collective revision. Platform, procedural, and multimodal mediation contributed 9 units (21.95%); examples included colour-coded provenance, rule-based recombination, cross-modal prompts, and the claim that “Twitter/X is positioned as a co-author in literary production.” Model-conditioned regularity and imitation comprised 7 units (17.07%). Its clearest formal marker occurred when “the requested limerick is realized as repeated quatrains,” indicating that recurrent model structures could override explicit genre instructions. Attribution, reception, and evaluation also comprised 7 units (17.07%), whereas traceability and conceptual boundaries contributed 5 units (12.20%). Perceived voice, actual provenance, and stylistic recurrence therefore remained empirically distinct dimensions.

These distributions support a relational account of algorithmic literary voice rather than an autonomous model-centred definition. Mikros (2025) demonstrates how lexical, syntactic, and structural recurrence can identify model-conditioned style, yet such regularity explains only part of the corpus. Human editing and authenticity management confirm that voice is repeatedly curated, while Gunkel’s (2025) distinction between prompting and generation clarifies why initiation does not exhaust authorship. Nadifah (2025) further shows that language, code, and platform jointly mediate electronic textuality, a position reinforced by bot circulation and visual interface traces. Porter and Machery (2024) complicate matters by showing that perceived humanness may diverge from production provenance. Algorithmic literary voice is therefore best understood as a traceable composite of model tendencies, human judgment, executable procedures, platform affordances, and reader attribution.

Table 4. Distribution of algorithmic literary voice profiles

Main Category (Category Total)	Voice Configuration	Operational Indicators	Frequency	Percentage	Uniformity or Variation	Authorship and Platform Trace	Textual or Documentary Evidence	Data Codes
Hybrid and human-curated voice (13; 31.71%)	Prompt-conditioned and edited hybridity	Topic conditioning; selection among alternatives; translation strategy; editing; dialogue; mixed contribution	7	17.07%	Predominantly variable	Human initiation, selection, revision, or integration remains visible	"Human prompt, multiple model options, acceptance, rejection, and editing"; writers "accepted them or dismissed them, and edited accepted text"	COD-001; COD-003; COD-008; COD-011; COD-018; COD-020; COD-026
	Human control, authenticity, and ownership	Voice preservation; craftsmanship; contribution scale; stage-specific intervention	4	9.76%	Variable	Human authorial priority or final control	"NLG technologies struggle to preserve style and authorial voice"; writers made deliberate decisions around authenticity and craftsmanship	COD-019; COD-021; COD-022; COD-023
	Collective and emergent revision	Multi-actor revision; interturn development; collective negotiation	2	4.88%	Variable	Distributed among human participants and AI interlocutors	Teams developed collaboration strategies around generative prompting; meaning emerged through iterative reformulation	COD-024; COD-033
Platform, procedural, and multimodal mediation (9; 21.95%)	Platform, code, and interface traces	Parameter regularity; colour-coded provenance; repetition; fragmentation; executable constraints	4	9.76%	Comparable, variable, or uniform	Script, interface, bot, and platform traces are materially present	"Human text, model suggestions, and prompt text are visually differentiated"; "algorithmic repetition and structural fragmentation"	COD-007; COD-009; COD-015; COD-038
	Multimodal, remediated, and procedural voice	Emotion-image alignment; visual-verbal convergence; screen layout; rule-based recombination	4	9.76%	Variable or procedurally recurrent	Voice crosses textual, visual, code-based, and screen-mediated forms	"Prompts encode emotions, visual elements, and themes"; visual-verbal design and screen layout shape interpretation	COD-013; COD-036; COD-039; COD-040
	Distributed platform authorship	Circulation-dependent meaning; infrastructural participation; platform visibility	1	2.44%	Variable	Platform explicitly functions as a co-producing actor	"Twitter/X is positioned as a co-author in literary production"	COD-016
Model-	Constraint-	Terminal-word regularity;	3	7.32%	Comparable or	Human constraints	End-word instructions, topic-form	COD-002;

