

Retell, Reward, Repeat: Reinforcement Learning for Narrative Theory-Informed LLM Story Retelling

David Y. Liu Xanthe Muston Dipankar Srirag
Aditya Joshi Paul Dawson Sebastian Sequoiah-Grayson
University of New South Wales, Sydney, Australia
Corresponding Author: david.liu2@unsw.edu.au

Abstract

Counterfactual story retelling exposes LLM shortcomings in constrained narrative solution spaces where they can no longer rely on recalling memorised training data. Ground-truth-based post-training, such as SFT, fails to teach LLMs how to generate logical and rational narrative events. In this paper, we introduce Retell, Reward, Repeat (RRR), an RL-based pipeline synthesising Structuralist Narratology with scalar narrativity to teach storytelling structure. We extend an existing counterfactual stories dataset with human-annotated stages of narrative equilibrium. By using d-RLAIF, RRR derives training signals from the narrativity of textual features without the need for reference outputs. Evaluations demonstrate that RRR-trained LLMs outperform few-shot and SFT baselines in logic, rationality, and completeness, with output quality additionally validated by blind human preference. Training on a small, query-only dataset, RRR provides a linguistically grounded, cost-effective post-training mechanism for storytelling—a domain currently lacking effective post-training methods. RRR highlights the continued relevance of integrating established linguistic theories into contemporary NLP.

1 Introduction

Storytelling differs from typical Natural Language Processing (NLP) tasks due to the lack of a universally preferable output. While statistical metrics¹ based on reference outputs have been a norm in NLP, they align poorly/negatively with human judgements of Automatic Story Generation (ASG) outputs (Netisopakul and Taoto, 2023).

Existing large language model (LLM) post-training paradigms such as Supervised Fine-tuning (SFT) and Reinforcement Learning from Human

Feedback (RLHF) restrict the diversity of generated stories (Lu et al., 2025; Sui, 2026), where models converge towards clichéd and unnecessary linguistic styles² (Chakrabarty et al., 2025).

TODOROV'S THEORY OF NARRATIVE EQUILIBRIUM
Equilibrium. Disruption. Recognition of disruption.
Attempt to resolve. New equilibrium.

ORIGINAL STORY Alex and Blair were classmates. They secretly hated each other. Alex was angry and started an argument. Then they were enemies.
COUNTERFACTUAL RETELLINGS (hated → liked)
Alex and Blair were classmates. They secretly liked each other. Their class had 20 students. NARRATIVITY: 3
Alex and Blair were classmates. They secretly liked each other. Nothing ever happened because it was a secret. NARRATIVITY: 4
Alex and Blair were classmates. They secretly liked each other. Alex couldn't wait anymore and asked Blair out. Then they were roommates. NARRATIVITY: 5

Figure 1: Todorov (1971)’s narrative theory applied to Qin et al. (2019)’s counterfactual storytelling task. Narrativity measures the presence of narrative structure.

We move towards a more narrative-aligned LLM post-training starts by redefining how we measure storytelling abilities. When we formulate the problem of ASG as selecting and representing a sequence of events that (A) can be told as a story and (B) make a good story (Li et al., 2013), we can see that current frameworks (Chhun et al., 2022; Chakrabarty et al., 2024) disproportionately prioritise (B). This drives an impossible search for ‘universally good storytelling’, as evidenced by the persistent absence of standardised benchmarks (Liu et al., 2026) and conflicting subjective human evaluations of ASG outputs (Marco et al., 2025).

¹e.g., BLEU (Papineni et al., 2002), ROUGE (Lin, 2004) and BERTSCORE (Zhang et al., 2020).

²ASG literature defaults to reductive universal criteria such as complexity, suspense, aesthetics and entertainment: but soothing bedtime stories are not complex or suspenseful; news stories depicting genocide are not aesthetic or entertaining.

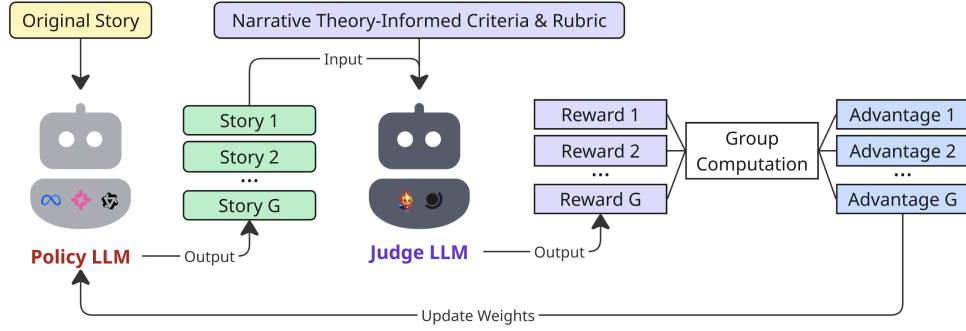


Figure 2: The Retell Reward Repeat (RRR) post-training pipeline uses d-RLAIF (Lee et al., 2024) with a narrative theory-informed LLM-as-judge to generate the reward signal for GRPO (Shao et al., 2024).

Proposing an alternative path, we shift our priority from (B) to (A) and seek to measure the narrativity of a text—measured using fundamental qualities shared by all narratives that provide an agnostic metric across diverse storytelling forms:

RQ1: How can narrativity in LLM outputs be measured with a set of definable properties?

RQ2: How can post-training align LLM outputs with narrativity and other storytelling criteria?

To truly evaluate an LLM’s storytelling abilities, constrained story retelling is a better approach than open-ended generation; while open-ended prompts allow LLMs to simply recall memorised narratives, constraints force them to dynamically sequence logically and rationally connected events—a process that empirically increases task difficulty (Atmakuru et al., 2024). Accordingly, we select Qin et al. (2019)’s counterfactual story retelling task (Figure 1) for this ASG case study, where our contributions are as follows:

- We construct a scheme that synthesises Todorov (1971)’s Structuralist Narratology with Prince (2003)’s scalar concept of narrativity, and additionally define 3 criteria (Logical, Rational, Complete_N) specific to counterfactual story retelling tasks. We show that human-authored stories satisfy the 3 criteria more than LLM-generations, and are also higher in narrativity.
- We release our code and dataset³ of 200 counterfactual retellings with textual features tagged and evaluated by three human annotators, empirically showing the usefulness and limitations of Todorov’s theory. We also identify Gemini-3-Flash (Google, 2025) as the evaluator most aligned with human judgements.

³<https://anonymous.4open.science/r/CounterfactualTodorov>

- We introduce a theory-informed ‘Retell Reward Repeat’ pipeline (RRR; Figure 2) and show that 7B/8B LLMs trained with RRR consistently achieve the highest narrativity and satisfy the 3 retelling criteria whereas SFT models copy surface syntax at the cost of logic and rationality.

Through engaging with debates in literary theory, our study synthesises linguistically grounded narrative frameworks to confront the fundamental challenges of aligning LLM post-training with storytelling abilities. While existing evaluation criteria for storytelling remain valuable, our work demonstrates the necessity of exploring alternative pathways to overcome limitations in current training paradigms.

2 Structuralist Narratology & Narrativity

Todorov (1969) coined the term *narratologie* to describe a systematic study of narrative which seeks to discover the abstract structure and essential textual features that universally and exclusively define a narrative. Structuralist Narratology treats an individual story as just one possible manifestation of narrative’s underlying universal structure.

Todorov (1971)’s plot-focused Theory of Narrative Equilibrium claims that narrative’s underlying universal structure can be represented by sequences of causally and contextually connected transformations, categorisable into 5 types of narrative stages. We illustrate these ‘Todorovian stages’ in Figure 1 where each stage is colour-coded by type.

