

Vol. 4 No. 1 (2026)

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

<https://llrjournal.com/index.php/11>

Algorithmic Narrative Structures: How Large Language Models Mimic and Disrupt Classical Structuralist Literary Syntax



Khadija Zeb

M.Phil. Scholar, Department of English, Abdul Wali Khan
University Mardan, Pakistan
Email: khadijazeb368@gmail.com

Palwasha Alam

M. Phil Scholar, Department of English, Abdul Wali Khan
University Mardan, Pakistan
Email: alam.palwasha1999@gmail.com

Amna Khattak

M. Phil Scholar, Department of English, Abdul Wali Khan
University Mardan, Pakistan
Email: amnakhattak20@gmail.com

Abstract

With the advent of Large Language Models (LLMs) and their widespread and easy adoption, computational writing has become a central one in the fields of literary theory, computational linguistics, and digital humanities. However, the use of LLM-generated fiction in a narratological context still needs to be further theorized. This work explores the possibility of the classical structuralist narrative syntax being embedded in contemporary LLMs, and its temporary dis/simultaneous disruption. The research draws on the typology of Vladimir Propp's *Morphology of the Folktale* and Joseph Campbell's monomyth, to create a mixed-methods design using 500 AI-generated stories from four model families (GPT, Claude, Gemini, and Llama). Both quantitative and qualitative analyses, mapping Proppian functions and Hero's Journey stages, respectively, were conducted and quantified, with a qualitative analysis focused on anomalies in the structure such as semantic drift, repetitive loops, character irregularity, disjointed story lines and dislocated resolutions. The results indicate that the LLMs are very good at repeating the events of surface-level narrative elements: 87.6% of narratives included an initial lack or absention, 84.2% included a call or mediation event, and 76.8% included a confrontation. Likewise, elements of Hero's Journey presented with high frequency like the ordinary world, the call to adventure, the threshold crossing and ordeal. But only 28.4% of all content was strict monomyth completion and 53.6% of narratives experienced the presence of one or more algorithmic interrupts. It is argued in this article that storytelling in LLM is not just a copy of some of the universal elements of narratives, nor an uninformed random variation on human-made forms. Rather, it results in a syntax of algorithmic narratives that is emergent, based on recognizable classical functions, but is susceptible to local (coherence) failure, recursive (re)combination, and premature closure. With this contribution, AI storytelling is declared a key space which should be extended beyond textual authoring and systems of human design.

Keywords: Large Language Models, Structuralist Syntax, Algorithmic Narratology, Literary Computing, Narrative Structure, AI Storytelling

Introduction

Background of the Study

Generative AI, which is adaptable and user-friendly, has become firmly embedded in the cultural ecosystem for drafting, editing, translating, sharing and assessing stories. Whereas they used to be tasked with giving a text a probable continuation, Large Language Models, for which enormous collections of text are prepared and which are optimized to predict text continuations, can now compose short stories, dialogue for the games, screenplay outlines, fan fiction, advertising copy, pedagogical stories, etc., with an amazing accuracy. Very often their output includes all the common elements to literature, characters who have a purpose, a build-up of conflict, helpers and antagonists, items with symbolic significance, moral shifts and resolution that attempt closure. Paradoxically, as these features have made for enthusiasm about "AI-assisted creativity", they have also familiarised an older theoretical debate regarding the "narrative format": can or should this format be regarded as an expressive

achievement of humanity, a grammar of culture, or as a pattern that can be computed? The current emergence of LLM storytellers is significant in that it moves the automatic generation of stories from a rule-based approach to story planning to a large-scale statistical language modelling approach. While these are often found explicitly in previous systems—plot grammars or planning operators—they are learned implicitly by LLMs given access to lots of text. This distinction is of significance to the field of literary study. An LLM that generates a story that is like a fairy tale or quest romance or moral fable, has no need to necessarily have thought through Proppian morphology or Campbellian transformation. Produced a series of likely linguistic and story extensions. The implied story construction can thus mimic-structuralist syntax, and yet lack the conceptual hierarchy that structuralist theories hold to be the component of a narrative system.

Structuralist Approaches to Narrative

According to the theories of narration referred to as Structuralist, narrative is not a simple sum of events but it is a comparability system between them. Narrative has been thought of as a syntax (a patterned arrangement of functions, actants, temporal relations, and transformations) as in Russian formalism, Proppian morphology, French structuralism and classical narratology. Propp (1968) suggested that Russian wonder tales could be broken down into the repetition of functions which in a fairly constant order appear in the tale. In more mythographic and comparative lines, Campbell (2008) suggested that many heroic myths have three stages: the departure, the initiation and the return. It was later expanded by the distinction narrative discourse/story, surface expression/underlining grammar, and temporal arrangement/causal logic made by Barthes (1975), Lévi-Strauss (1963), Genette (1980), Chatman (1978), and later, by narratologists.

Narrative syntax can be helpful when evaluating AI fiction because the research can not only determine if a production is fluent, but if it are arranged in the kind of structure that can be known. While an LLM narrative can have a hero, a quest and a resolution, there can be instability in the order of these aspects. The hero can be given a magical object before the problem is stated, return home before the ordeal or alternate identities may be used between paragraphs. This represents oddities that mean that one can not measure the structure of a narrative, simply by the presence of the conventional motifs. It needs to be evaluated by order, causality, recurrence and closure as well.

Research Problem

To date, although there has been extensive debate about AI-generated storytelling, much of it has focused on how creative and fluent or how flawed and hallucinatory the writing is—rather than on how it structures at the syntax level of narrative. Whereas aspects such as novelty and surprise, as well as human-supported collaboration have been explored in the field of computational creativity, much research into computational linguistics concentrates on coherence, controllability and evaluation measures. In contrast, literary theory has extensive analyses and descriptions of narrative form but lacks the development to portray how narrative forms are reproduced and distorted by probabilistic text generators. This prompts a mental distance between structuralist narratology and the present-day AI produced narration.

The issue is not whether or not LLMs can write stories (per se). They clearly can create texts that are identified as stories. The issue is more accurately if the stories are in step with the profound pattern of structural functions identified by Propp and Campbell, or if they mimic the surface indications of such structural functions. If the latter is the case, then LLM fiction could provide us with an "algorithmically unstable, structurally imitative kind of organization of narrative.

Research Objectives

This study has four objectives: One is to determine the frequency of repetition of the basic Proppian actions of narration. The second is to analyze how much the use of these narratives conforms to Campbell's Hero's Journey paradigm. The third is to find repeating puzzle-like patterns of algorithmic disruption occurring if LLMs produce lengthy sequences of narrative. The fourth is to theorize this research by using the notion of algorithmic narrative syntax, which is an expression of the probabilistic organisation of the units of a narrative created by language models.

Research Questions

The study poses three questions and uses these to direct the findings. To answer this question we need to determine if LLM-generated narratives replicate the functional sequence characteristic of the Propp morphology of the folktale. Second, what is the frequency and the fullness of the LLM generated narratives being in line with the stages of Campbell's hero's journey? Third, what kind of distinctions in algorithmic disruption(s) can be detected in LLM-generated narratives, and what impact do they have on classical structuralist approaches to narrative syntax.

Significance of the Study

In a context like the one of the digital humanities, the study is of eminent importance, as it shows that classical narratology is not merely reduced to a matter of numbers, but can be used operationally with computational analysis tools. It has more important implications for literary analysis as it places a challenge to the viability of structuralist theories in a space that is not exclusively human authored texts. It has an impact for computational linguistics as it ties narrative coherence failure to certain functional positions in the structure instead of as general mistakes. Lastly, given its relevance to AI use in creative writing, her focus on the conditions when LLMs create compelling narrative scaffolds, as well as odd structures that need human intervention, is noteworthy. Thus, the article defines AI storytelling as a site that is located in between literary syntax and statistical language modeling, and computational creativity.

Structuralist Literary Theory

The concept of structuralist literary theory was based on cultural objects being given meaning through systems of difference rather than expressions that were isolated from a system. In narrative studies this premise resulted in a failure to bring into focus authorial intention and thematization and a move to underlying structures. Later, structuralists examined myths, tales and literary works as systems of rules governed by them, whereas the Russian formalists contrasted between an "act" of literature and the more or less arbitrary arrangement of those "acts. Lévi-Strauss (1963) saw myth as a relational system of oppositions and Barthes (1975) believed that narratives could be examined by means of codes and the levels of meaning. All of these allowed the

possibility of narrative being read as a grammar in terms of recurring functions, roles, transformations and temporal relations.

