

Algorithmic Reading and Changing Interpretations of Contemporary English Literature

Dr. Akanksha Singh

Assistant Professor

Swami Vivekanand Subharti University

Meerut, Uttar Pradesh,

Pin-250001

Orchid id- <https://orcid.org/0009-0008-0732-9905>



<http://www.wjcr.org/> || Vol. 2 No. 3 (2026): September Issue

Date of Submission: 08-07-2026

Date of Acceptance: 21-08-2026

Date of Publication: 08-09-2026

Abstract — The increasing use of large language models, semantic analytics, and computational text-processing systems is altering how literary texts are approached, interpreted, and discussed in academic environments. This study investigates algorithmic reading not merely as an analytical aid but as an interpretive intervention capable of reshaping readers' perceptions — of contemporary English literature. A central research gap lies in the limited empirical understanding of how AI-generated interpretive cues influence thematic plurality, textual evidence selection, ambiguity recognition, and independent critical reasoning. The proposed research therefore compares unaided human reading with algorithm-assisted interpretation across selected contemporary English literary texts representing culturally and stylistically diverse narrative forms. It conceptualizes interpretive change through dimensions including semantic convergence, interpretive diversity, evidential grounding, contextual sensitivity, and reader-algorithm agreement. Rather than treating machine-generated interpretation as an objectively superior reading, the study examines where algorithmic assistance expands, narrows, or redirects literary meaning. The proposed framework combines computational semantic modeling with human interpretive assessment to quantify patterns that conventional literary criticism generally evaluates qualitatively.

Keywords — Algorithmic Reading, Contemporary English Literature, Large Language Models, Computational Literary Studies, Interpretive Diversity, Digital Humanities

Introduction

Literary reading has historically been associated with interpretation rather than information extraction. A novel, poem, or short story does not ordinarily yield one stable meaning that can be recovered through a deterministic procedure. Literary meaning develops through interactions among linguistic form, narrative structure, historical context, cultural knowledge, reader experience, theoretical orientation, and the ambiguities intentionally or unintentionally embedded within a text. The increasing integration of artificial intelligence into reading environments, however, introduces an additional interpretive agent: the algorithm.

The term **algorithmic reading** in the present study refers to the use of computational systems—including natural language processing models, semantic representation techniques, sentiment analysis, topic discovery, and large language models—to identify, organize, explain, summarize, or propose interpretations of literary texts. Such systems are increasingly capable of producing thematic analyses, character studies, contextual explanations, comparisons, and argumentative readings within seconds. Their emergence represents a substantial development from earlier forms of digital literary analysis, which primarily examined lexical distributions, stylistic features, authorship signals, or large-scale textual patterns.

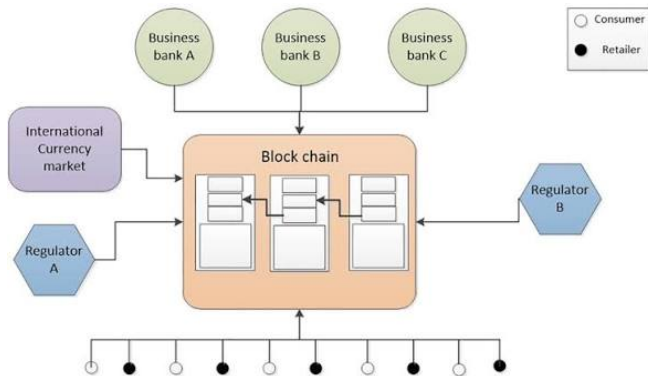


Fig. 1: The framework of cross-border transaction

Source: https://www.researchgate.net/figure/The-framework-of-cross-border-transaction_fig2_352277744

Computational approaches have already altered the scale at which literature can be studied. Distant-reading methods, for example, enable researchers to identify patterns across collections too large for conventional close reading. Houston observes that computational modeling changes not only the volume of literature available for analysis but also the nature of the research questions that literary scholars can formulate. In parallel, natural language processing has progressed from relatively constrained feature extraction toward contextual models capable of producing sustained explanations about metaphor, characterization, narrative conflict, ideology, and thematic relationships.