While Classical Narratology treated a text’s narrativeness as a boolean property, Narratology has increasingly embraced the concept of scalar narrativity (Prince, 2003), where narrativeness is a matter of degree rather than an absolute. Todorov (1971) implied this continuum in his original work—a reader must be able to recognise most types of the

five equilibrium stages for a text to be considered a narrative.

We synthesise these perspectives to operationalise a measurement of narrativity to evaluate the completeness of a counterfactual retelling—shown in Figure 1—where a text containing 5 types of Todorovian stages exhibits higher narrativity than one containing 4 or 3.

By testing Todorov’s assumption that narrativity relies on universal textual features, our case study actively engages with broader debates in literary theory. Post-structuralists like Barthes (1967) argue that interpretation relies on more than just textual features, as narratives exist within fluid cognitive and cultural structures inhabited and contested by ideas. **RQ1** grounds the issue empirically by investigating how *narrativity* itself can be measured from textual features. Ultimately, our annotation process serves to test the boundaries of Todorov’s structuralist framework.

3 Story Annotation

Qin et al. (2019) show that for the task of rewriting the ending of a short story after a provided ‘Counterfactual’⁴ event replaces the ‘Initial’ event (Table 1), an AI-generated output’s similarity⁵ to the human-written ground truths correlates weakly and sometimes negatively with human given ratings of coherence and relevance. To overcome the unreliability of ground truth based evaluation, our case study applies Todorov’s theory to provide a linguistically grounded approach to automatic metrics for training and evaluation.

The functional purpose of the annotation dataset is solely evaluative: i.e., to test agreement between human annotators and LLM-evaluators to inform our decisions in selecting LLM-as-judge models for training and evaluation. The annotated data is not directly used for model training. Additionally, through quantitatively and qualitatively examining the inter-annotator agreement, we can empirically observe whether and how we can measure narrativeness using definable conditions (**RQ1**).

3.1 Curation

We curate a pool of human and AI-generated retellings to reflect the distribution of texts the

LLM-evaluator will see during the reinforcement learning process and final evaluation.

We start the curation process from the first 3000 items in the TimeTravel supervised training split (Qin et al., 2019), where each item (Table 3) provides a ‘query’ consisting of the premise, initial state, original ending, and counterfactual, and a ‘response’ which is the edited ending (human-written ground truth). For each query, we produce three AI-generated responses (edited endings) using Llama-3.1-8B-Instruct (Grattafiori et al., 2024), Olmo-3-7B-Instruct (Ettinger et al., 2025), and Qwen-3-8B (Yang et al., 2025) with thinking disabled.

Because the three modern instruction-tuned LLMs typically produce similar responses to the same query, we introduce a filter aiming to diversify the data. We do not include the human ground truth in the calculation of our diversity metric to avoid biasing our selection towards ‘harder’ queries where the LLM output is distant from the ground truth (even if it is a weak correlation).

For each query’s three AI-generated responses, we use DeBERTa-v3-small (He et al., 2023) to compute all pairwise distances (defined as $1 - \text{BERTScore}_{F1}$) and define the diversity metric as $\text{diversity} = \min(\text{distances}) \times \text{mean}(\text{distances})$. We then select the 50 inputs with the highest diversity and, for each, retain four retellings (the three LLM responses plus the ground truth human response), yielding an annotation set of $n = 200$.

3.2 Scheme

We operationalise our scalar measurement of narrativity: the number of recognisable Todorovian stages in a text (Table 1). Since disruption (i.e., change) is essential for narrative transformation, any text lacking a disruption is assigned a score of 1. For all other cases, the narrativity score is defined as $\text{Narrativity} = \min\{s + 1, 5\}$, where s is the number of identifiable Todorovian stages. Our conditions only require 4 recognisable narrative stages to achieve the maximum scores since we do not wish to penalise stories which begin in medias res or end on a cliffhanger (Appendix E.2).

In addition to scalar narrativity, we define specific criteria to measure quality for the counterfactual story retelling task, which we use to investigate **RQ2**. These criteria are designed to align with those originally proposed by Qin et al. (2019), which primarily measure relevance and coherence. Our approach differs by grounding these definitions in Todorov’s narrative theory wherever possible:

⁴The ‘Counterfactual’ is a counterfactual antecedent while the ‘Edited Ending’ is a counterfactual consequent.

⁵Measured using BLEU (Papineni et al., 2002), ROUGE (Lin, 2004) & BERTScore (Zhang et al., 2020).

Narrativity (Equilibrium→Disruption→Recognition→Attempt→New Equilibrium)

A: (4) Evan was ready to buy himself a bicycle. He went to the store but forgot his money. Then he looked around some more, trying to find something else he could afford or a way to get home.

B: (4) Evan was ready to buy himself a bicycle. He went to the store but forgot his money. Then he looked around some more, trying to find something else he could afford or a way to get home.

C: (3) Evan was ready to buy himself a bicycle. He went to the store but forgot his money. Then he looked around some more, trying to find something else he could afford or a way to get home.

Is the story logical?

A: (3) I interpret ‘but forgot his money’ as Evan forgetting a large sum of money intended for buying a bicycle. Evan may still have the usual money he carries everyday, and is able to logically find something else he could afford.

B: (3) The edited text is logical. Forgetting his money and then looking for an alternative or a way home is internally consistent.

C: (1) There is a logical impossibility in the story because: 1. Evan forgets his money, but 2. proceeds to look around the store for something he can “afford,” which would be nothing.

Table 1: An example of 3 human annotators (A/B/C) applying our framework. More examples are in Appendix J.

Logical: The text is free from inconceivable scenario(s) that contradict the text’s internal logic.

Rational: The text can be rationally interpreted using Todorovian stage(s) in its entirety without contextual/causal disconnection.

Complete_N: The text includes all key Todorvian stage(s) from the original that ensure the text is as narratively complete as the original.

min_{LRC}: Since any desirable story needs to satisfy all 3 of the aforementioned criteria, we use min(Logical, Rational, Complete_N) to represent a text’s overall quality instead of the average. This reflects the principle that failing even one criterion makes the story completely undesirable.

3.3 Human Annotation

Annotator	Logical	Rational	Complete _N	Narrativity
A	0.31	0.53	0.46	0.43
B	0.24	0.37	0.38	0.53
C	0.17	0.46	0.37	0.49
Average	0.24	0.46	0.40	0.48

Table 2: Correlation (Spearman’s ρ) of each annotator vs the median of the other two ($n = 200$). All $p < 0.05$.

Three human annotators, blind to whether the stories are human- or LLM-authored, annotate the texts using a tag-and-evaluate process: (1) Tag parts of text with Todorovian stages. The variety of tags is automatically counted and computed into a 1-5 Narrativity score, then (2) use 3-point Likert scales to rate the Logical/Rational/Complete_N criteria. (Additional instructions and examples given to the annotators are included in Appendix E.)

After acquiring some additional rationales from the annotators, a qualitative review of the disagreements (Table 1) shows that annotators often formed differing yet valid interpretations of the same text.

Readers frequently disagreed on structural starting points—such as whether a narrative begins in a state of equilibrium or in disruption—and applied subjective assumptions to resolve ambiguities.

Story quality and narrativity arise from subjective interactions between readers and texts (Marco et al., 2025). This inherent subjectivity, consistent with reader-centric post-structuralist literary theory (Barthes, 1967), accounts for the fair to moderate agreement observed from annotators (Table 2). Textual features capture only one dimension of narrative understanding, which is ultimately completed by the reader’s cognitive processing.