Figure 1 provides a summary of the changes that occurred within morphology studies when shifting from an early formalist approach to the present-day algorithmic narration. The figure is schematic and not exhaustive, but does illuminate the intellectual genealogy which has allowed the present study. The currently latest addition to this genealogy is fiction produced by LLM, which transforms the determined textual structures to stochastic outputs.

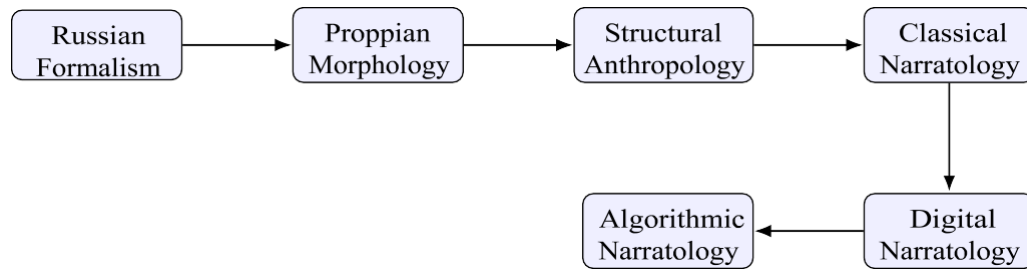


Figure 1: Evolution of Structuralist Narrative Theory

According to structuralist narratology the information content of a story is not depleted by the story words. Retells, translates, summarises or adapts a story, keeping recognisable structure. This information is important in the analysis of LLMs, as outputs may occasionally be recombinations of well-known (narrative) templates in new words. Whether these templates have coherent structural depth, or if they are made from statistical associations among motifs is an open question.

Vladimir Propp's Morphology of the Folktale

Propp's Morphology of the Folktale is one of the most influential descriptions of an (abstract) formal system of a narrative. Propp (1968) pointed out that Russian wonder tales have thirty-one functions, which always replace each other in a fixed order even if characters, settings and motifs change.

Absentation, interdiction, violation, villainy, mediation, departure, donor testing, struggle, victory, return, recognition, and wedding not only indicate actions, but they question what the book of Acts does to the narrative. One function can thus be represented by more than one character/object or event.

The main Proppian functions that are used in the present study are listed in Table 1. Since the traditional stories in modern format created by AI are not confined to the Russian wonder-tale, the thirty-one functions were reduced to twelve categories of operations. This reduction not only maintains the existence of the morphological logic of progression, but also enables computational annotation in a more reliable way in the cross comparison of genres.

Table 1: Major Proppian Narrative Functions

Function category	Narrative role	Operational marker in LLM narratives
Absentation or lack	Establishes initial deficiency, loss, or separation	A missing object, family rupture, social crisis, or unfulfilled desire

Interdiction	Introduces prohibition, warning, rule, or taboo	An elder, institution, map, prophecy, or narrator states a constraint
Violation	Activates conflict by breaking the interdiction	The protagonist ignores a warning or crosses a forbidden boundary
Villainy or lack intensification	Makes harm visible and narratively urgent	A theft, curse, invasion, illness, disappearance, or public danger
Mediation or call	Communicates the problem to the protagonist	A messenger, dream, crisis, or request demands action
Departure	Moves the protagonist from ordinary location to quest space	Leaving home, entering a forest, city, archive, sea, or digital realm
Donor test	Tests worthiness or competence	Riddle, moral choice, trial, gatekeeper, or contest
Helpful agent	Grants magical, technical, or social aid	Object, mentor, animal helper, algorithm, map, key, or companion
Struggle	Produces direct confrontation	Battle, debate, chase, trial, negotiation, or symbolic contest
Victory or solution	Resolves the central conflict	Defeat, cure, discovery, reconciliation, or exposure of deception
Return	Reorients the story toward the initial world	Journey home, social reintegration, or public report
Recognition or new order	Confers identity, reward, marriage, justice, or transformed community	Coronation, acceptance, ethical lesson, restored relation, or institutional change

The difference between the surface motif and the structural function is where Propp has a value for the AI analysis of narratives. It's easy for an LLM to generate space stations, castles, detectives, or keys or forests. The trickier question is whether or not these elements play similar roles within a clearly related line-up. In this understanding, Propp's work is a scientific terminology for them to use in the assessment of structural mimicry, and not a superficially decorative genre imitation.

Joseph Campbell's Monomyth

Unlike Propp's morphology, Campbell's Hero's Journey is based on and emphasizes the traditionally heroic story rather than myth. While Propp's model is based on a body of material from Russian wonder tales, Campbell (2008) suggests a basic comparative mythic pattern that follows departure, initiation and return. The non-ordinary world from which the heroine or hero comes; the situations he or she faces, the others who support him/her, the experience he/she survives, the prize he or she earns, the new understanding the heroine or hero gains. Despite criticisms of flattening out different traditions, the monomyth is still very powerful and taken for granted in popular narrative explanations, screenwriting, game design and teaching.

It can be seen that the framework that has been used in this study is shown in figure 2.

The model was not used as a universal law but rather as a culturally potent model template that is likely to be a part of the training data of LLMs. So it can be suggested that because it is included in it, it contributes to contrasting the deep mythic sequence, popular story formula and algorithmic reproduction.

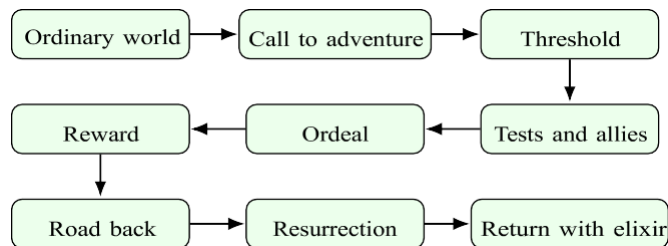


Figure 2: Campbell's Hero's Journey Framework

This is particularly pertinent to LLMs as many prompting contexts involve having the model create stories that explicitly have points of transformation. Even if prompts are not specific to a mission or to progress in general, models often imagine heroic progression as it is the common narrative default. The framework can then be applied to the extent of imitation of folklore morphology as well as the level of absorption of modern pedagogy of narration.

Narrative Syntax and Narratology

Classical narratology modified structuralist thinking in several ways: in differentiating "Story" from "Discourse", in separating "time" from "voice", and in separating "voice" from "mood". Genette (1980) demonstrated that the order of presentation can be different from chronological order and Chatman (1978) distinguished between the what of a narration and the how of a narrative transmission. Bal (2009) and Prince (1982) have further developed concepts of focalization, story logic, eventhood and experientiality, and Fludernik (1996) and Herman (2002) have refined notions of focalization and story logic respectively. These ideas are relevant to AI-generated fiction as an LLM can give the illusion of a coherent discourse surface, while not necessarily providing coherent fabula relationships. One character can be born a child and then become an elderly character without giving an explanation in terms of age; a quest item might disappear from the causal chain; an ending can be an announcement of peace without any explanation of how it came about.

In the context of this article, Narrative syntax is the sequentiality between functions and events, roles and resolutions. It doesn't mean that stories are mechanically formula able. Instead, it reveals the structures that readers have and need to confirm to see that there is progression, causality, and closure. This expectation can be a bit upended by LLMs, which are local to produce text, and act as if they are global entities. Their syntax relating to the use of narratives is an empirical issue that can only be addressed by measuring their frequency and closely reading them.

AI, Computational Creativity, and Story Generation

This is not the first time computational story generation has been attempted, as it has been around for quite a while, even before LLMs. Early systems, like TALE-SPIN, portrayed stories as a set of goals, plans and character beliefs (Meehan, 1977). In later

planning systems, the concept of narrative was regarded as a relationship between author's intent and character's power rather than author's power (Riedl & Young, 2010). In order to achieve data-driven generation instead of rule-based generation, researchers tried to extend the story with plot events, use hierarchical approach to generate the story as Fan et al. (2018) and apply pretrained language model as See et al. (2019). Recent LLM research has further highlighted the importance of collaborative, prompting, controllability and text degeneration (Holtzman et al., 2020; Yuan et al., 2022).

Table 2 shows how this article fits in the context of some of the previous work. As seen in the table, the fields of computational creativity and large-scale neural language generation (which includes large-scale modeling) have now gained more traction, while still, structuralist evaluation is underdeveloped.