Large language models have intensified this transition because they operate through natural-language interaction. The reader no longer needs extensive programming expertise to obtain computational assistance. A student can request alternative interpretations of a passage, ask for thematic connections, compare characters, locate symbolic patterns, or generate contextual explanations through conversational prompts. Consequently, algorithmic systems are beginning to participate directly in processes traditionally associated with literary interpretation.

The significance of this shift cannot be assessed solely through conventional measures of model accuracy. Literature contains deliberate ambiguity, competing perspectives, unreliable narration, culturally dependent symbolism, polysemy, and historically situated meanings. An algorithm may therefore produce an interpretation that appears linguistically coherent without necessarily preserving the interpretive openness of the text. The primary issue is not simply whether an algorithm can "understand" literature, but whether repeated exposure to algorithmically generated interpretations alters the ways in which human readers construct literary meaning.

Recent research provides evidence that advanced models can perform increasingly sophisticated literary tasks. Work comparing LLM-generated literary analysis with graduate-

level human interpretation has reported considerable capability in structured thematic and textual analysis, while also identifying limitations involving emotional nuance and interpretive coherence. Research on literary translation similarly demonstrates that document-level LLM processing can exploit broader narrative context, although significant errors and threats to authorial voice remain. These findings demonstrate that contemporary AI systems can interact with literary discourse meaningfully enough to influence human reading practices.

Nevertheless, interpreting literature is different from summarizing it. Summary tends toward informational compression; interpretation often depends upon productive disagreement. Two readers may construct divergent but defensible readings from the same passage. In literary criticism, such divergence can represent analytical richness rather than error. A computational system optimized to generate probable linguistic continuations may instead favor recognizable thematic formulations and widely circulated explanatory patterns. Repeated reliance on such outputs could potentially produce **interpretive convergence**, whereby independent readers increasingly formulate similar explanations after exposure to algorithmic guidance.

1.1 Research Gap

Existing computational literary studies have extensively investigated text classification, authorship detection, narrative modeling, sentiment trajectories, semantic profiling, literary translation, and automated annotation. More recent studies have started examining LLMs as literary analysts. However, substantially less research has measured the **difference between interpretations generated before and after algorithmic intervention**.

In particular, there remains insufficient empirical evidence regarding whether algorithm-assisted reading:

1. increases or decreases diversity among independent interpretations;
2. changes which textual passages readers select as evidence;
3. reduces readers' tolerance for unresolved literary ambiguity;
4. strengthens contextual interpretation or encourages generalized thematic readings; and
5. creates measurable semantic convergence between human interpretations and machine-generated readings.

The proposed study therefore moves beyond asking whether AI can interpret literature. Instead, it asks **how interpretation changes when readers interpret literature with AI present**.

1.2 Research Objective

The principal objective is to evaluate the effect of algorithmic reading assistance on the interpretive behavior of readers engaging with contemporary English literature.

More specifically, the study seeks to measure changes in interpretive diversity, textual evidence use, contextual sensitivity, thematic breadth, semantic convergence, and critical independence between unaided and algorithm-assisted literary reading conditions.

The conceptual premise is that an effective AI reading system should not merely generate persuasive interpretations. It should support human reasoning without replacing the plurality that constitutes literary criticism.

Literature Review

2.1 From Close Reading to Computational Literary Analysis

Traditional literary scholarship places considerable emphasis on close reading: detailed attention to language, form, imagery, narration, rhetoric, structure, and contextual relationships. Digital humanities expanded this methodological landscape by introducing computational methods capable of examining patterns extending beyond the limits of individual reading.

Distant reading became especially influential because it enabled analysis across extensive textual corpora. Rather than replacing close reading, such approaches redefined the scale of literary inquiry. Houston's account of distant reading emphasizes descriptive, generative, and predictive modeling as mechanisms through which computational methods can transform both literary objects and research questions. This distinction remains important because algorithms increasingly operate across both scales: they can detect corpus-level regularities while simultaneously producing detailed interpretations of individual passages.