For context, Chhun et al. (2022) previously used ICC2k to report agreement between 3 human annotators with an average of 0.40 across 6 criteria: relevance (0.48), coherence (0.29), empathy (0.34), surprise (0.28), engagement (0.46) and complexity (0.56). In comparison, the ICC2k of our criteria—Narrativity (0.64), Logical (0.36), Rational (0.67), complete_N (0.58) and min_{LRC} (0.56)—align with the upper range of expected inter-annotator agreement when subjectively evaluating narratives only based on their textual features.

Our evaluation scheme reliably differentiates between ground truths and different LLM generations. When ranking the authors (Human/Olmo/Qwen/Llama) by their average received ratings, the blind evaluations consistently yield the same rankings of authors across all four criteria: Logical, Rational, Complete_N, and Narrativity (see Appendix Table 10). Human ground truths are consistently rated higher in both narrativity and overall quality than LLM generations. Ultimately, these annotations validate our framework, proving that our criteria establish a meaningful, non-trivial benchmark for the counterfactual story retelling task (Table 3).

(query) Original Story + Counterfactual						
Premise	Kimmy was on a trip to Paris.					
Initial	Her plane had just landed and she and her classmates her excited.					
Original Ending	They couldn't wait to start seeing museums and churches! They spent lots of time shopping and seeing historical sites. Kimmy was very excited.					
Counterfactual	Kimmy's plane crashed over the ocean and killed everyone on board.					
(response) Premise → Initial → Counterfactual → Edited Ending						
Source/Model	Counterfactual Story	L	R	C	min _{LRC}	Our Comment
Ground Truth	Kimmy was on a trip to Paris. Kimmy's plane crashed over the ocean and killed everyone on board. Kimmy had just wanted to see the museums and churches! Her parents spent a lot of time planning her funeral. Kimmy was dead.	3	3	3	3	The scenarios are logical, the sequence is rationally connected, and the story feels complete.
Zero-Shot	Kimmy was on a trip to Paris. Kimmy's plane crashed over the ocean and killed everyone on board. They never made it to Paris. Kimmy's plane crashed over the ocean and there were no survivors.	3	1	3	1	The crash is described twice. Irrationally repetitive.
Few-Shot	Kimmy was on a trip to Paris. Kimmy's plane crashed over the ocean and killed everyone on board. Kimmy's plane crashed over the ocean and killed her.	1	3	1	1	The scenario of Kimmy's plane crashing and killing her again is illogical. The story does not feel complete either.
SFT	Kimmy was on a trip to Paris. Kimmy's plane crashed over the ocean and killed everyone on board. They couldn't wait to start seeing museums and churches! They spent lots of time shopping and seeing historical sites. Kimmy was very excited	1	1	3	1	The contradiction of the plane crash and then the Paris visit is illogical. There is no rational causal/contextual connection between the events.
RRR _N	Kimmy was on a trip to Paris. Kimmy's plane crashed over the ocean and killed everyone on board. The plane wreckage was found at the bottom of the ocean, a tragic reminder of the devastating crash that took Kimmy and her classmates' lives.	3	3	3	3	The scenarios are logical, the sequence is rationally connected, and the story feels complete.

Table 3: Example test outputs from models based on Llama-3.1-8B-Instruct. Ratings are given by LLM-as-judge (Gemini-3-Flash): L=Logical, R=Rational, C=Complete_N, M=min_{LRC}. Additional examples in Appendix K.

3.4 LLM-as-Judge

Model	Logical	Rational	Completeness _N	Narrativity
Selene	−0.05 [†]	0.31	0.23	0.39
Prometheus	0.24	0.29	0.33	0.32
Gemini	0.34	0.39	0.52	0.50

Table 4: Correlation (Spearman’s ρ) of LLM-as-Judge vs the median of 3 human annotators ($n = 200$). All $p < 0.05$ except Selene Logical ($^{\dagger}p = 0.47$).

We test 3 LLM-as-judge evaluators: Selene-1-mini-8B (Alexandru et al., 2025), M-Prometheus-14B (Pombal et al., 2025) and Gemini-3-Flash (Google, 2025). The first two are open-weight models that specialise in evaluation tasks, while Gemini-3-Flash is chosen as a state-of-the-art (SOTA) proprietary model. We compute Spearman’s ρ correlations between each LLM-judge and the median of three human annotators (Table 4).

Gemini achieves an average correlation of 0.44,

which is comparable to the average human annotator—the correlation between an average human annotator and the median of the remaining two is 0.40 (Table 2). In contrast, open-weight models achieve lower correlations, 0.25 for Selene and 0.30 for Prometheus. These results show that Gemini serves as a good proxy for human judgment. We select Gemini for our final evaluation as it is the most reliable and human-aligned.

3.4.1 Comparison with Previous Metrics

Qin et al. (2019)’s dataset and task are evaluated in the existing literature (Liu et al., 2023, 2025) using automatic-metrics—BLEU-4, ROUGE-L and BERTscore—out of which only BERTscore has a positive correlation with all the original 3 human-annotated criteria. In comparison, our chosen Gemini generated automatic metric minLRC and the 3 human rated criteria have an average pearson’s correlation of 0.30 (vs BERTscore’s 0.17). Therefore, our automatic evaluation is an improvement.

4 Retell, Reward, Repeat (d-RLAIF)

As shown in Figure 2, we adapt Wei et al. (2025)’s implementation of d-RLAIF to use an LLM-as-judge to transform an input story into a theory-informed reward signal. We use GRPO (Shao et al., 2024) and LoRA (Hu et al., 2022) to optimise the RRR training process using TimeTravel’s unsupervised training split. To study the effects of our RRR training on a wide range of models, 3 instruction-tuned models (Llama-3.1-8B, Qwen-3-8B, and Olmo-3-7B) are chosen as policy models for generating retellings.

The policy model generates 16 responses for each query, which are then given to the LLM-as-judge to evaluate and generate 16 reward scores. The relative advantages are then calculated from the reward scores and used to update the weights (Figure 2) of the policy model’s LoRA adapter.

Based on the results in Table 4, we select Selene to generate the narrativity-based reward score R_N (Figure 3). For comparison, we select Prometheus to compute an alternative reward signal, based on the Logical, Rational and Complete_N criteria. To maximise speed and token efficiency during d-RLAIF, we generate R_O (‘overall’) using a prompt with criteria identical to the $\min_{LRC}$ score. Although this configuration aligns less closely with human preferences than the three-inference aggregated $\min_{LRC}$ score (Appendix Table 9), we consider the minor performance trade-off acceptable given the $\frac{1}{3}$ reduction in token costs.

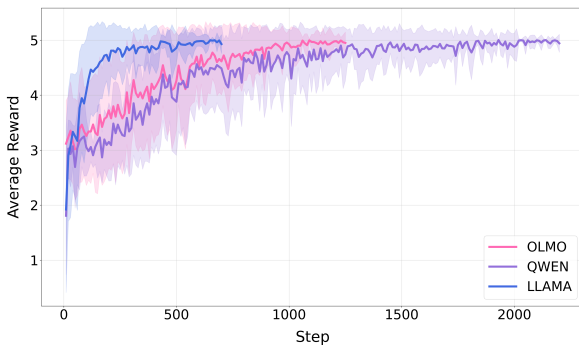


Figure 3: Mean and standard deviation of the reward signal R_N generated by Selene-1-mini during d-RLAIF.

We apply early stopping once training reaches a reward plateau, defined as 200 consecutive optimiser steps for which the average reward remains within 0.1 of the maximum achievable reward, indicating no further measurable improvement. Additionally, we introduce a length-based penalty for

retellings that exceed 3 times the length of the original to prevent reward hacking by lengthy outputs.