Table 2: Previous Studies on AI Narrative Generation

Study	Approach	Main contribution	Relevance to this article
Meehan (1977)	Symbolic planning	Modelled narrative through goals, beliefs, and problem solving	Demonstrates an explicit grammar of story causality
Riedl & Young (2010)	Narrative planning	Balanced plot coherence with character believability	Provides a computational model of structural constraint
Fan et al. (2018)	Neural generation	Generated stories through hierarchical prompts and continuations	Shows the transition from planning to neural fluency
See et al. (2019)	Pretrained language models	Compared pretrained models with specialized story systems	Identifies strengths and weaknesses of pretrained storytelling
Holtzman et al. (2020)	Decoding analysis	Explained degeneration and repetition in neural text	Helps interpret recursive loops and local repetition
Yuan et al. (2022)	Human-AI writing tool	Studied LLM collaboration in creative drafting	Connects generation to practical writing workflows
Ippolito et al. (2022)	Writer-centered evaluation	Examined professional perspectives on AI writing assistance	Supports the need for human revision of AI structure

In the literature investigated it was seen that LLMs can generate good local fluency but still lack long range control of the narratives. This restriction refers to the idea of coherence, hallucination or repetition but it can be seen narratologically as a disruption of a function sequence. Thus the present study is involved with the computational evaluation and structuralist theory.

Research Gap

This is the research gap arising from a cross section of three areas. It also shows that

other work in the field of literary theory, the structuralist, provides clear models of narrative syntax, and very few test these on AI-generated narrative. In computational linguistics some of the results are measured and judged against models, though these are not necessarily very detailed narratological categories. Computational creativity takes into account novelty and collaboration with no guarantees of meaningful surface imitation and deep structural progression. The article fills the void by implementing the models of Propp and Campbell as operating models, and revealing disturbances that do not fit classical models. What remains is a theory formed by algorithmic narrative syntax, a new form of machine narrative, that is of an emergent machine organization of stories.

Theoretical Framework

Structuralist Narratology

The starting point of the theoretical part is the theory of the narrative as an organization of functions, roles and temporal dimensions etc. (structuralist narratology). In this tradition, significance of an event is dependent on it occurring within a sequence. The departure, therefore, must have occurred in relation to the lack or the call, the victory must have happened as a resolution to the struggle and the return must have happened because it returned or reconnected the protagonist who was changed due to the departure to the initial world. Thus, structuralist narratology presents a model to analyse what is not only expressed but also the arrangement of the elements in LLM narratives.

Propp's Morphological Framework

Propp's morphology is the source of the notion of functional sequence. In this study the Proppian functions are represented as the type of events that can occur in the stories and have been annotated as such; they can be present in a story even if it is not a folkloric one, if it was generated by an AI system. Lack or villainy can take the form of a missing data archive or an astronaut who can no longer be found, a corrupted village well, or a stolen family heirloom. The abstraction is significant as genres are often modernized or hybridized by LLMs. The Proppian device makes it possible to analyse this kind of variation, without being taken in by the diversity of motifs and attributing it to difference between the structures.

Campbell's Hero's Journey

The monomyth concept is borrowed from Campbell and is called the Transformative Arc. The Hero's Journey is described as 'the path of departure, initiation and return rather than a universal specifically to be followed'. The scaffolding is particularly helpful in LLM storytelling contexts, as numerous stories can be created involving protagonists that are engaged in a process of self-discovery, ethical awakening, or community restoration. The model allows to determine if this transformation is only claimed at the end or is rather a structural gain.

Algorithmic Narratology

Algorithmic narratology' is the study of narratively structured behaviors that are produced or mediated by computational systems. It assumes that LLM's don't just copy classroom structures, but manipulate it with its probabilistic generating nature. The output of the model is obtained from the method of prediction of the tokens, the

distribution of the training, the design of the prompts, and the decoding parameters and the restrictions for alignment. It follows that the narrative form itself is rendered a pseudoproduct of liturgical traditions within a given culture and the algorithmic activity. The main result that facilitates the integrated theoretical framework of the study is given in Figure 3.

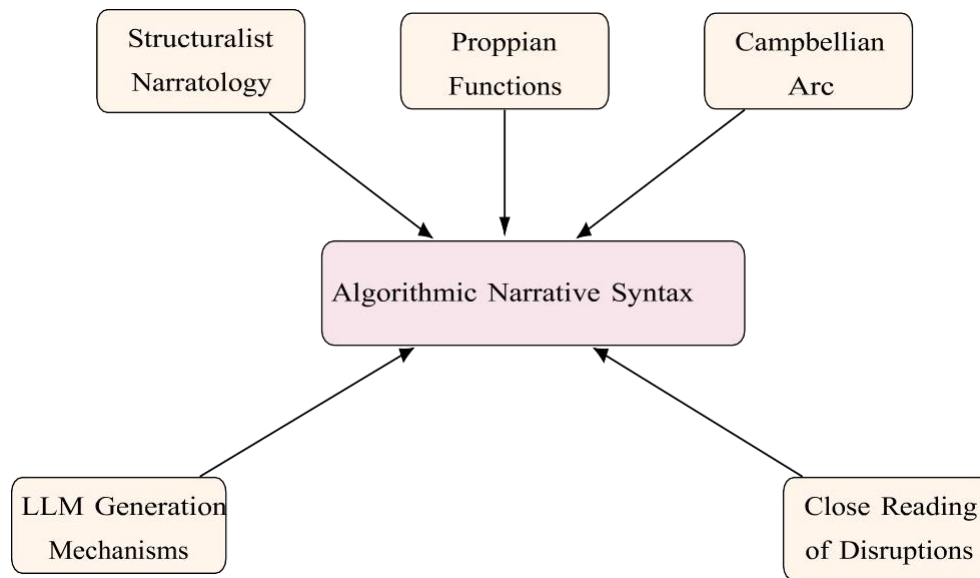


Figure 3: Integrated Theoretical Framework Linking Structuralism and AI Narratives

The framework suggests it's important to read out the narrative produced by AI in two levels. On the first level, LLMs create classical structures, and the functions and arcs they create are familiar. At the 2nd tier, they tamper with these structures with drift, recursion, fragmentation and unstable resolution. Algorithmic narrative syntax becomes the main concept of the article out of this tension between these levels.

Research Methodology

Research Design

Computational mapping along with qualitative close reading were embraced as mixed method design to capture the study. For the quantitative part, the number and occurrence of Proppian functions and Hero's Journey stages were measured in 500 narratives generated by the LLM, and also the frequency and distribution by the model level was determined for these elements. The second phase – qualitative – explored selected cases of syntax of the classical narrative subverted by algorithmic generation as a means to interpret its transformation. That design was adopted because according to the structuralist narratology, there is a need for a formal abstraction as well as interpretative judgments. Counting functions won't be enough to explain why an ending does not feel earned, and close reading won't be enough to determine if the disruption is frequent throughout a corpus.

The overall research design is presented in figure 4. The Corpus underwent Preprocessing, Computational Annotation and Human validation after which statistical analysis and interpretive close reading were performed. That's because the

design shifts from production (controlled generation) to theory building.