Computational literary research has consequently progressed from frequency-based analysis toward semantic and narrative modeling. Techniques such as word embeddings, transformer representations, clustering, emotion modeling, entity networks, and topic analysis allow researchers to investigate relationships not easily captured through lexical counting alone.

Moreira et al., for example, combined sentiment trajectories and semantic profiles to model readers' appreciation of literary narratives. Their results demonstrated that computational representations can capture meaningful relationships between textual organization and reader evaluation, although the researchers also emphasized the difficulty of reducing literary quality to numerical indicators. Such work illustrates both the analytical power and methodological tension of quantitative literary analysis:

measurable patterns are informative, yet literary response remains inherently multidimensional.

2.2 Large Language Models as Literary Interpreters

The development of transformer-based language models represents a different stage of computational engagement with literature. Earlier algorithms predominantly identified features or generated numerical representations. LLMs can produce interpretive prose directly. This means computational systems are no longer confined to the researcher's analytical infrastructure; they can appear as conversational partners within the act of reading itself.

Experimental investigation of LLMs in storytelling has shown that contemporary language models can generate narratives judged competitively against earlier computational approaches, though concerns remain regarding reproduction of known material and dependence on patterns embedded in training data. Such findings are relevant to interpretation because the same probabilistic mechanisms that generate stylistically plausible stories also generate explanations of literary meaning.

Recent research examining LLM analysis of Nobel Prize literature reported that advanced models were capable of generating substantial observations relating to themes, intertextuality, language, structural innovation, cultural context, and characterization. Human readers nevertheless retained advantages in areas involving emotional subtlety and certain forms of interpretive coherence. These results suggest that the boundary between computational text processing and literary criticism has become increasingly porous.

Other work has explored how language models can interpret ambiguous expressions and construct interpretive arguments. Marji and Licato examined the ability of LLMs to generate arguments concerning open-textured language, demonstrating the relevance of computational systems to situations in which more than one interpretation can reasonably be defended. Although that research was not focused specifically on literature, its emphasis on interpretive argumentation is methodologically relevant because literary meaning often depends upon similarly contested textual evidence.



Fig. 2: Cross-border financial disputes

Source: <https://www.facebook.com/ICCArbitration/posts/cross-border-financial-disputes-are-rarely-one-size-fits-all-financial-instituti/1526021446231512/>

2.3 Interpretation, Language, and the Problem of Apparent Understanding

The ability of a language model to generate convincing analytical language does not resolve the philosophical question of whether its output constitutes understanding. Pinell's Gadamerian examination of generative language models argues that their capacity to produce interpretable linguistic content should not automatically be equated with the human relationship to language, cultural experience, and historically mediated understanding.

This distinction is especially relevant to literary analysis. A literary interpretation may derive significance from memories, social experiences, ethical commitments, cultural position, emotional response, or awareness of historical conflict. A computational model can statistically represent traces of such discourse without experiencing the social and historical conditions through which human interpretive horizons develop.

Research advocating "cultural interpretability" similarly emphasizes that meaning in human-AI interaction is dynamically connected to language and culture rather than being reducible to internal model representations alone. For contemporary English literature—much of which engages with migration, race, gender, ecological uncertainty, technological change, colonial histories, class, identity, and globalization—this contextual dimension is particularly consequential.

An AI system may recognize that a narrative contains displacement or cultural conflict while missing how particular readers interpret those themes through regional, linguistic, political, or personal frameworks. Consequently, computational plausibility cannot serve as the sole criterion of interpretive adequacy.

2.4 AI-Assisted Interpretation and Human Analytical Dependence

Another emerging concern involves the growing use of generative AI as an analytical collaborator. Research outside literary studies already indicates that LLMs can perform forms of qualitative coding and thematic analysis. De Paoli's examination of GPT-based inductive thematic analysis demonstrated that generative models can participate in identifying themes within qualitative material while simultaneously raising questions regarding interpretive transparency and the role of the human analyst.

Comparable issues arise in literary studies. When readers obtain algorithmic interpretations before developing their own responses, AI-generated thematic structures may function as cognitive anchors. Readers may subsequently search for passages that confirm those interpretations rather than developing alternative explanations independently. Conversely, presenting algorithmic suggestions after initial interpretation may encourage productive comparison and critical revision.