We detail additional ablation experiments on the reward design in Appendix H.

5 Evaluation

We evaluate our models on TimeTravel’s test split ($n=1871$) using the most aligned LLM-as-judge we found in Section 3.4: Gemini-3-Flash. Each item in the test split contains 3 slightly different human-written edited endings—we evaluate each and use their average for the results in Table 5. We also use BLEU-4 and ROUGE-L to calculate the average similarity of each model’s generation to the nearest-of-3 human-written endings to measure similarity.

Human authored ground truths receive higher average ratings than LLMs when considering $\min_{LRC}$, which more better represents overall quality than the average, as all desirable stories must score perfectly across the 3 criteria.

For baseline comparisons against our method, we compare zero-shot and few-shot for each LLM. We also evaluate the outputs produced by the Program-Prompt method (Liu et al., 2023). Per base model, few-shot does not result in consistent or noticeable performance over zero-shot.

For training comparison, we conduct standard SFT with LoRA using TimeTravel’s supervised training split, also with early stopping at convergence—defined by when loss decreases < 0.01 for 3 consecutive training steps (Appendix Figure 7). SFT achieves the highest Complete_N, BLEU-4 and ROUGE-L for all models. SFT results in an increase of $\min_{LRC}$ performance for Qwen and Llama but degrades the performance of the best-performing base model Olmo. Stories generated by the SFT models are highest in similarity to the human-written ground-truths, as well as the most structurally complete when compared to the corresponding original story (measured by Complete_N).

When using RRR, we find that d-RLAIF with M-Prometheus’s R_O improves the overall performance of all tested models, where Selene-1-mini’s narrativity-based reward signal R_N reliably results in the best $\min_{LRC}$ performance, second only to humans (significance testing in Appendix I). Their outputs also contain the highest narrativity and dissimilarity to human ground truths.

Model	Method	Logical	Rational	Completeness	$\min_{\text{LRC}}$	Narrativity	BLEU-4	ROUGE-L
GROUND TRUTHS	Crowd Workers	2.884	2.927	2.933	2.811	4.718	-	-
LLAMA-3.1-8B-INST.	BASE	2.735	2.385	2.528	2.080	4.347	0.340	0.502
	few-shot	2.768	2.615	2.505	2.213	4.303	0.306	0.473
	SFT	2.320	2.501	2.934	2.260	4.692	0.810	0.835
	RRR _O -Prometheus	2.801	2.825	2.594	2.400	4.534	0.179	0.386
	RRR _N -Selene	2.785	2.860	2.740	2.554	4.733	0.087	0.299
QWEN-3-8B	BASE (non-thinking)	2.429	2.132	2.759	1.940	4.551	0.572	0.667
	few-shot	2.543	2.219	2.686	2.006	4.559	0.492	0.607
	SFT	2.402	2.532	2.925	2.317	4.693	0.790	0.823
	RRR _O -Prometheus	2.717	2.336	2.566	2.024	4.356	0.337	0.491
	RRR _N -Selene	2.781	2.890	2.705	2.516	4.894	0.003	0.192
OLMO-3-7B-INST.	BASE	2.762	2.812	2.400	2.222	4.328	0.097	0.315
	few-shot	2.745	2.837	2.377	2.208	4.415	0.072	0.287
	SFT	2.295	2.441	2.910	2.211	4.647	0.787	0.821
	RRR _O -Prometheus	2.770	2.857	2.491	2.325	4.460	0.090	0.309
	RRR _N -Selene	2.793	2.893	2.595	2.434	4.899	0.002	0.178
CODEx	Program Prompt	2.452	2.484	2.758	2.181	4.514	0.617	0.695

Table 5: Evaluation of post-trained LLMs using the test split (n=1871) from the TimeTravel dataset using Gemini-3-Flash. For each base LLM, the highest score in each column is bolded. Significance testing in Appendix I.

6 Discussion

6.1 $\min_{\text{LRC}}$ and Human Preference

	Llama (SFT)	Llama (RRR _N)
Narrativity	4.820	4.580
$\min_{\text{LRC}}$	2.220	2.480
Win-Rate	28%	72%

Table 6: Human preference test on a subset of 100.

As a sanity check to see whether our primary quality metric $\min_{\text{LRC}}$ correlates with human preference, we instruct an additional annotator D (who is not involved in the initial annotations) to choose their preferred story on a subset of the first 50 queries and 100 shuffled responses, while being blind to which model generated it (Table 6).

Out of 100 stories generated by Llama-SFT and Llama-RRR_N, a blinded annotator D strongly prefers Llama-RRR_N at 72% of the time (Table 6). When looking at qualitative outputs (Table 3), it is also self-evident that stories which are logical, rational, and narratively complete are more desirable than stories which are not. This shows that, for these micro-stories, the Gemini-generated $\min_{\text{LRC}}$ is a valid measurement of story quality which positively correlates with human preference.

6.2 Reinforcement Learning vs SFT

Intuitively, it does not make sense to evaluate a student’s storytelling ability by comparing their stories with the teacher’s. When we consider that

humans naturally learn storytelling as a communicative practice through reinforcement learning (Hineline, 2018), our RRR pipeline using d-RLAIF is an intuitive choice for Automatic Story Generation (ASG).

The increases in $\min_{\text{LRC}}$ and Narrativity across all models trained using our RRR_N method (Table 5) indicate that Todorov’s abstract and theoretical narrative structure can be learned through reinforcement learning (since Narrativity is measured by the presence of Todorovian structures) and is beneficial in generating logical, rational, and structurally complete stories even though they differ the most to human ground truths (Table 5).

While SFT models achieve the highest Completeness, the highest similarity to ground truths, and the second highest Narrativity, this does not translate to higher $\min_{\text{LRC}}$ scores or human preference. Example outputs (Table 3) show that SFT models often generate high-narrativity story endings by illogically and irrationally duplicating the original ending. It appears that SFT forces LLMs to over-fit to the training data without learning functional narrative abilities for the counterfactual story retelling task.

The results also highlight that, while optimising for narrativity during training (RRR_N) succeeds in our case study, evaluating based on narrativity alone is insufficient. Recalling our formulation of ASG—selecting and representing a sequence of events that (A) form a story and (B) make a good story—both components are essential. While our findings demonstrate the value of prioritising

(A) during training, evaluation must still assess the criteria (B) required by the specific task.

6.3 d-RLAIF for Storytelling

The lack of consistent or significant performance increases of few-shot over zero-shot (Table 5) shows that, for counterfactual story retelling, post-training is necessary for improving storytelling quality. Since our RRR pipeline uses less than a single epoch from the dataset to reach convergence, it suggests that massive storytelling datasets may not be necessary for post-training. Efforts to curate datasets for d-RLAIF may instead focus on comparatively smaller and query-only datasets. This approach potentially circumvents the bottleneck of limited shared datasets (Liu et al., 2026).

Using d-RLAIF, methods of effective narrative understanding can provide great assistance in providing the reward signals needed for the training of future ASG models. We believe it to be beneficial to move towards unified knowledge and methods between generation and understanding.

Although we apply RRR solely to the counterfactual story retelling task, the method is highly adaptable. By prompting an LLM-as-judge within the d-RLAIF pipeline, the training process allows the student policy model to learn behaviours compatible with the teacher model’s conceptual representations. For instance, models can be trained to adhere to stylistic genre conventions or generate empathetic narratives. This process requires minimal training data, provided the LLM-as-judge possesses a robust representation of the target behavior. In this way, the pipeline itself could potentially as a postclassical representation of narrative that models the interaction between text and the reader (Piper et al., 2021).