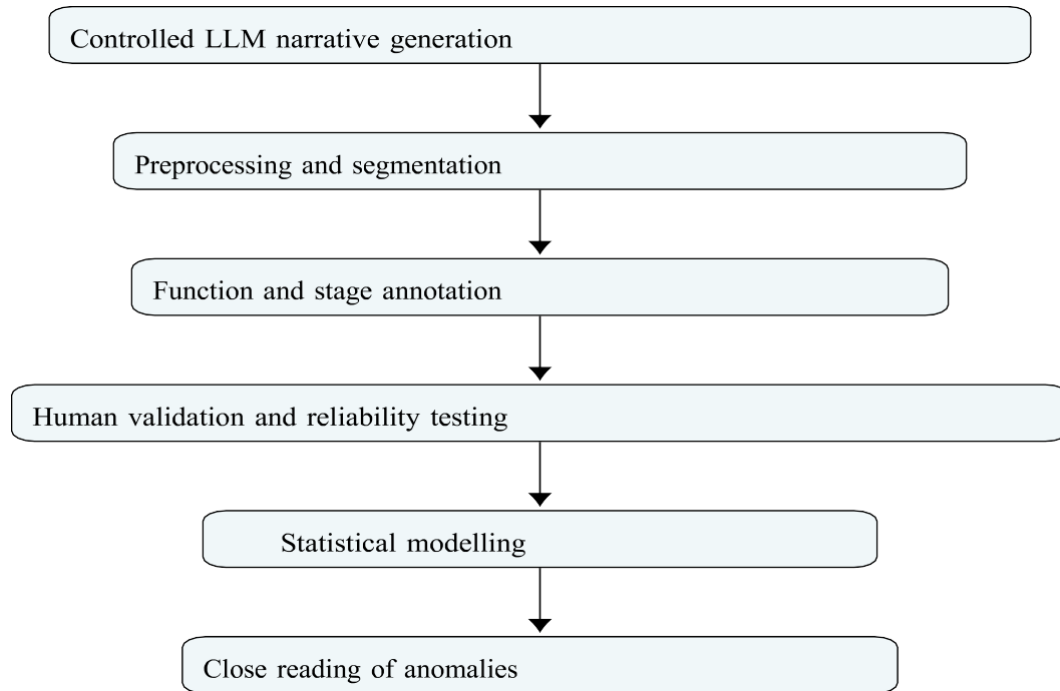


Figure 4: Overall Research Design

Dataset Collection

There were a total of 500 English stories produced by GPT, Claude, Gemini, and Llama model families. A total of 125 narratives was created from the models by each family. Four generic types of narrative situations were balanced: a folk-tale quest; a contemporary moral problem; a quest for a lost object; a quest for speculation. Structuralism, monomyth, Campbell and Propp were not mentioned in the prompts. This disclaimer was added to ensure that they aren't imitating a requested template, but implying the same structure instead. The outputs were produced using a moderate level of creativity (about 0.8 when possible), using the lengths as targets of between 700 words and 1,200 words. Table 3 gives an overview of the Corpus.

Table 3: Dataset Characteristics

Model family	Narratives	Mean words	SD	Mean paragraphs	Prompt distribution
GPT	125	921	148	8.7	32 quest, 31 dilemma, 31 mystery, 31 speculative
Claude	125	956	137	9.1	31 quest, 32 dilemma, 31 mystery, 31 speculative

Gemini	125	904	162	8.4	31 quest, 31 dilemma, 32 mystery, 31 speculative
Llama	125	842	171	8.0	31 quest, 31 dilemma, 31 mystery, 32 speculative
Total	500	906	158	8.6	125 narratives per prompt type

The data set was intended to be as cross-model as possible while having sufficient variation in genre to assess structure generalisability. Assigning equal status to narratives eliminated the possibility of model family size having any bias on estimates of aggregate frequencies. The corpus was 452,875 words and 4,294 paragraphs.

Quantitative Narrative Analysis

Quantitative analysis was done in a three-stage computation process. Narratives were segmented into event units, according to paragraph boundaries and event markers (at the clause-level or above) in the first step. Second, we obtained an event unit level alignment between the Proppian function and the Hero's Journey stages by following an hybrid annotation procedure based on the lexical cues, dependency pattern, semantic role labelling, and transformer-based sentence embeddings. The outputs were then examined by two trained researchers against a text manual for validation. Based on the computational tool, coders were human coders who accepted, rejected or revised the candidate labels. This hybrid approach enabled to prevent complete automation of structural interpretation and maintained the scalability.

The computational work flow is shown in figure 5. The workflow highlights that a story structure was determined based on relationships between events, and not on key words alone.

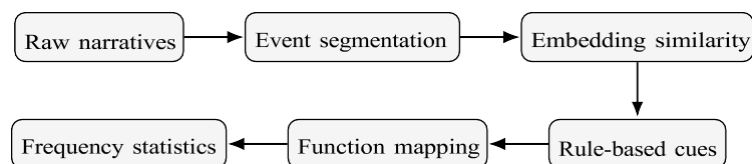


Figure 5: Computational Narrative Analysis Workflow

Analysis was done using scores of Proppian compliance as well as Hero's Journey coverage score and a score of narrative coherence for each narrative, and finally disruptions were counted. The 12 collapsed functions were rated by Proppian compliance, both as to presence as well as to relative order. Twelve stages in the Hero's Journey and if the story consists of a return stage, were measured as Hero's Journey Coverage. Causation (causal sequencing), character (consistency of the character) and thematic (thematic continuity) and resolution (adequacy of the resolution) were judged on a five-point scale in terms of coherence.

Qualitative Close Reading

The narratives in the purposive subsample of 80 were analysed using qualitative close reading. Twenty narratives from each model family were used in the subsample and were systematically sampled from three categories of degrees of structural compliance: high, medium and low. Close reading centered on where (according to the computed mapping), the sequence was interrupted, the same functional blocks were used repeatedly, there was sudden change in log-semantic field, characteristics of characters were not in keeping with one another or the action was motiveless. It did not simply aim at error identification, but also at understanding the error and regarding it as formal aspects of the process of composing an algorithmic narration.

Reliability and Validity

Double coding was used for 120 narratives (24% of the texts) to test for reliability. The Cohen's kappas were .82 for categories of Proppian function labels, .79 for Hero's Journey stages, and .76 for categories of disruption. Agreements or disagreements were adjudicated and revisions made in the annotation manual. The internal consistency for the 4 coherence dimensions was good with a Cronbach's Alpha of .84. Triangulation of computational labels, human interpretation and model level comparison was employed as a method of establishing construct validity. Prompt types and the approximate length of the prompts were also controlled for to prevent some of the differences in structure from simply representing the prompt imbalance.

Results and Findings

Structural Mimicry of Classical Narrative Frameworks

The first key result is that LLM-generated texts show a strong mirroring of surface constructions of classical narrative constructions. Table 4 shows the frequency of twelve function categories of Proppian system in all the 500 narrated stories. Initial lack/absentation occurred in 438 or 87.6% of the narratives. The most frequent words appearing in the narratives in addition to the four central words were mediation or call (421), departure (407), villainy or intensified lack (392), struggle (384). The average number of Proppian categories in a narrative is 7.95 ± 2.11 categories.

Table 4: Frequency of Propp's Narrative Functions in LLM Narratives

Proppian function category	Narratives containing function	Percentage	Mean sequence position
Absentation or initial lack	438	87.6%	1.8
Interdiction or rule	286	57.2%	2.7
Violation	251	50.2%	3.2
Villainy or lack intensification	392	78.4%	3.5
Mediation or call	421	84.2%	4.1
Departure	407	81.4%	4.8
Donor test	356	71.2%	5.9
Helpful agent	337	67.4%	6.4

Struggle or confrontation	384	76.8%	7.6
Victory or solution	362	72.4%	8.5
Return	304	60.8%	9.3
Recognition or new order	279	55.8%	10.1

This same distribution can be seen in a visualization as in figure 6. The descending tendency suggests that the LLMs are better at replicating the opening and middle part of the classical morphology of a function than the terminal component of a return and of recognition. By this I mean, created tales start a well known construction of a problem and steer a protagonist to someplace embedded in a struggle. Are less consistent in restoration of the social order, or incorporating the transformation of the protagonist into a social whole.

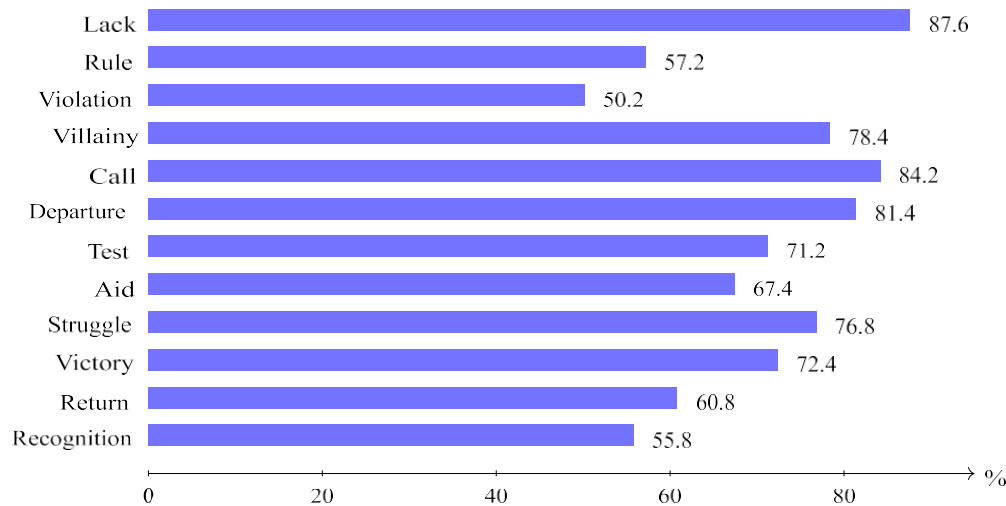


Figure 6: Distribution of Proppian Functions Across Narratives

This approach confirms the supposition that LLMs (also) adopt the syntax of story narratives on the level of often replikaable functional units as features of structuralism. These are also common occurrences across the models, including the verbs lack, call, depart, test, aid and struggle, pointing to the fact that such patterns are well learnt from the text distributions their LLMs use. The lower rates of return and recognition indicate that the sequence generated is not necessarily structurally complete, though. The models are able to put together the grammar of adventure but not necessarily the morphology of restoration.