conditioned regularity and imitation (7; 17.07%)	conditioned model signatures	form adherence; lexical avoidance			variable	shape model realization	templates, and lexical prohibitions produced measurable formal conditions	COD-004; COD-005
	Default, imitative, comparative, and stylistic regularity	Recurrent stanza form; rhyme; meter; author-correlated markers; rewriting differences; lexical and syntactic recurrence	4	9.76%	Uniform or comparable	Model version, training data, and stylistic conditioning remain identifiable	"The requested limerick is realized as repeated quatrains"; stylometry tests recurrent linguistic patterns	COD-006; COD-014; COD-017; COD-035
Attribution, reception, and evaluation (7; 17.07%)	Evaluated and perceived voice	Creativity assessment; self-efficacy; perceived AI contribution; human/AI recognition	4	9.76%	Variable	Voice is assigned through human evaluation or perception	"Authorship attribution does not reliably reveal production origin to readers"	COD-012; COD-025; COD-027; COD-032
	Attribution, legal distance, and authorship labels	Separation of prompting and generation; transformative distance; attributed authorship	3	7.32%	Attribution-dependent	Actual provenance and attributed origin may diverge	"Prompted by me; generated by ChatGPT"; authorship labels can influence poetic perception	COD-031; COD-034; COD-041
Traceability and conceptual boundaries (5; 12.20%)	Parameter, workflow, and methodological traceability	Temperature; query and selection patterns; source-to-code linkage; combined feature analysis	3	7.32%	Comparable or not applicable	Metadata and analytic workflow make production conditions auditable	Session metadata includes prompt code, temperature, query count, and selection count	COD-010; COD-029; COD-037
	Non-optimization and complementarity boundaries	Refusal of automatic superiority claims; ambiguity; excess; noncompliance	2	4.88%	Variable or not applicable	Literary value is not assigned solely to model performance	"Poetry resists reduction to optimization"	COD-028; COD-030

DISCUSSION

Prompting matters because it converts literary intention into an explicit, inspectable layer of composition. The prominence of dialogic, evaluative, and formal functions indicates that prompts do more than initiate generation; they organize choices about subject, form, imagery, revision, and acceptable deviation. This supports Chakrabarty, Padmakumar, and He's (2022) account of attribute-controlled poetry, yet extends it by showing that poetic value is negotiated through constraints, redirections, and evaluative criteria rather than delivered by specification alone. Henrickson and Meroño-Peñuela (2023) similarly frame prompting as interpretive negotiation, while Elam (2023) warns against equating literary achievement with optimization. Together, these positions explain why prompt precision can function productively when it sharpens compositional attention, but dysfunctionally when it reduces poetry to compliant output. Prompting therefore becomes poetically significant not because it guarantees quality, but because it externalizes aesthetic judgment and makes visible where human intention, machine regularity, and literary uncertainty intersect within composition in practice.

This pattern arises because generative systems respond probabilistically to linguistic framing while literary writers evaluate outputs through aesthetic norms. Formal prompts constrain lexical and structural possibilities, yet recurrent model tendencies may override requested genres, as illustrated by repeated quatrain formation under a limerick instruction. Such tension reflects the difference between statistical continuation and culturally embedded poetic competence. Mikros (2025) shows that model outputs can exhibit recurring stylistic signatures, whereas Ranade, Saravia, and Johri (2024) emphasize that rhetorical prompt design shapes usefulness without eliminating model indeterminacy. The correlation between greater prompt specificity and more inspectable output does not imply causation between specificity and literary quality. Instead, specificity narrows the space of responses, making compliance easier to assess while potentially limiting ambiguity and surprise. Prompting works best as a negotiated structure: human writers specify enough to orient generation, but retain openness for revision, rejection, and reinterpretation when defaults conflict with poetic aims.

The interactional findings matter because they relocate authorship from final ownership claims to observable sequences of decision-making. Human control appeared not only through writing or editing, but also through initiating, evaluating, refusing, calibrating, and validating model contributions. This supports Lee, Liang, and Yang's (2022) process-based account of co-writing and Dhillon et al.'s (2024) finding that scaffolding scale shapes perceived contribution and ownership. However, it complicates celebratory co-creation models by showing that collaboration is frequently asymmetric: models propose, transform, or generate, while humans retain responsibility for selection and legitimacy. McGuire, De Cremer, and Van De Cruys (2024) emphasize co-creation and self-efficacy, yet this corpus suggests that perceived partnership must be distinguished from documented intervention. The implication is methodological and ethical. Authorship should be assessed through traceable actions, revision depth, and attribution practices, rather than through broad labels such as collaborator or tool, which conceal unequal control and responsibility in practice today.

Such asymmetry emerges from underlying differences in agency, accountability, and system design. Language models generate text without intentions, legal responsibility, or stable aesthetic commitments, whereas writers must evaluate coherence, originality, voice, and consequences across work. Interfaces reinforce this difference by presenting suggestions, alternatives, and controls that position human users as selectors, even when model outputs redirect later decisions. Ippolito et al. (2022) found that professional writers valued AI for brainstorming while struggling to preserve voice, and Guo et al. (2024) showed that authenticity and craftsmanship governed selective adoption. These findings support the observed link between stronger human intervention and clearer authorial control, but they do not establish that more editing always produces better poetry. Rather, revision functions as a mechanism for aligning probabilistic output with situated literary judgment. Collaboration is therefore structured by unequal capacities: models expand possibilities quickly, while humans supply continuity, normative evaluation, and responsibility for textual commitment.