6.4 Implications to Narrative Theory

While we use narrative theory as a foundation for storytelling, the fair-to-moderate agreement empirically highlights that the inherent subjectivity of storytelling may be incompatible with universal benchmarks, and that Todorov’s structural theory alone cannot model all narrative qualities (e.g., the ‘Logical’ and ‘Rational’ criteria also draw on cognitive considerations).

7 Related Work

Qin et al. (2019)’s TimeTravel dataset and its counterfactual story retelling task have been approached

by a number of past works, experimenting with: inference-time optimisations (Qin et al., 2020; Chen et al., 2022), multi-step generation (Hao et al., 2021), Program-prompting (Liu et al., 2023) and ConceptNet (Ashwani et al., 2024). These past works position storytelling as a task that tests the ability of a LLM to reason and understand causal logic and do not focus on developing a training method which aligns with storytelling principles.

Wei et al. (2025) showed that d-RLAIF can perform well at improving creative-writing for generating Chinese greetings, a tasks which involves open-ended creative writing.

8 Conclusion

We address the persistent limitations of LLM post-training in Automatic Story Generation through a narrative theory-informed approach. First, we demonstrated that narrativeness can be measured through definable (although not universal/objective) properties—synthesising Todorov’s Theory of Narrative Equilibrium with Prince’s scalar narrativity into a quantifiable metric based on structural stages (RQ1). Building on this, we introduced “Retell Reward Repeat” (RRR) to align LLMs with these criteria (RQ2) using direct Reinforcement Learning from AI Feedback (d-RLAIF).

Our results across Llama-3.1, Qwen-3, and Olmo-3 show that models trained with our narrativity reward (R_N) significantly outperform SFT and few-shot baselines in output story quality, as measured in the counterfactual story retelling task using criteria—Logical, Rational and Complete_N—which we show to correlate with human preference. Our findings establish d-RLAIF as a highly viable post-training paradigm for story generation.

Our case study demonstrates both the benefits and limitations of applying Todorov (1969)’s narrative theory to model storytelling, highlighting the need to move beyond textual features. Postclassical narratology expands beyond the limitations of structuralist narratology by emphasising the historical, cultural, and cognitive parameters of storytelling and the reader’s role in constructing narratives from linguistic artifacts (Meister, 2011).

Limitations

To ensure correct understanding and application of narrative theory, the authors served as annotators in this case study, which introduces potential bias. While a large pool of external annotators is optimal

practice, we believe this limitation does not undermine our core findings. Among the annotators: only A was involved in modelling and collecting examples. B, C and D were involved in evaluating the outputs, and drafting the paper.

Our research demonstrates that human evaluation of stories is highly subjective. Supported by narrative theory (Barthes, 1967), existing empirical findings (Marco et al., 2025), and our own results, we argue that annotator disagreement is a feature, rather than a bug, of storytelling tasks. Specifically, our data shows: (1) There are inherent disagreements between human annotators regarding story understanding. (2) Story understanding is a fundamentally subjective process. (3) Blinded human annotators consistently rate human-written ground truths higher than LLM generations.

For an expanded annotator pool to contradict our claims, new data would have to demonstrate that story understanding is objective (showing little to no disagreement) or that humans consistently rate LLM generations higher than human-written ground truths. Given existing evidence across multiple disciplines, we consider this unlikely.

Apart from Liu et al. (2023)’s program-prompt, we do not compare our models with those developed by others using the TimeTravel dataset (Qin et al., 2020; Chen et al., 2022; Hao et al., 2021; Liu et al., 2023; Ashwani et al., 2024). This is primarily because none of them are reasonable baselines for our experiments, particularly in the context of post-training with narrative-based metrics (the others primarily optimise for statistical metrics based on reference outputs).

To avoid the massive costs, we did not compare the final accuracy of Gemini-3-Flash’s evaluations with human judgements on the whole test set (we only compare on a small subset of $n = 100$). Our reliance on LLM-as-judge to approximate human judgement is an assumption and limitation.

As a novel case study for reinforcement learning for story generation, we do not test how our approach scales to longer stories. We only focused on training 7/8B models, the effectiveness of our method for smaller/larger LLMs is untested. These can be explored in future studies.

Ethics Statement

We use a benchmark dataset (Qin et al., 2019), and do not have any additional ethical considerations to report. Our human evaluators are authors on the

paper. GitHub Copilot and Google AI studio were used to assist with coding tasks and debugging. Microsoft Copilot was used for formatting the latex in the manuscript. Outputs generated by these tools were carefully reviewed and validated by the authors to maintain accuracy and correctness.

Our dataset, tasks and generated stories are in English and refer to Western contexts. We use a European narrative theory. It is well known in the narrative studies community that cultural differences can affect how we perceive narratives and their underlying themes and structures (Aziz, 2023; Phillips et al., 2025). As such, our paper’s representation of storytelling is biased and does not reflect all cultures and is not universal.

Storytelling is a highly effective communication tool that can be used for influence and manipulation. The development of Automatic Story Generation methods, with the goal of matching human storytelling, introduces the risk of unintended or even malicious harm.

Acknowledgments

Masked for blind review.

References

- Andrei Alexandru, Antonia Calvi, Henry Broomfield, Jackson Golden, Kyle Dai, Mathias Leys, Maurice Burger, Max Bartolo, Roman Engeler, Sashank Pisupati, and 1 others. 2025. Atla selene mini: A general purpose evaluation model. *arXiv preprint arXiv:2501.17195*.
- Swagata Ashwani, Kshiteesh Hegde, Nishith Reddy Mannuru, Dushyant Singh Sengar, Mayank Jindal, Krishna Chaitanya Rao Kathala, Dishant Banga, Vinija Jain, and Aman Chadha. 2024. Cause and effect: can large language models truly understand causality? In *Proceedings of the AAAI Symposium Series*, volume 4, pages 2–9.
- Anirudh Atmakuru, Jatin Nainani, Rohith Sidhartha Reddy Bheemreddy, Anirudh Lakkaraju, Zonghai Yao, Hamed Zamani, and Haw-Shiuan Chang. 2024. Cs4: Measuring the creativity of large language models automatically by controlling the number of story-writing constraints. *arXiv preprint arXiv:2410.04197*.
- Sardar Khawar Aziz. 2023. Cross-cultural narratology: A comparative study of storytelling techniques in eastern and western literature. *Journal of Asian Development Studies*, 12(4):742–753.
- Roland Barthes. 1967. The death of the author. *Aspen*, (5–6).