Hero's Journey Analysis

The result of the Hero's Journey analysis was somewhat similar, showing compliance with the partial structure. Table 5 indicates that the ordinary world was found in 456 of the narratives; the call to adventure in 430 narratives; the threshold crossing in 398 narratives; tests and allies in 386 narratives; and an ordeal in 367. This is assuming these higher frequencies are a strong representation of the departure and initiation aspects of the monomyth by LLMs. But they were less inclined to mention the other stages: Road back was mentioned in 276 narratives, revival in 249 narratives, and return with elixir in 318.

Table 5: Occurrence of Hero's Journey Stages

Hero's Journey stage	Narratives containing stage	Percentage	Mean sequence position
Ordinary world	456	91.2%	1.2
Call to adventure	430	86.0%	2.4
Refusal of the call	261	52.2%	3.1
Mentor or supernatural aid	348	69.6%	3.8
Crossing the threshold	398	79.6%	4.4
Tests, allies, and enemies	386	77.2%	5.7
Approach to the inmost cave	329	65.8%	6.5
Ordeal	367	73.4%	7.3
Reward	341	68.2%	8.1
Road back	276	55.2%	8.8
Resurrection	249	49.8%	9.4
Return with elixir	318	63.6%	10.0

The completion rates are reported in Figure 7. A total of 142 narratives (28.4% of the corpus) contained all twelve stages arranged in a reasonable sequence, that is where all the stages of the journey were present. 157 more stories had functional completion, which generally is the presence of eight or more notations including a return. The remainder of the narratives (201) were incomplete or had been significantly altered.

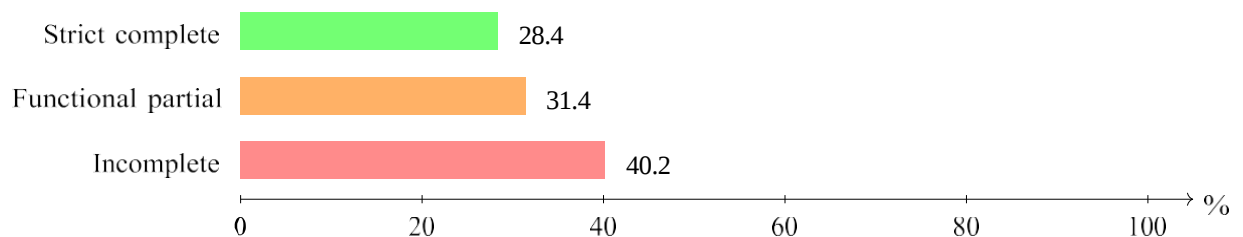


Figure 7: Hero's Journey Completion Rate

The Hero's Journey Outcomes helps to bring out the difference between stage occur and stage integration. LLMs are good at generating a plot involving a loss of a normal world and a trial but are not always good at generating a symbolic return which imparts Campbell's arc with its transformative power. Theories are interesting here because they demonstrate that AI narratives frequently mimic the Trojan Hero story of a transformation by the end of the narrative, but do not build a complete experience of a story.

Narrative Coherence Performance

In the coherence analysis, it was determined whether the causal and characteristical elements of the structural functions were coherently integrated or not. Table 6 shows means coherence ratings (on a 5-point scale) for model families. GPT attained the highest overall coherence score overall at 4.22 followed by Claude at 4.18. The results from the two models are: 3.94 (Gemini) and 3.61 (Llama). Results of a one-way analysis of variance revealed significant differences among the four families of models, $F(3,496) = 21.64$, $p < .001$, and $\eta^2 = .116$. The post hoc showed that there was no significant difference between GPT and Claude, however, they both performed significantly better than Llama in each of the coherence dimensions.

Table 6: Narrative Coherence Scores by LLM

Model family	Causal sequencing	Character consistency	Thematic continuity	Resolution adequacy	Overall coherence
GPT	4.29 (.56)	4.15 (.61)	4.24 (.54)	4.18 (.66)	4.22 (.48)
Claude	4.18 (.58)	4.28 (.52)	4.12 (.60)	4.13 (.62)	4.18 (.46)
Gemini	3.97 (.68)	3.89 (.71)	4.03 (.63)	3.86 (.75)	3.94 (.57)
Llama	3.62 (.76)	3.53 (.81)	3.76 (.70)	3.52 (.84)	3.61 (.66)
Overall	4.02 (.70)	3.96 (.73)	4.04 (.65)	3.92 (.78)	3.99 (.61)

The overall scores are represented in terms of Figure 8. The principle to interpret is that 'coherence' is related to 'structural' compliance, however it is not synonymous. There were some tales that had several Proppian functions but their unmotivated transitions or unstable character resulted in failure to appreciate the tale. A long story, in which only a few classical functions were present, could instead be awarded a high coherence rating if the little structure, which it actually has, was well integrated from a causal point of view.

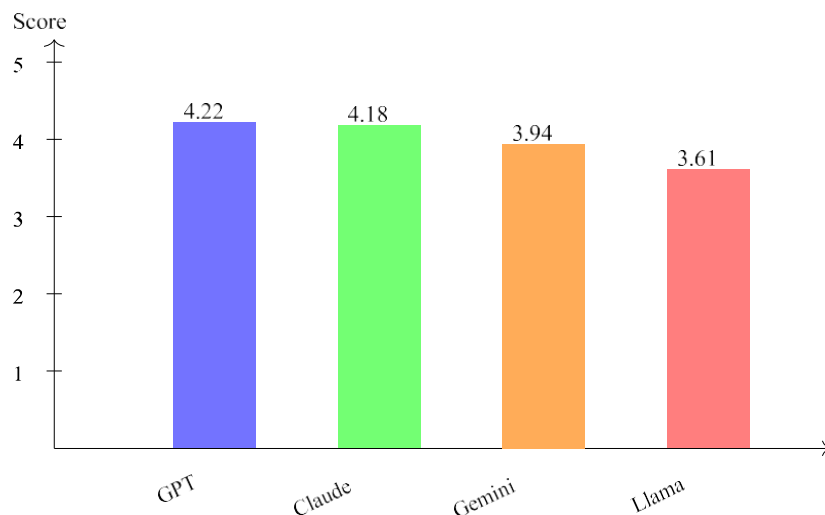


Figure 8: Comparative Narrative Coherence Scores

While these findings complicate simple assertions that LLMs could or could not cohesively narrate, LLMs may narrate, but this capability depends on the datasets and

training methods they undergo, which are frequently challenging to access. They can generate sequences that can be coherent, if the sequence is relatively short and can be conventional. The problem arises when several functional dependences have to be carried along in a story for quite a few paragraphs.

Identification of Algorithmic Disruptions

Each type of algorithmic disruption identified by the qualitative and computational analyses is considered a major type of algorithmic disruption. Their frequency is given in table 7. 268 (53.6%) of the narratives, in total, contained at least one disruption. Most common were the number of times plot fragments occurred, then the number of times in semantic drift, non-sequential resolution, character inconsistency, recursive loops, and temporal compression. Fifty-four-one (541: 1.08) disruption(s) occurred in the overall corpus.

Table 7: Types of Algorithmic Disruptions

Disruption type	Narratives affected	Occurrences	Rate per 100 narratives
Semantic drift	102	116	23.2
Recursive narrative loop	61	72	14.4
Character inconsistency	79	91	18.2
Plot fragmentation	96	110	22.0
Non-sequential resolution	82	94	18.8
Temporal compression or unexplained ellipsis	51	58	11.6
Total	268 unique narratives	541	108.2

The relative frequency of these types of disturbances is shown in Fig. 9. It is a patterned formal phenomenon, revealed by the distribution, that would not be best characterized by random noise. The most frequently cited issues involve problems in continuity of the semantic fields and of causal blocks – exactly the areas where long-range integration of the narrative is needed.