The voice-profile findings matter because they challenge the assumption that algorithmic literary voice is simply a model's recognizable style. Hybrid and human-curated configurations were more prominent than purely model-conditioned regularity, indicating that voice emerges through interaction among prompts,

outputs, edits, interfaces, and circulation. This aligns with Nadifah's (2025) view that language, code, and platform jointly mediate electronic textuality and with Abdillah's (2026) account of repetition and fragmentation as platform-conditioned aesthetics. It also qualifies Mikros's (2025) stylometric approach: recurrent lexical or structural patterns remain analytically valuable, but they capture only one layer of voice. Porter and Machery (2024) further demonstrate that perceived humanness may diverge from production provenance. Consequently, algorithmic literary voice should be treated as a composite trace rather than an autonomous subject. Its significance lies in revealing how style, attribution, technical infrastructure, and editorial intervention converge to produce recognizable yet unstable literary effects across digital environments in contemporary practice.

This composite structure arises because literary voice is produced simultaneously at textual, interactional, and infrastructural levels. Model regularities contribute repeated diction, stanza forms, and syntactic preferences; prompts condition those tendencies; human revisions interrupt or redirect them; and platforms mark provenance, visibility, and circulation. Gunkel's (2025) distinction between human prompting and machine generation clarifies why contribution cannot be collapsed into singular authorship, while Mazzi (2024) separates originality in prompts from originality in outputs. Reception studies by Stańko-Kaczmarek, Dera, and Koscielska (2024) and Porter and Machery (2024) show that attribution labels further shape perceived voice independently of actual production. These layered relations collectively explain why algorithmic voice appears stable in some controlled comparisons yet variable in edited or multimodal contexts. This study therefore supports a relational model: voice is not located solely in human intention or machine output, but in the patterned transformations connecting instruction, generation, revision, platform mediation, and reader interpretation.

CONCLUSION

This study concludes that prompting functions as a poetic practice when it translates aesthetic intention into formal, rhetorical, dialogic, and multimodal operations. Its principal contribution lies in repositioning human–AI collaboration as a traceable sequence of instruction, generation, selection, refusal, revision, and platform mediation rather than as undifferentiated co-creation. By integrating a Prompt–Rhetoric Coding Matrix, Human–AI Negotiation Sequence Analysis, and an Algorithmic Literary Voice Profile, this study offers a layered method for distinguishing prompt design, observable agency, stylistic recurrence, attribution, and technological infrastructure. Algorithmic literary voice consequently emerges as a relational formation shaped by human judgment, model tendencies, and mediated circulation.

This study remains limited by its reliance on publicly available documents, repositories, interfaces, and reported interaction records, which restrict access to drafting histories, private prompts, unsuccessful iterations, and unpublished revisions. Its heterogeneous corpus also compares platforms, models, genres, and research settings that cannot be treated as equivalent. Future research should therefore construct longitudinal author-based datasets containing complete prompt–response–revision chains, model-version metadata, keystroke records, and reflective interviews. Comparative studies across languages, poetic traditions, and interface designs would clarify how cultural conventions shape prompt practice. Experimental work could test relationships among prompt structure, revision intensity, perceived ownership, stylistic distinctiveness, and reader evaluation.

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

Muhammad Khairul Umam: conceptualization, methodology, formal analysis, writing – original draft, writing – review and editing; data curation, formal analysis, writing – original draft, writing – review and editing.

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 publicly accessible documents, publisher paratexts, and platform policy materials.

ETHICAL APPROVAL

This research did not involve human or animal subjects; it analyzed publicly available literary texts, creator accounts, publisher paratexts, and platform policy documents, 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>
- Assyabani, R. (2025). Multimodal poetics in digital literature: A corpus-based analysis of visual-verbal design, screen-based textuality, and reader interaction. *Lingua Technica: Journal of Digital Literary Studies*, 1(1), 41–53. <https://doi.org/10.64595/cy0q2n14>
- Atikurrahman, M. (2025). Reimagining textuality: Digital convergence and literary adaptation in Indonesia. *Lingua Technica: Journal of Digital Literary Studies*, 1(2), 63–71. <https://doi.org/10.64595/lingtech.v1i1.30>
- Bozkurt, A. (2024). Tell me your prompts and I will make them true: The alchemy of prompt engineering and generative AI. *Open Praxis*. <https://doi.org/10.55982/openpraxis.16.2.661>
- Chakrabarty, T., Padmakumar, V., & He, H. (2022). Help me write a poem: Instruction tuning as a vehicle for collaborative poetry writing. *arXiv*. <https://doi.org/10.48550/arxiv.2210.13669>
- Dang, H., Mecke, L., Lehmann, F., Goller, S., & Buschek, D. (2022). How to prompt? Opportunities and challenges of zero- and few-shot learning for human–AI interaction in creative applications of generative models. *arXiv*. <https://doi.org/10.48550/arxiv.2209.01390>
- Dhillon, P. S., Molaei, S., Li, J., Golub, M., Zheng, S., & Robert, L. P. (2024). Shaping human–AI collaboration: Varied scaffolding levels in co-writing with language models. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*. Association for Computing Machinery. <https://doi.org/10.1145/3613904.3642134>
- Elam, M. (2023). Poetry will not optimize, or what is literature to AI? *American Literature*. <https://doi.org/10.1215/00029831-10575077>
- Gunkel, D. J. (2025). Prompted by me. Generated by ChatGPT. *Human-Machine Communication*. <https://doi.org/10.30658/hmc.10.2>
- Guo, A., Sathyanarayanan, S., Wang, L., Heer, J., & Zhang, A. X. (2024). From pen to prompt: How creative writers integrate AI into their writing practice. In *Proceedings of the 2025 Conference on Creativity and Cognition*. Association for Computing Machinery. <https://doi.org/10.1145/3698061.3726910>
- Han, Y., Qiu, Z., Cheng, J., & Lc, R. (2024). When teams embrace AI: Human collaboration strategies in generative prompting in a creative design task. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*. Association for Computing Machinery. <https://doi.org/10.1145/3613904.3642133>
- Henrickson, L., & Meroño-Peñuela, A. (2023). Prompting meaning: A hermeneutic approach to optimising prompt engineering with ChatGPT. *AI & Society*, 40, 903–918. <https://doi.org/10.1007/s00146-023-01752-8>