- Tuhin Chakrabarty, Philippe Laban, Divyansh Agarwal, Smaranda Muresan, and Chien-Sheng Wu. 2024. Art or artifice? large language models and the false promise of creativity. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, pages 1–34.
- Tuhin Chakrabarty, Philippe Laban, and Chien-Sheng Wu. 2025. Can ai writing be salvaged? mitigating idiosyncrasies and improving human-ai alignment in the writing process through edits. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, pages 1–33.
- Jiangjie Chen, Chun Gan, Sijie Cheng, Hao Zhou, Yanghua Xiao, and Lei Li. 2022. Unsupervised editing for counterfactual stories. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 36, pages 10473–10481.
- Cyril Chhun, Pierre Colombo, Fabian Suchanek, and Chloé Clavel. 2022. Of human criteria and automatic metrics: A benchmark of the evaluation of story generation. In *Proceedings of the 29th International Conference on Computational Linguistics*, pages 5794–5836.
- Allyson Ettinger, Amanda Bertsch, Bailey Kuehl, David Graham, David Heineman, Dirk Groeneveld, Faeze Brahman, Finbarr Timbers, Hamish Ivison, and 1 others. 2025. Olmo 3. *arXiv preprint arXiv:2512.13961*.
- Google. 2025. Introducing Gemini 3 Flash: Benchmarks, Global Availability. <https://blog.google/products/gemini/gemini-3-flash/>. Accessed 2025-12-27.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, and 1 others. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.
- Kilem L Gwet. 2014. *Handbook of inter-rater reliability: The definitive guide to measuring the extent of agreement among raters*. Advanced Analytics, LLC.
- Changying Hao, Liang Pang, Yanyan Lan, Yan Wang, Jiafeng Guo, and Xueqi Cheng. 2021. Sketch and customize: A counterfactual story generator. In *Proceedings of the AAAI conference on artificial intelligence*, volume 35, pages 12955–12962.
- Pengcheng He, Jianfeng Gao, and Weizhu Chen. 2023. *Debertav3: Improving deberta using electra-style pre-training with gradient-disentangled embedding sharing*. In *International Conference on Learning Representations (ICLR)*.
- Philip N Hine. 2018. Narrative: Why it’s important, and how it works. *Perspectives on Behavior Science*, 41(2):471–501.
- Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, Weizhu Chen, and 1 others. 2022. Lora: Low-rank adaptation of large language models. *ICLR*, 1(2):3.
- Yash Kumar Lal, Vanya Cohen, Nathanael Chambers, Niranjan Balasubramanian, and Ray Mooney. 2024. Cat-bench: Benchmarking language model understanding of causal and temporal dependencies in plans. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 19336–19354.
- Andrew Lampinen, Ishita Dasgupta, Stephanie Chan, Kory Mathewson, Mh Tessler, Antonia Creswell, James McClelland, Jane Wang, and Felix Hill. 2022. *Can language models learn from explanations in context?* In *Findings of the Association for Computational Linguistics: EMNLP 2022*, pages 537–563, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Harrison Lee, Samrat Phatale, Hassan Mansoor, Thomas Mesnard, Johan Ferret, Kellie Lu, Colton Bishop, Ethan Hall, Victor Carbune, Abhinav Rastogi, and 1 others. 2024. Rlaif vs. rlhf: scaling reinforcement learning from human feedback with ai feedback. In *Proceedings of the 41st International Conference on Machine Learning*, pages 26874–26901.
- Boyang Li, Stephen Lee-Urban, George Johnston, and Mark Riedl. 2013. Story generation with crowd-sourced plot graphs. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 27, pages 598–604.
- Chin-Yew Lin. 2004. *Rouge: A package for automatic evaluation of summaries*. In *Text Summarization Branches Out (Workshop at ACL)*, pages 74–81, Barcelona, Spain. Association for Computational Linguistics.
- David Y Liu, Aditya Joshi, and Paul Dawson. 2026. Narrative theory-driven llm methods for automatic story generation and understanding: A survey. *arXiv preprint arXiv:2602.15851*.
- Xiao Liu, Da Yin, Chen Zhang, Yansong Feng, and Dongyan Zhao. 2023. *The magic of IF: Investigating causal reasoning abilities in large language models of code*. In *Findings of the Association for Computational Linguistics: ACL 2023*, pages 9009–9022, Toronto, Canada. Association for Computational Linguistics.
- Xiao Liu, Da Yin, Chen Zhang, Dongyan Zhao, and Yansong Feng. 2025. Eliciting and improving the causal reasoning abilities of large language models with conditional statements. *Computational Linguistics*, 51:467–504.
- Ximing Lu, Melanie Sclar, Skyler Hallinan, Niloofar Mireshghallah, Jiacheng Liu, Seungju Han, Allyson Ettinger, Liwei Jiang, Khyathi Chandu, Nouha Dziri, and Yejin Choi. 2025. *AI as humanity’s salieri*:

- Quantifying linguistic creativity of language models via systematic attribution of machine text against web text. In *The Thirteenth International Conference on Learning Representations*.
- Guillermo Marco, Julio Gonzalo, and Víctor Fresno. 2025. The reader is the metric: How textual features and reader profiles explain conflicting evaluations of ai creative writing. In *Findings of the Association for Computational Linguistics: ACL 2025*, pages 25432–25449.
- Jan Christoph Meister. 2011. *Narratology*. The Living Handbook of Narratology. Rev. 19 January 2014.
- Nasrin Mostafazadeh, Nathanael Chambers, Xiaodong He, Devi Parikh, Dhruv Batra, Lucy Vanderwende, Pushmeet Kohli, and James Allen. 2016. *A corpus and cloze evaluation for deeper understanding of commonsense stories*. In *Proceedings of the 2016 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 839–849, San Diego, California. Association for Computational Linguistics.
- Ponrudee Netisopakul and Usanisa Taoto. 2023. Comparison of evaluation metrics for short story generation. *IEEE access*, 11:140253–140269.
- Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. 2002. *Bleu: A method for automatic evaluation of machine translation*. In *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics*, pages 311–318, Philadelphia, Pennsylvania, USA. Association for Computational Linguistics.
- Louise H Phillips, Louisa Lawrie, Zuzana Suchomelova, Sara Heinämaa, Amy O’Dwyer, and Min Hooi Yong. 2025. Age and cultural differences in the relationship between reading and theory of mind. *Poetics*, 109:101984.
- Andrew Piper, Richard Jean So, and David Bamman. 2021. *Narrative theory for computational narrative understanding*. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 298–311, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- José Pombal, Dongkeun Yoon, Patrick Fernandes, Ian Wu, Seungone Kim, Ricardo Rei, Graham Neubig, and André FT Martins. 2025. M-prometheus: A suite of open multilingual llm judges. *arXiv preprint arXiv:2504.04953*.
- Gerald Prince. 2003. *A dictionary of narratology*. U of Nebraska Press.
- Lianhui Qin, Antoine Bosselut, Ari Holtzman, Chandra Bhagavatula, Elizabeth Clark, and Yejin Choi. 2019. *Counterfactual story reasoning and generation*. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 5043–5053, Hong Kong, China. Association for Computational Linguistics.
- Lianhui Qin, Vered Shwartz, Peter West, Chandra Bhagavatula, Jena D. Hwang, Ronan Le Bras, Antoine Bosselut, and Yejin Choi. 2020. *Back to the future: Unsupervised backprop-based decoding for counterfactual and abductive commonsense reasoning*. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 794–805, Online. Association for Computational Linguistics.
- Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, YK Li, Yang Wu, and 1 others. 2024. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300*.
- Peiqi Sui. 2026. Llms exhibit significantly lower uncertainty in creative writing than professional writers. *arXiv preprint arXiv:2602.16162*.
- T. Todorov. 1969. *Grammaire du Décaméron*. Approaches to semiotics. Mouton.
- Tzvetan Todorov. 1971. The 2 principles of narrative. *diacritics*, pages 37–44.
- Xiaolong Wei, Bo Lu, Xingyu Zhang, Zhejun Zhao, Dongdong Shen, Long Xia, and Dawei Yin. 2025. *Igniting creative writing in small language models: LLM-as-a-judge versus multi-agent refined rewards*. In *Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing*, pages 17171–17197, Suzhou, China. Association for Computational Linguistics.
- An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, and 1 others. 2025. Qwen3 technical report. *arXiv preprint arXiv:2505.09388*.
- Slavica Zec, Nicola Soriani, Rosanna Comoretto, and Ileana Baldi. 2017. High agreement and high prevalence: the paradox of cohen’s kappa. *The open nursing journal*, 11:211.
- Tianyi Zhang, Varsha Kishore, Felix Wu, Kilian Q. Weinberger, and Yoav Artzi. 2020. *Bertscore: Evaluating text generation with bert*. In *International Conference on Learning Representations (ICLR) 2020*. OpenReview.net.