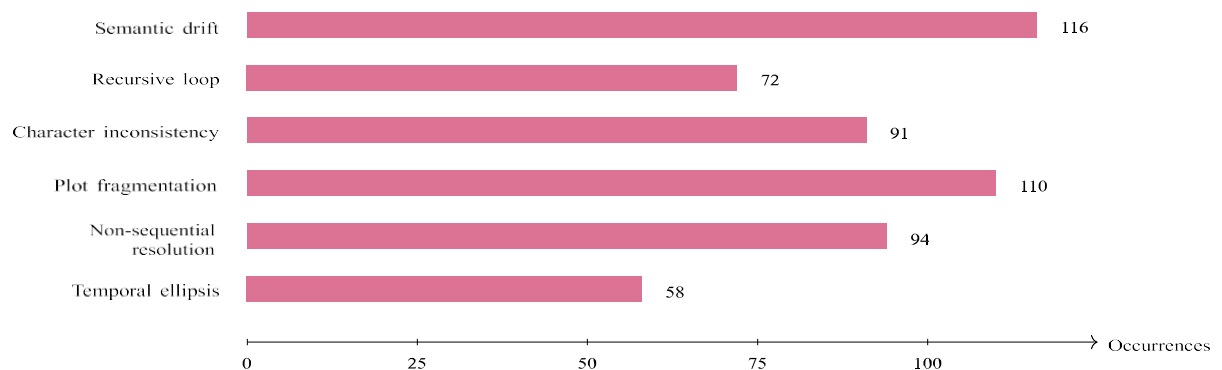


Figure 9: Frequency of Algorithmic Disruptions

The discovery sustains the article's main thesis; that LLMs replika classical functions

whilst producing a unique syntax of breakdown. These turns around aren't just punctuation or style errors. They take place at the plot, semantic continuity and resolution sequence levels.

Semantic Drift Analysis

When a story started in one semantic area and drifted to another semantic area, but no mediation was offered, this was considered to be semantic drift. Table 8 offers some sample examples from the closereading sample. In one GPT story a medical apprentice looking for a cure morphed into a cartographer figuring out a lost map with the illness not being referenced until the end, at the moral. A missing violin mystery is turned into a courtroom drama without a legal battle, when both characters are engaged in a tale of rehabbing a stolen classical instrument. The examples show that LLMs can keep the theme without the original causal object of the story, which involves it in the organization of the story.

Table 8: Examples of Semantic Drift

Case	Model	Initial semantic field	Drift pattern	Drift score
SD-014	GPT	Village medicine and cure quest	Cure narrative becomes map-decoding quest; illness returns only as moral image	.38
SD-037	Claude	Archive mystery and missing manuscript	Manuscript search shifts into family reconciliation without evidential bridge	.31
SD-052	Gemini	Missing violin in urban school	Investigation becomes courtroom redemption despite no trial setup	.46
SD-071	Llama	Desert pilgrimage and water scarcity	Pilgrimage turns into mechanical repair story; sacred spring disappears	.49
SD-088	Llama	Space mission rescue	Rescue plot becomes diplomatic election narrative with no transition	.52

Semantic drift is shown as a trajectory in Figure 10, out from the initial narrative vector. The line itself is not one of stories, but rather a conceptual model to which embedding-distance analysis was used. Narrative coherence was low with a small semantic distance between beginning, middle and ending in the stable narratives. Drifted narratives accounted for a significant amount of the stories afterwards.

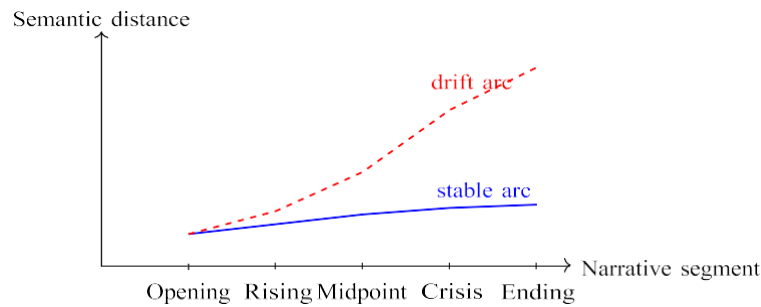


Figure 10: Semantic Drift Trajectory Model

How there is semantic drift is narratologically important because it will show the difference between "thematic association" and "structural causality. The drifting stories were often well written or well balanced in emotion, but the "main thing" desired/conflicted was changed. This creates a narrative with layers of significance and punctuates each sentence but that, as a whole, subverts the sense of stability.

Recursive Narrative Loops

Cyclical rehashes where some part of a story was repeated but otherwise not contributing to the story. The number of times that loops have been encountered for each model family are shown in Table 9. The most number of loop occurrences was seen in llama with 34 loop occurrences out of 28 narratives, while the least number of occurrences was seen in Claude with 8 loop occurrences out of 7 narratives. The mean loop length is 2.9 segments, which means that generally loops contained a number of paragraphs, but not repeated phrases.

Table 9: Recursive Loop Occurrences by Model

Model family	Narratives with loops	Occurrences	Mean loop length	Predominant trigger
GPT	9	10	2.4 segments	Reflective moral restatement
Claude	7	8	2.1 segments	Reassuring concession and return to premise
Gemini	17	20	2.8 segments	Quest reinitialization after partial success
Llama	28	34	3.3 segments	Repeated threshold crossing or renewed warning
Total	61	72	2.9 segments	Repeated function without changed state

The recursive loop is the typical loop structure as illustrated in figure 11. The main character has a call, leaves and faces a test, comes back to a warning or new call then leaves again, never is it repeated. If the repetition is performed in a text produced by human beings (as in fictional writing), it can lead to a sense of ritual, of ironization, and/or of psychological emphasis. However, it normally occurred as structural redundancy without marking, as in the case of the LLM corpus.

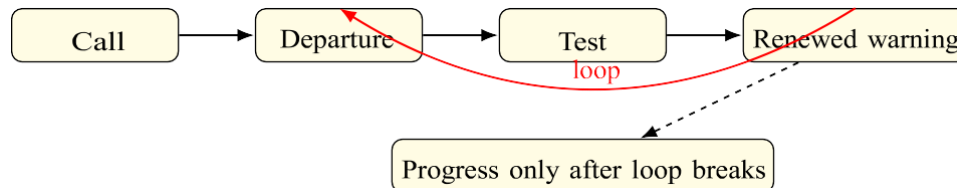


Figure 11: Example of Recursive Narrative Loop Structure

Analysis of recursive loops (RLs) shows that LLM's can generate local narrative momentum while lacking a suitable model of the local story state change. The model is able to associate warnings, departures and tests together, but could repeat the regeneration of the same cluster if it were unsure.

Non-Sequential Narrative Resolutions

Non-sequential resolution was when a narrative solved a conflict out of order (without achieving the necessary structural elements) or used closure without solving/omitting the event that ultimately would settle the conflict. There are five resolution-error patterns reported in Table 10. The most common pattern was an error of premature victory before ordeal, that is, 30.9% of the resolution errors were such errors. One also lost a lot of returns and reversions after the closing as well.

Table 10: Non-Sequential Resolution Patterns

Resolution pattern	Occurrences	Share of resolution errors	Narrative effect
Premature victory before ordeal	29	30.9%	Conflict ends before confrontation is narratively earned
Unexplained return	21	22.3%	Protagonist reappears in ordinary world without transition
Reversal after closure	17	18.1%	Ending is contradicted by a renewed crisis
Resolution without agent	15	16.0%	Problem disappears without protagonist, helper, or antagonist

			action
Causal gap between crisis and reward	12	12.8%	Reward is granted without visible solution
Total	94	100.0%	Closure loses functional sequence

Figure 12 visualizes the distribution. These results clarify why return and recognition had lower frequencies in the Proppian analysis. LLMs often know that a story should end in restoration, but they may not maintain the causal chain required to justify restoration. The result is a non-sequential resolution: closure as stylistic declaration rather than structural consequence.