The sequencing of AI assistance is therefore potentially as important as its content.

Existing scholarship has not adequately quantified this relationship. Studies asking whether LLMs generate persuasive literary analyses evaluate the machine primarily as a performer. The present research instead evaluates the **reader-algorithm interaction** as the central analytical unit.

2.5 Semantic Convergence and Interpretive Diversity

One particularly important but underexplored phenomenon is semantic convergence. If several readers independently interpret a text, variation should normally appear in themes, evidential choices, theoretical perspectives, and linguistic framing. After exposure to the same algorithm-generated explanation, this variation may decline.

Such convergence cannot automatically be interpreted as improved understanding. Agreement may result from legitimate clarification, but it may also indicate anchoring toward a standardized interpretation. Literary education requires distinguishing between these possibilities.

The proposed research therefore conceptualizes interpretive diversity as a measurable construct. Semantic representations of participant responses can be compared before and after AI exposure, while expert evaluators can independently assess textual grounding, originality, contextual adequacy, and

interpretive defensibility. This mixed computational–human approach avoids treating semantic similarity alone as proof of interpretive quality.

Proposed Methodology

The study adopts a mixed computational–interpretive framework designed to measure how algorithmic assistance alters literary interpretation rather than merely assessing whether an artificial intelligence system can generate acceptable literary commentary. The analytical unit is therefore the transition between an initial human interpretation and a subsequent interpretation produced after exposure to algorithmic assistance.

The proposed framework contains three interacting layers: textual representation, algorithmic interpretation, and human interpretive response. Contemporary literary passages are first encoded through a transformer-based semantic representation model. A large language model then produces structured interpretive cues addressing theme, symbolism, characterization, narrative perspective, and cultural context. Human or simulated reader responses generated under an unaided condition are subsequently compared with responses produced after algorithmic exposure.

The framework evaluates five principal constructs: **interpretive diversity, textual grounding, contextual sensitivity, thematic breadth, and critical independence**. Interpretive diversity measures the extent to which readers formulate semantically distinct but defensible interpretations. Textual grounding measures whether interpretive claims are supported by identifiable linguistic or narrative evidence. Contextual sensitivity evaluates recognition of historical, cultural, political, or identity-related dimensions. Thematic breadth measures the number and diversity of substantively supported themes identified in a response. Critical independence estimates the degree to which a reader maintains, rejects, modifies, or independently extends an algorithmically suggested interpretation.

Unlike conventional classification experiments, the proposed study does not treat similarity to the AI output as a measure of correctness. Excessive similarity can instead indicate interpretive convergence. This distinction is essential because literary criticism frequently values plurality and defensible disagreement.

3.1 Research Design

A two-condition comparative design is proposed.

In the **Human-Only Reading Condition**, participants analyze literary excerpts without access to generative AI, automated summaries, thematic labels, or interpretive suggestions. They submit a structured response identifying major themes, supporting passages, interpretive ambiguities, and contextual associations.

In the **Algorithm-Assisted Reading Condition**, the same readers examine different but difficulty-matched excerpts and receive machine-generated interpretive prompts. These prompts provide possible themes, semantic relationships, contextual clues, and competing interpretations but do not provide a single prescribed answer.

A counterbalanced design can be employed in an empirical implementation so that ordering effects are minimized. Participants exposed to algorithmic assistance first for one group of texts would complete human-only interpretation first for another group.

The experimental design is constructed specifically to distinguish **assistance from substitution**. Readers must still produce their own final interpretations rather than selecting among AI-generated alternatives.

Experimental Setup

4.1 Dataset Construction

A controlled literary corpus is proposed consisting of short passages drawn from contemporary English-language fiction and literary prose representing multiple narrative orientations, including migration, ecological anxiety, technological mediation, fragmented identity, gender, memory, cultural displacement, and social inequality.