A License

The TimeTravel dataset is used under the MIT License. Hugging face’s Transformers and TRL libraries are used under the Apache License 2.0 license. Llama 3.1 is used under the LLaMA 3.1

Community License. Olmo 3, Qwen 3 and Selene-1-mini are used under the Apache License 2.0.

We use these artifacts in a way that’s consistent with intended use for the purpose of academic research.

B TimeTravel Dataset

The TimeTravel dataset is constructed from the ROCStories corpus (Mostafazadeh et al., 2016) which consist of over 100k human written five-sentence short stories. In TimeTravel, there are 1.87k items in the test set, 1.87k items in the validation set, and 16.8k items in the supervised training set, all with human written counterfactual endings. Furthermore, the entire 100k ROC stories can be used for unsupervised training.

The dataset includes narratives that may allude to violent themes, and some generated outputs could be distressing. To the best of our knowledge, these stories are entirely fictional and do not represent real individuals.

C Computational Budget

Inference and d-RLAIF and SFT training was conducted on a single H200 GPU. The duration of each training process was between 30 minutes and 4 hours.

Evaluation using Gemini-3-Flash was done using OpenRouter API, which cost roughly \$3 per testset ($n = 1871$).

D Hyperparameters

We detail the hyperparameters applied for supervised fine-tuning (SFT) and GRPO-based deep reinforcement learning from AI feedback (d-RLAIF). All training and inference were performed using Hugging Face’s Transformers and TRL libraries. The same hyperparameter configuration was maintained across all experimental settings. Early stopping was implemented for both SFT and d-RLAIF.

D.1 Supervised Fine-Tuning (SFT)

For SFT, we train the model using the AdamW optimiser with a fixed learning rate and Low-Rank Adaptation (LoRA) parameterisation (Hu et al., 2022). Training is conducted for a single fine-tuning stage, with evaluation performed periodically on a held-out validation set. The best model checkpoint is selected based on validation loss using early stopping.

LoRA-specific hyperparameters are:

- Rank $r = 64$
- Alpha $\alpha = 128$
- Dropout = 0.05

Other training hyperparameters are:

- Per-device batch size = 8
- Gradient accumulation steps = 8
- Learning rate = 1×10^{-4}
- Number of epochs = 1
- Maximum sequence length = 640 tokens

Mixed-precision training is enabled using bfloat16 (bf16). Gradient check pointing is used to reduce memory usage. No extensive hyperparameter tuning was performed; values were chosen heuristically.

D.2 d-RLAIF with GRPO

For d-RLAIF, we employ GRPO with LoRA parameterisation for the policy model. At each training step, 16 candidate completions are generated per prompt, and early stopping is used based on evaluation metrics.

LoRA-specific hyperparameters for d-RLAIF:

- Rank $r = 64$
- Alpha $\alpha = 128$
- Dropout = 0.05

Other training hyperparameters are:

- Per-device batch size = 24
- Gradient accumulation steps = 2
- Learning rate = 5×10^{-6}
- Number of epochs = 1
- Maximum prompt and completion length = 512 tokens

Mixed-precision training using bfloat16 (bf16) was enabled. No extensive hyperparameter search was performed; hyperparameters were selected heuristically for stability. All experiments use fixed hyperparameters across conditions to ensure comparability. Sensitivity to hyperparameter choices is left to future work.

E Annotator Guidelines

To ensure consistent understanding of definitions, annotators are provided with instructions and examples before continuing into the main workflow, which consist of two primary phases: (1) Reading and Tagging, and (2) Rating.

E.1 Narrative State Tagging

Annotators are instructed to read both the original and rewritten stories, highlighting text according to Todorov's Narrative Theory. They are provided with the following criteria for the five narrative stages:

Equilibrium: An initial state before any transformation that is then changed by a disruption.

Disruption of Equilibrium: The event/condition that disrupts equilibrium and initiates transformations.

Recognition of the Disruption: A character or narrator recognises that the disruption has occurred.

Attempt to Resolve Disruption: Actions taken in an attempt to address or resolve the disruption.

Return to New Equilibrium: A new equilibrium state reached after the transformation.

Key Tagging Guidelines

- 1: Tagging every single word or sentence is not required; only highlight parts specifically relevant to that stage.
- 2: Some stages may be implicit and cannot be tagged.
- 3: Filler or background detail can be left untagged.

E.2 Common Variations and Visual Examples

We provide annotators with annotated examples of different narrative structural variations.

Variation 1: In Medias Res The narrative begins with the Disruption of Equilibrium, skipping the initial equilibrium.

Example: "The kitchen pipe burst, flooding the tiles in seconds. I saw the water reaching the carpet and knew the basement was next. I grabbed a wrench and tightened the main valve. Within an hour, the plumber arrived and the house was dry again."

Variation 2: Non-Linear Narrative Events are revealed out of order, often starting with the resolution or final effort.

Example: "I crossed the finish line and collapsed onto the grass. Years of regular jogs had kept me healthy until a sudden ankle sprain last winter. Sitting in the doctor's office back then, I decided I would run a marathon to prove I could heal. Today, my medal finally felt real."

Variation 3: Cliffhanger Endings The text ends abruptly during an Attempt to Resolve Disruption, with no resolution shown.

Example: "The exam room was quiet as the students wrote. Suddenly, the fire alarm screeched throughout the hall. The proctor stood up and realized there was smoke beneath the door. She grabbed the evacuation log and shouted for everyone to follow her into the stairwell..."

E.3 Rating/Evaluation Guide

After tagging, annotators evaluate the edited text based on its adherence to specific criteria. They are instructed to focus strictly on whether the criteria are met, not on the stylistic quality of the writing.

Logical: The text must contain no internal contradictions. For example:

Logical: "Alex lost her job, but since she's an optimist, she walked home smiling."

Not Logical: "Alex lost the job that meant everything to her. She walked home with a smile." (Illogical without context.)

Rational: The text must follow a coherent Todorovian progression, with no contextual or causal disconnection. All events should be interpretable as belonging to compatible narrative stages. For example:

Rational: "The oven broke. I ordered pizza. We had a great meal." (Clear causal and contextual continuity.)

Not Rational: "I was eating dinner. The waiter spilt wine. I finally got my car fixed." (Narrative jumps without coherent stage progression.)

Complete_N: The text must preserve all key Todorovian stage(s) from the original narrative. An edited text is complete only if it reproduces all essential stages, even if the specific events differ. For example:

Complete: Orig: Health → Flu → Recovery.
Edit: Health → Accident → Surgery → Death.
(All original stages preserved.)

Incomplete: Orig: Health → Flu → Recovery.
Edit: Health → Accident → Surgery. (Missing the New Equilibrium stage.)

E.4 Handling Ambiguity

Annotators were given an **Unsure** (rating=2) option, to be used when the text is ambiguous.

Metric	Rating	A (%)	B (%)	C (%)
Logical	Agree	85.5	72.0	95.5
	Neutral	7.0	1.0	0.0
	Disagree	7.5	27.0	4.5
Rational	Agree	60.5	72.0	67.5
	Neutral	19.5	6.0	0.0
	Disagree	20.0	22.0	32.5
Complete	Agree	56.0	64.0	49.0
	Neutral	7.5	2.5	0.0
	Disagree	36.5	33.5	51.0
min _{LRC}	Agree	36.5	47.0	32.5
	Neutral	18.0	2.5	0.0
	Disagree	45.5	50.5	67.5
Narrativity	Agree+	49.0	53.5	14.0
	Agree	31.5	27.5	42.0
	Neutral	8.0	9.0	25.5
	Disagree	0.0	1.0	2.5
	Disagree-	11.5	9.0	16.0

Table 7: Distribution: 3 human annotators (A, B, C), percentages shown (200 per metric per annotator). Agree+ is Strongly Agree; Disagree- is Strongly Disagree.