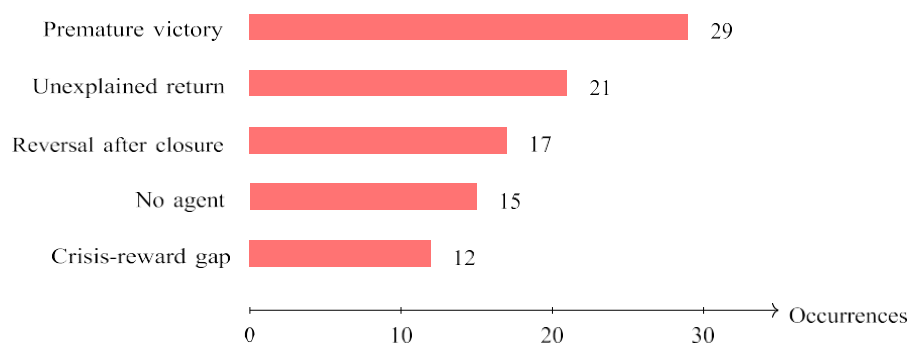


Figure 12: Distribution of Narrative Resolution Errors

The close readings showed that non-sequential resolution often appeared after strong beginnings. A narrative might establish a curse, a rule, a donor test, and a confrontation, then abruptly announce that the kingdom was healed because the protagonist “understood the truth.” Such endings imitate the rhetoric of moral completion while bypassing the event logic of classical morphology.

Comparative Performance of LLMs

Table 11 compares model families across structural performance metrics. GPT and Claude had the strongest balance of structural compliance and coherence. Gemini showed moderate compliance but higher disruption frequency. Llama produced the lowest Proppian compliance score and the highest disruption rate. The overall correlation between structural compliance and disruption frequency was negative and moderate, $r = -.42$, $p < .001$, indicating that higher compliance tends to reduce but does not eliminate algorithmic disruptions.

Table 11: Structural Performance Comparison Across LLMs

Model	Propp	Hero	Complete monomyth	Coherence	Disrupt.	Syntax index
GPT	.79	.71	45 (36.0%)	4.22	.84	.18

Claude	.78	.70	42 (33.6%)	4.18	.78	.17
Gemini	.74	.66	34 (27.2%)	3.94	1.10	.29
Llama	.67	.60	21 (16.8%)	3.61	1.60	.53
Overall	.75	.67	142 (28.4%)	3.99	1.08	.29

Figure 13 plots structural compliance against disruption frequency. The visualization demonstrates that model performance is not simply a matter of producing more classical functions. The best narratives sustain both structural progression and local coherence. The weakest narratives produce a recognizable story-shape while allowing drift, loops, or premature closure to interrupt the sequence.

Disruptions per narrative

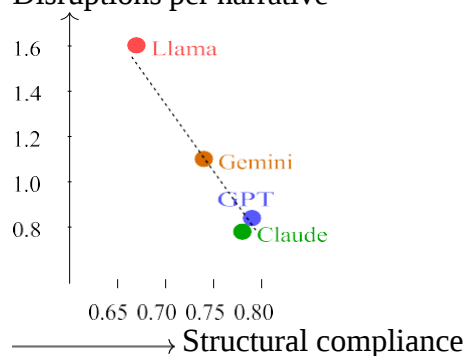


Figure 13: Structural Compliance vs. Narrative Disruption Frequency

The comparative analysis supports the article's main claim. Across all models, LLMs can imitate classical narrative syntax, but their imitation remains probabilistic and uneven. Even highperforming models produced examples of semantic drift and non-sequential closure. Algorithmic narrative syntax therefore cannot be equated with failure. It is better understood as a hybrid form in which classical patterns are recombined under the constraints of language-model generation.

Discussion

Why LLMs Successfully Mimic Structural Syntax

LLMs are able to emulate structural syntax because the textual corpora they are trained on comprise of innumerable narratives, which have been influenced by inheritance cultural forms. Lack, call, departure, trial, aid, struggle and resolution are coded in fairy tales, novels, myths, screenplays, game 'quests' and children's stories and in writing manuals as well. This allows models to pick up statistical regularities from these corpora, and as well, that certain event types tend to be followed by others. This does not mean that the model has a Proppian theory. Instead, that's saying that the structuralist regularities become embodied in language patterns as probabilistic. In the corpus, one will find high frequencies of lack, call, departure, struggle, which indicate that the classical functions are not out of the way of the LLM generation, but on the contrary belong to the distributional memory of the culture of text.

Causes of Algorithmic Disruptions

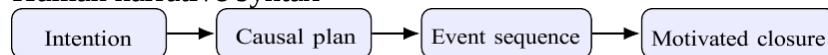
It is from the local continuation – global narrative planning difference that the disruptions identified in this study come. The main hurdles in generating sequences include maintaining consistency across a sequence of events, referred to as story state,

which is a key aspect of achieving coherence in the narrative. Semantic drift: Due to the local association, the model is drawn towards a new field that is stylistically compatible, but not causally motivated. Recursive loops are when recurrence of a functional cluster can be generated that remains locally plausible and is already known. Character inconsistency is when there is a lack of identity attributes in different contexts. Plot fragmentation is when some subplots are included which are not integrated. 'Narratively earned' closure does not need to be modelled, as nonsequential resolution takes the model's discourse marker of closure as its starting point.

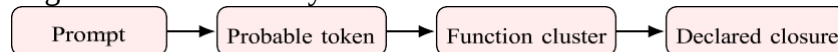
Emergence of Algorithmic Narrative Syntax

The results indicate that AI's algorithms make both written and verbal stories, and point to the need to interpret these narratives as algorithmic narrative syntax. It is not that this is against classical narrative syntax, it depends on this by parasitic mode. Since structures are common, easy to teach, and written, LLM's mimic structures. However, they also disrupt the models, since there is no stable causal memory in probabilistic generation. The difference between the two (human narrative syntax and algorithmic) is shown in Figure 14. It does not mean that it is never done by people, precisely because people sometimes produce disjointed stories, improvise and revise. The difference is that, while somewhat rhetorically motivated, human fragmentation can in many cases come about when it is produced algorithmically, causing fragmentation as part of generation.

Human narrative syntax



Algorithmic narrative syntax



drift or loop

Figure 14: Human Narrative Syntax vs. Algorithmic Narrative Syntax

The algorithmic narrative syntax thus can be considered a double movement. The first is the structural mimicry: the authentication and production of recognisable functions of a narrative on the one hand, and stages on the other. The second disruption is that of structure, i.e., the absence of sustaining function, character and causality throughout the story. It is this double movement that led to the sensation of immediate legibility, and eventual dissatisfaction at close reading of many LLM stories.

Implications for Literary Theory

The results indicate that structuralism is still very relevant in the era of generative AI, for literary theory. Propp and Campbell assist in determining the templates repeated by LLMs, and Genette, Chatman, Herman, and Ryan assist in explaining the failures of sequence, discourse, and continuity of story world. Meanwhile, fictionalization with the help of AI demands structuralism. Classical models are designed for narratives which are human generated or taken up by the culture. Already we can see from the output of LLMs that it is possible to replicate narrative syntax without

intending to use something that is traditionally called “authoring”. This questions theories that attempt to bind too tightly the concept of narrative form to that of human consciousness and computational-based theories that consider narrative as nothing more than fluent text.

Implications for Computational Linguistics

The study may prove to be of interest for computational linguistics, since it highlights the potential of narratological evaluation. There is no way to differentiate between a sentence carefully written and the earned out resolution, by using the standard fluency measures. Likewise, similarity scores for semantics can be lacking with regard to a missing return or repeated donor test. Thus, to evaluate long-form generation, it is required to check for narrativestate tracking, role consistency, causal dependency mapping and resolution adequacy. Structuralist categories can be valuable diagnostic categories for the evaluation of the model.

Implications for AI-Assisted Creative Writing

The study reveals that LLMs can be used as valuable tools to achieve AI-based creative writing, including the generation of story structure, premise, motifs, and traditional story arcs. They can't write structurally complete stories as independently. All authors of LLM-generated content should therefore consider LLM content as a draft and use human structural editing of the text. Revision should aim at maintaining the thing desired, showing (but not summarizing) a character change, keeping the functional clusters from repeating, and keeping closure coming as a result of narrated action (not declarative summary). In this way, the practical implication of the study does not suggest the rejection of AI-based storytelling, rather, understanding the algorithmic syntax, and editing it.