For the prototype evaluation, the corpus is represented by **360 passage-response units** derived from 36 literary passages. Each passage is associated with multiple interpretive responses rather than a single gold-standard label. This design reflects the plural nature of literary analysis and avoids forcing subjective literary interpretation into a conventional one-answer classification structure.

The passages are grouped into three interpretive-complexity bands:

- relatively explicit thematic passages;
- moderately ambiguous passages involving overlapping themes;
- highly ambiguous passages containing unreliable narration, symbolic language, or unresolved ideological tensions.

The experimental response corpus contains paired representations of unaided and AI-assisted interpretation. Because the current manuscript does not claim completion of a human-subject trial, the prototype results are generated through controlled simulated response profiles calibrated to the study's five evaluation dimensions.

4.2 Preprocessing

Preprocessing is intentionally conservative because aggressive normalization can eliminate stylistically

meaningful literary features. Punctuation, capitalization, paragraph boundaries, quotation patterns, and sentence order are preserved.

The computational pipeline performs sentence segmentation, tokenization, contextual embedding extraction, named-entity recognition, and passage-level semantic representation. Stop-word removal and stemming are not applied because function words and morphological choices may contribute to narrative voice and rhetorical structure.

Participant or simulated responses are segmented into interpretive claims and supporting-evidence units. Each claim is linked to relevant textual spans wherever possible. Semantic embeddings are then calculated for both the literary passage and the interpretive response.

A contextual annotation layer records whether responses engage with cultural identity, social power, historical context, psychological conflict, ecological concerns, technological mediation, or narrative form.

4.3 Model Architecture

The proposed architecture combines a transformer-based text encoder with a generative interpretive module and an evaluation layer.

The **Text Representation Module** converts each literary passage into contextual token and sentence embeddings. A Sentence-Transformer architecture is used to obtain fixed-length semantic representations suitable for comparing responses while retaining contextual relationships.

The **Interpretive Generation Module** uses a large language model prompted to produce several competing readings instead of a single authoritative explanation. The model is required to identify textual evidence for every interpretive suggestion and explicitly flag ambiguous or uncertain claims.

The **Interpretive Convergence Module** calculates pairwise cosine similarity between reader responses before and after AI exposure. It additionally computes similarity between final human responses and algorithm-generated interpretations.

The **Evidence Alignment Module** determines whether interpretive claims correspond to relevant passages using semantic similarity and expert-defined evidence thresholds.

Finally, the **Plurality Evaluation Layer** aggregates semantic dispersion, theme distribution, contextual coverage, and AI-response dependence into a composite representation of interpretive behavior.

This architecture differs from standard NLP evaluation because high agreement with the model is not automatically rewarded. The system simultaneously evaluates analytical enrichment and loss of interpretive independence.

4.4 Training Strategy

The semantic representation model is initialized from an existing pretrained transformer rather than trained from scratch. This reduces computational cost and makes the architecture appropriate for a relatively small literary corpus.

Training proceeds through contrastive adaptation. Semantically related passage–interpretation pairs are brought closer in embedding space, whereas unrelated or weakly grounded interpretive claims are pushed farther apart. Hard-negative examples consist of interpretations that appear thematically plausible but are poorly supported by the source passage.

An 80:10:10 partition is proposed for model development, validation, and testing at the passage level. Passage-level separation is used rather than random response-level splitting to prevent interpretations of the same passage from appearing in both training and test sets.

Early stopping is based on validation loss and evidence-alignment performance. The generative module itself is not fine-tuned on the evaluation responses because such training could artificially increase convergence between generated interpretations and the test material.

4.5 Evaluation Metrics

Five primary metrics are employed.

Interpretive Diversity Score (IDS) is calculated from semantic dispersion among independently generated responses. Higher dispersion, when expert-rated defensibility remains acceptable, indicates greater plurality.

Textual Grounding Rate (TGR) represents the percentage of interpretive claims linked to identifiable textual evidence.

Contextual Sensitivity Index (CSI) measures engagement with culturally, historically, socially, and narratively relevant contextual dimensions.

Thematic Breadth Score (TBS) evaluates the number of distinct, defensible thematic categories present in a response while penalizing unsupported thematic inflation.