- Huang, Y., Shea, J., Howe, D., & Holopainen, J. (2025). Lyric poetry in the face of posthumanism: An analysis of generative AI-assisted poetry writing. In *Proceedings of the 2025 Conference on Creativity and Cognition*. Association for Computing Machinery. <https://doi.org/10.1145/3698061.3726919>
- Ippolito, D., Yuan, A., Coenen, A., & Burnam, S. (2022). Creative writing with an AI-powered writing assistant: Perspectives from professional writers. *arXiv*. <https://doi.org/10.48550/arxiv.2211.05030>
- Jamil, S., Reddy, B. A., Kumar, R., Saha, S., Joseph, K., & Goswami, K. (2025). Poetry in pixels: Prompt tuning for poem image generation via diffusion models. *arXiv*. <https://doi.org/10.48550/arxiv.2501.05839>
- Lee, M., Liang, P., & Yang, Q. (2022). CoAuthor: Designing a human–AI collaborative writing dataset for exploring language model capabilities. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*. Association for Computing Machinery. <https://doi.org/10.1145/3491102.3502030>
- Mazzi, F. (2024). Authorship in artificial intelligence-generated works: Exploring originality in text prompts and artificial intelligence outputs through philosophical foundations of copyright and collage protection. *The Journal of World Intellectual Property*. <https://doi.org/10.1111/jwip.12310>
- McGuire, J., De Cremer, D., & Van De Cruys, T. (2024). Establishing the importance of co-creation and self-efficacy in creative collaboration with artificial intelligence. *Scientific Reports*, 14. <https://doi.org/10.1038/s41598-024-69423-2>
- 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/lc/fqaf035>
- Nadifah, S. A. (2025). Language, code, and platform-mediated textuality in electronic literature. *Lingua Technica: Journal of Digital Literary Studies*, 1(1), 28–40. <https://doi.org/10.64595/80xngp06>
- Park, J., & Choo, S. (2024). Generative AI prompt engineering for educators: Practical strategies. *Journal of Special Education Technology*, 40, 411–417. <https://doi.org/10.1177/01626434241298954>
- Porter, B., & Machery, E. (2024). AI-generated poetry is indistinguishable from human-written poetry and is rated more favorably. *Scientific Reports*, 14. <https://doi.org/10.1038/s41598-024-76900-1>
- Prabowo, B. A., & Asmarani, R. (2025). Generative literature: The role of artificial intelligence in the creative writing process. *Allure Journal*. <https://doi.org/10.26877/allure.v5i1.19959>
- Ranade, N., Saravia, M., & Johri, A. (2024). Using rhetorical strategies to design prompts: A human-in-the-loop approach to make AI useful. *AI & Society*, 40, 711–732. <https://doi.org/10.1007/s00146-024-01905-3>
- Reza, M., Thomas-Mitchell, J., Dushniku, P., Laundry, N., Williams, J., & Kuzminykh, A. (2025). Co-writing with AI, on human terms: Aligning research with user demands across the writing process. *Proceedings of the ACM on Human-Computer Interaction*, 9, 1–37. <https://doi.org/10.1145/3757566>
- Stańko-Kaczmarek, M., Dera, L., & Koscielska, H. (2024). “Between the lines”: Perceptions of poetry with authorship attributed to artificial intelligence or humans—A comparative analysis. *The Journal of Creative Behavior*. <https://doi.org/10.1002/jocb.1513>