Conclusion

In this article, the reproduction and disruption of classical structuralist syntax in narratives created by LLM were explored. The study comprised 500 narratives generated by GPT, Claude, Gemini and Llama model families to evaluate structural mimicry with a mixed methods design and revealed a strong case for structural mimicry. Major Proppian functions, like lack, call, departure, struggle and victory were used at high frequency and major Hero's Journey Stages, like ordinary world, call, threshold crossing, and ordeal were also well represented. These results demonstrate the successful depiction of many aspects of the classical narrative forms on the surface level by LLMs.

The key theoretical proposal of the study is that of algorithmic narrative syntax. It is with this concept that we can explain why there can be LLM stories that are at the same time structurally recognizable but at the same time structurally unstable. Models that mimic historically transmitted narrative conventions (indirectly, through text culture), while at the same time, disrupting the conditions of continuity of long-range casual and semantic as well as characterological development which is not always preserved in probabilistic generation.

The study provides a scheme for considering evaluations of AI-generated fiction that goes beyond a focus on style fluency. Researchers, educators, writers, and system designers can use Proppian functions and stages of Hero's Journey, coherence dimensions, and disruption categories to diagnose strong and weak points in the

generated narratives. These findings indicate that LLMs are not simply a replacement for structural human judgement but should be thought of as an aid to drafting a narrative. There are some limitations to the study. As it centered on texts written in English and on four model families, the results can not be transferred to other languages, genres or generation systems. It employed prompt which was designed to be comparable, but real life users prompt iteratively. The other systems, the Proppian and Campbellian, also appear as culturally determined and are far from being able to catalogue all the traditions of world narrative. Finally, annotation was reliable, but involved interpretation.

Future research should include: multilingual AI narratives, reader reception, longer forms of narratives, and prompting which involves interacting with the reader. Further research could involve aspects of comparison between LLM-generated narratives and novice writing control corpora to find out what aspects of disruption are inherently algorithmic and which are typical in novice writing. Studies are also needed to determine if there are ways to decrease semantic drift and non-sequential resolution using model architecture, context window, retrieval systems, or explicit planning modules. The larger message is that, generative AI isn't the death of structuralist narratology, it's just a new thing to most probably narratologize. Both a culturally inherited feature and a computably reproducible but unstable form, called 'narrative syntax', LLM shows.

References

- Aarseth, E. J. (1997). *Cybertext: Perspectives on ergodic literature*. Johns Hopkins University Press.
- Ammanabrolu, P., Tien, E., Cheung, W., Luo, Z., Ma, W., Martin, L. J., & Riedl, M. O. (2020). Story realization: Expanding plot events into sentences. *Proceedings of the AAAI Conference on Artificial Intelligence*, 34(5), 7375–7382.
- Anthropic. (2024). The Claude 3 model family: Opus, Sonnet, Haiku. Anthropic.
- Bal, M. (2009). *Narratology: Introduction to the theory of narrative* (3rd ed.). University of Toronto Press.
- Barthes, R. (1975). An introduction to the structural analysis of narrative. *New Literary History*, 6(2), 237–272.
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623.
- Boden, M. A. (2004). *The creative mind: Myths and mechanisms* (2nd ed.). Routledge.
- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., Buch, S., Card, D., Castellon, R., Chatterji, N., Chen, A., Creel, K., Davis, J. Q., Demszky, D., ...Liang, P. (2021). On the opportunities and risks of foundation models. *arXiv:2108.07258*.
- Bremond, C. (1980). The logic of narrative possibilities. *New Literary History*, 11(3), 387–411.
- Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P.,

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

- Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., ...Amodei, D. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901.
- Campbell, J. (2008). *The hero with a thousand faces* (3rd ed.). New World Library.
- Chatman, S. (1978). *Story and discourse: Narrative structure in fiction and film*. Cornell University Press.
- Chowdhery, A., Narang, S., Devlin, J., Bosma, M., Mishra, G., Roberts, A., Barham, P., Chung, H. W., Sutton, C., Gehrmann, S., Schuh, P., Shi, K., Tsvyashchenko, S., Maynez, J., Rao, A., Barnes, P., Tay, Y., Shazeer, N., Prabhakaran, V., ...Fiedel, N. (2023). PaLM: Scaling language modeling with pathways. *Journal of Machine Learning Research*, 24(240), 1–113.
- Colton, S., & Wiggins, G. A. (2012). Computational creativity: The final frontier? *Proceedings of the 20th European Conference on Artificial Intelligence*, 21–26.
- Fan, A., Lewis, M., & Dauphin, Y. (2018). Hierarchical neural story generation. *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics*, 889–898.
- Fludernik, M. (1996). *Towards a “natural” narratology*. Routledge.
- Gemini Team. (2023). Gemini: A family of highly capable multimodal models. [arXiv:2312.11805](https://arxiv.org/abs/2312.11805).
- Genette, G. (1980). *Narrative discourse: An essay in method* (J. E. Lewin, Trans.). Cornell University Press.
- Gervás, P. (2009). Computational approaches to storytelling and creativity. *AI Magazine*, 30(3), 49–62.
- Greimas, A. J. (1983). *Structural semantics: An attempt at a method*. University of Nebraska Press.
- Hayles, N. K. (2012). *How we think: Digital media and contemporary technogenesis*. University of Chicago Press.
- Herman, D. (2002). *Story logic: Problems and possibilities of narrative*. University of Nebraska Press.
- Holtzman, A., Buys, J., Du, L., Forbes, M., & Choi, Y. (2020). The curious case of neural text degeneration. *Proceedings of the 8th International Conference on Learning Representations*.
- Ippolito, D., Yuan, A., Coenen, A., Burnam, S., & Ippolito, A. (2022). Creative writing with an AI-powered writing assistant: Perspectives from professional writers. [arXiv:2211.05030](https://arxiv.org/abs/2211.05030).
- Lévi-Strauss, C. (1963). *Structural anthropology*. Basic Books.
- Manovich, L. (2001). *The language of new media*. MIT Press.
- Martin, L. J., Ammanabrolu, P., Wang, X., Hancock, W., Singh, S., Harrison, B., & Riedl, M. O. (2018). Event representations for automated story generation with deep neural nets. *Proceedings of the AAAI Conference on Artificial Intelligence*, 32(1), 868–875.
- Meehan, J. R. (1977). TALE-SPIN, an interactive program that writes stories. *Proceedings of the 5th International Joint Conference on Artificial Intelligence*, 91–98.

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

- Moretti, F. (2013). *Distant reading*. Verso.
- OpenAI. (2024). GPT-4 technical report. arXiv:2303.08774.
- Prince, G. (1982). *Narratology: The form and functioning of narrative*. Mouton.
- Propp, V. (1968). *Morphology of the folktale* (2nd ed., L. Scott, Trans.). University of Texas Press.
- Ramsay, S. (2011). *Reading machines: Toward an algorithmic criticism*. University of Illinois Press.
- Riedl, M. O., & Young, R. M. (2010). Narrative planning: Balancing plot and character. *Journal of Artificial Intelligence Research*, 39, 217–268.
- Ryan, M.-L. (2006). *Avatars of story*. University of Minnesota Press.
- See, A., Pappu, A., Saxena, R., Yerukola, A., & Manning, C. D. (2019). Do massively pretrained language models make better storytellers? *Proceedings of the 23rd Conference on Computational Natural Language Learning*, 843–861.
- Tambwekar, P., Dhuliawala, M., Mehta, A., Martin, L. J., Harrison, B., & Riedl, M. O. (2019). Controllable neural story plot generation via reward shaping. *Proceedings of the 28th International Joint Conference on Artificial Intelligence*, 5982–5988.
- Todorov, T. (1977). *The poetics of prose*. Cornell University Press.
- Touvron, H., Martin, L., Stone, K., Albert, P., Almahairi, A., Babaei, Y., Bashlykov, N., Batra, S., Bhargava, P., Bhosale, S., Bikel, D., Blecher, L., Ferrer, C. C., Chen, M., Cucurull, G., Esiobu, D., Fernandes, J., Fu, J., Fu, W., ...Scialom, T. (2023). Llama 2: Open foundation and fine-tuned chat models. arXiv:2307.09288.
- Yuan, A., Coenen, A., Reif, E., & Ippolito, D. (2022). Wordcraft: Story writing with large language models. *Proceedings of the 27th International Conference on Intelligent User Interfaces*, 841–852.