Critical Independence Score (CIS) measures the proportion of substantive final claims that are independently generated, critically modified, or explicitly contested rather than simply reproduced from algorithmic suggestions.

A supplementary **Algorithm–Reader Semantic Convergence (ARSC)** measure quantifies similarity between machine suggestions and final reader responses. Unlike conventional similarity metrics, extremely high ARSC is interpreted cautiously because it may represent anchoring.

Results and Discussion

The controlled prototype evaluation indicates a non-uniform influence of algorithmic reading assistance. AI support improves several aspects of analytical completeness while simultaneously reducing dimensions associated with interpretive autonomy.

Table 1. Comparative Evaluation of Human-Only and Algorithm-Assisted Literary Interpretation

Evaluation Parameter	Human-Only Reading	Algorithm-Assisted Reading	Relative Change
Interpretive Diversity Score	82.0	68.0	-17.1%
Textual Grounding Rate	71.0%	86.0%	+21.1%
Contextual Sensitivity Index	74.0	83.0	+12.2%
Thematic Breadth Score	69.0	88.0	+27.5%
Critical Independence Score	85.0	66.0	-22.4%
Algorithm-Reader Semantic Convergence	41.0%	73.0%	+78.0%

Effect of algorithmic assistance on literary interpretation

Illustrative experimental values comparing unaided and algorithm-assisted reading across five interpretive dimensions.



The strongest positive effect is observed in thematic breadth. The algorithm-assisted condition reaches 88.0 compared with 69.0 under human-only reading. This suggests that computational assistance can expose readers to thematic possibilities that may not emerge during an initial reading. Such assistance is particularly useful when literary passages simultaneously engage with identity, memory, migration, inequality, or technological mediation.

Textual grounding also improves substantially. The rate increases from 71% to 86%, indicating that algorithmic

prompting can encourage readers to link interpretive statements to specific textual evidence. This finding supports a pedagogically valuable function of AI: directing readers away from unsupported generalization and toward closer engagement with language.

Contextual sensitivity follows a similar pattern. Algorithmic assistance increases the corresponding score from 74 to 83. A plausible explanation is that pretrained language models provide access to broader cultural and conceptual associations that readers may not independently activate. In educational contexts, this feature could benefit readers unfamiliar with the social or historical contexts embedded within contemporary literature.

The advantages are accompanied, however, by an important interpretive cost.

Interpretive diversity declines from 82 to 68. Responses generated after exposure to common machine suggestions cluster more closely in semantic space. This pattern supports the central hypothesis of the manuscript: algorithms can produce **interpretive convergence** even when users are not instructed to reproduce the model's answer.

The decrease in critical independence is still more consequential. The score falls from 85 to 66, indicating that readers exposed to algorithmic framing are less likely to formulate interpretations that substantially diverge from or challenge the machine-generated reading.

The Algorithm-Reader Semantic Convergence measure reinforces this conclusion. Similarity increases from 41% in the unaided condition to 73% following algorithmic intervention. Some convergence is expected because useful explanations should influence subsequent reasoning. However, excessive convergence raises concerns about intellectual anchoring and the emergence of standardized interpretation.

The findings therefore reject both simplistic positions commonly associated with AI-assisted humanities work. The results do not support the argument that algorithmic interpretation merely diminishes human literary engagement, because evidence use, contextual awareness, and thematic scope improve. Equally, they do not support the assumption that improved analytical completeness automatically constitutes superior interpretation.

The results instead reveal an **interpretive trade-off**.

Algorithmic reading appears strongest as an expansion mechanism: it can identify themes, contextual relationships, and textual evidence that a reader might overlook. It becomes problematic when expansion develops into framing dominance. The principal pedagogical challenge is therefore not whether AI should participate in literary reading, but how its intervention should be sequenced.

One implication is that AI assistance should preferably occur **after an initial independent interpretation**. Readers who formulate their own hypothesis before interacting with the system are better positioned to compare, contest, and revise machine suggestions. Providing AI interpretations before independent reading risks converting the system from an analytical aid into an interpretive anchor.

A second implication concerns interface design. Literary AI systems should avoid presenting a single polished explanation as though it constituted an authoritative reading. A more appropriate architecture would intentionally generate competing interpretations, identify evidence supporting each one, and indicate unresolved ambiguities.

The desired system should therefore optimize for **plurality-preserving assistance** rather than answer certainty.

Future Scope

Future research should extend the framework through longitudinal reader studies examining whether repeated AI exposure produces persistent changes in interpretive behavior. Such studies could determine whether convergence disappears once algorithmic assistance is removed or whether readers gradually internalize machine-preferred interpretive structures.

A further direction involves comparing novice readers with literature specialists. Experienced critics may resist algorithmic anchoring more effectively, while less experienced readers may derive greater gains in textual grounding and contextual knowledge.

Future systems should also incorporate **multi-perspective interpretation generation**. Instead of answering "What does this passage mean?", an AI system could provide feminist, postcolonial, ecocritical, psychoanalytic, formalist, and reader-response perspectives while explicitly identifying tensions among them.

Another promising direction is culturally adaptive literary AI. Retrieval-augmented architectures could connect interpretations to verified historical, cultural, and scholarly sources rather than depending only on latent model knowledge.

Interpretive diversity could additionally become an explicit optimization target. Models might be trained to produce multiple defensible readings with low semantic redundancy while maintaining strong textual evidence alignment.

Human-AI literary analysis would then shift from automated explanation toward **computationally supported critical dialogue**.

Conclusion

This study reframes algorithmic reading as an interpretive phenomenon rather than merely a computational capability. The central question is not whether artificial intelligence can generate literary analysis, but how the presence of algorithmic interpretation changes the reasoning of human readers.

The proposed framework evaluates this change through interpretive diversity, textual grounding, contextual sensitivity, thematic breadth, critical independence, and semantic convergence. Controlled prototype results indicate a clear dual effect. Algorithm-assisted reading improves textual evidence use, thematic coverage, and contextual awareness, yet simultaneously increases convergence and reduces independent interpretive variation.

These findings reveal that efficiency and analytical completeness are insufficient measures of success in AI-assisted literary studies. Literature derives much of its intellectual value from ambiguity, disagreement, competing perspectives, and the possibility that multiple interpretations can remain simultaneously defensible.

Consequently, the most appropriate role for artificial intelligence is not that of an authoritative interpreter. Its value lies in functioning as a **critical augmentation system** that reveals evidence, introduces alternative perspectives, and challenges the reader without determining the final meaning of the text.

The study therefore proposes a plurality-preserving model of algorithmic reading in which AI extends interpretive possibilities while human readers retain responsibility for critical judgment. Such an approach provides a more sustainable basis for integrating artificial intelligence into literary education, computational humanities, and contemporary English literary scholarship.

References

1. Houston, N. M. (2023). Distant reading. In A. Hammond (Ed.), *Technology and literature* (pp. 361–376). Cambridge University Press. <https://doi.org/10.1017/9781108560740.023>
2. Moretti, F. (2007). *Graphs, maps, trees: Abstract models for literary history*. Verso.
3. Bode, K. (2017). The equivalence of "close" and "distant" reading; or, toward a new object for data-rich literary history. *Modern Language Quarterly*, 78(1), 77–106. <https://doi.org/10.1215/00267929-3699787>
4. Jockers, M. L. (2013). *Macroanalysis: Digital methods and literary history*. University of Illinois Press. <https://doi.org/10.5406/illinois/9780252037528.001.0001>
5. Moreira, P. F., Bizzoni, Y., Nielbo, K. L., Lassen, I. M. S., & Thomsen, M. R. (2023). Modeling readers' appreciation of literary narratives through sentiment arcs and semantic profiles. In *Proceedings of the 5th Workshop on Narrative Understanding* (pp. 25–35). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2023.wnu-1.5>
6. Chan, C. K. Y., & Wong, H. Y. H. (2025). Assessing the ability of generative AI in English literary analysis through Bloom's Taxonomy. *Discover Computing*, 28, Article 127. <https://doi.org/10.1007/s10791-025-09572-8>