E.5 Additional Agreement Statistics

Since the collection of human-authored and LLM-generated stories is logical and rational in most cases, the rating distributions are skewed (e.g., the average rating for Logical is 2.7/3 while the average for narrativity is 3.9/5. See Appendix Table 7), leading to the well-known prevalence/bias paradox (Zec et al., 2017) for κ (e.g. $\kappa = 0.18$ despite 80% agreement, see Appendix Table 8).

We want the dataset to test whether LLM-evaluators can distinguish between genuine human-level narratives and modern LLM generations. To accurately reflect the expected distribution of scores that the LLM-evaluator will see during d-RLAIF and evaluation, we do not alter the distributions and instead report Gwet (2014)’s AC2 alongside κ for inter-annotator reliability (top row of Appendix Table 9).

Unlike κ , which mathematically conflates high class prevalence with a high probability of chance agreement, AC2 estimates chance agreement under the assumption that raters do not guess randomly when they are certain of a prevalent class, thereby providing a useful complementary metric for skewed distributions.

Inter-annotator agreement across the evaluation criteria ranged from fair to moderate (Appendix Table 9), with AC2 scores between 0.55–0.76 and

κ values between 0.18–0.41 (or 0.30–0.41 if we treat the Logical criteria as an outlier since the distribution is skewed, leading to a low κ).

F Few-Shot Method

To select examples for our few-shot method, we use TimeTravel’s supervised training set. For each example in the training set, we use DeBERTa-v3-small (He et al., 2023) to calculate the similarity between the retelling and the original (using BERTScore (Zhang et al., 2020)).

We then rank the examples from most different to least different. We select the 1st, 500th and 1000th examples to diversify our 3 shot prompt to include examples that require high, medium, and low changes to the original ending.

G LLM-as-judge

In addition to using Spearman’s ρ (Section 3.4), we also evaluate the models’ alignment based on AC2 & κ , measured by examining the agreement scores when the LLM-as-judge is included in a 4-way calculation with the 3 human annotators (Table 9).

We observe that Gemini is the only model which is able to consistently improve alignment measured by both metrics in Complete_N, Overall, min_{LRC} and Narrativity.

Additionally, when we again categorise the ratings by the author, we see that only Gemini rates human authored stories with the highest scores, while Selene and Prometheus only rank human authored stories first in narrativity (Appendix Table 11). We choose to save Gemini for final evaluation, given the evidence that it is our most reliable and aligned LLM-as-judge.

While we ultimately only use a ‘no-reasoning’ prompt for the LLM-as-judge for training and evaluation, we initially tested two different output formats:

1. Generate the reasoning and then the score
2. Generate the score only⁶

We found that while reasoning increases alignment, it increases token costs far too greatly for it to be economically used in a d-RLAIF setting. So

⁶We instruct the LLM to output the score then the reasoning because post-hoc reasoning can improve accuracy (Lampinen et al., 2022; Lal et al., 2024). We stop inference after the score is generated, since it is impossible for the subsequent tokens to change the already-generated score.

Criteria	AC2	% Agree	κ
Logical	0.7648 [0.71, 0.82]	0.7967	0.1797 [0.0567, 0.2968]
Rational	0.7082 [0.65, 0.76]	0.7792	0.4068 [0.3024, 0.5035]
Complete _N	0.5611 [0.50, 0.62]	0.6750	0.3180 [0.2322, 0.4129]
min _{LRC}	0.5544 [0.49, 0.61]	0.6767	0.2981 [0.2116, 0.3851]
Narrativity	0.7415 [0.69, 0.79]	0.8660	0.3764 [0.2690, 0.4796]

Table 8: Inter-annotator agreement measured by Gwet’s AC2, % Agreement, and Quadratic Weighted κ . Values in brackets represent the 95% Confidence Interval (CI).

Annotator	Logical	Rational	Complete _N	Overall	min _{LRC}	Narrativity
Humans Only	0.76/0.18	0.71/0.41	0.56/0.32		0.55/0.30	0.74/0.38
w/ Selene-1-mini	0.80/0.09	0.71/0.27	0.58/0.25	0.49/0.20	0.54/0.20	0.73/0.31
w/ M-Prometheus	0.81/0.12	0.71/0.30	0.56/0.25	0.47/0.21	0.48/0.23	0.73/0.28
w/ Gemini-3-Flash	0.78/0.22	0.71/0.36	0.59/0.36	0.53/0.29	0.55/0.31	0.77/0.40

Table 9: Inter-annotator agreement, reported by AC2/ κ . “w/” indicates when an LLM-evaluator is added.

reasoning was excluded from being used in LLM-as-judge models.

H Training

Shown in Appednix Table 12, we conduct additional variations of post-training on Llama-3.1-8B-Instruct. Those trained by M-Prometheus achieve better performance when trained using R_O and worse when trained using R_N compared to those trained by Selene-1-mini. When the LLM-as-judge model is instructed to output a 5-point Likert score instead of the original 3-point and the group relative advantage is calculated from R_{O5} instead of R_O , the policy model’s outputs improve at the Logical and Rational criteria but significantly degrade at the Complete_N criteria. We also additionally run GRPO with ROUGE-L to train Llama 3.1-8B-instruct and find it worsens the performance of the model.

Author	Logical	Rational	Complete _N	min _{LRC}	Narrativity
human	2.98 / 2.54 / 2.96	2.72 / 2.66 / 2.52	2.82 / 2.72 / 2.28	2.58 / 2.30 / 1.96	4.70 / 4.62 / 3.72
Olmo	2.80 / 2.50 / 2.96	2.52 / 2.66 / 2.48	1.92 / 2.04 / 2.00	1.70 / 1.84 / 1.72	3.84 / 4.00 / 3.36
Qwen	2.64 / 2.40 / 2.92	2.36 / 2.52 / 2.36	2.50 / 2.54 / 2.12	2.12 / 2.12 / 1.68	4.22 / 4.32 / 3.50
Llama	2.70 / 2.36 / 2.80	2.02 / 2.16 / 2.04	1.54 / 1.92 / 1.52	1.24 / 1.60 / 1.24	3.50 / 3.68 / 2.84

Table 10: Human and LLM author comparison across evaluation dimensions with 3 human annotators (A / B / C).

Model	Logical	Rational	Complete _N	Overall	min _{LRC}	Narrativity
human	2.80 / 2.94 / 3.00	2.92 / 2.94 / 2.96	2.96 / 2.38 / 2.82	2.52 / 2.46 / 2.62	2.72 / 2.38 / 2.80	4.90 / 3.62 / 5.00
Qwen3	2.60 / 3.00 / 2.96	2.80 / 3.00 / 3.00	2.36 / 2.54 / 2.86	2.40 / 2.74 / 2.76	2.04 / 2.54 / 2.82	4.78 / 3.42 / 4.88
Olmo3	2.72 / 3.00 / 3.00	2.80 / 2.98 / 3.00	1.88 / 2.14 / 2.72	2.20 / 2.40 / 2.76	1.76 / 2.14 / 2.72	4.38 / 3.26 / 4.76
Llama3.1	2.68 / 2.92 / 2.94	2.24 / 2.72 / 2.52	1.52 / 1.90 / 2.22	1.60 / 2.10 / 2.28	1.36 / 1.86 / 2.12	3.82 / 2.76 / 4.18

Table 11: Model comparison across evaluation dimensions with LLM judges (Gemini / Selene / Prometheus).

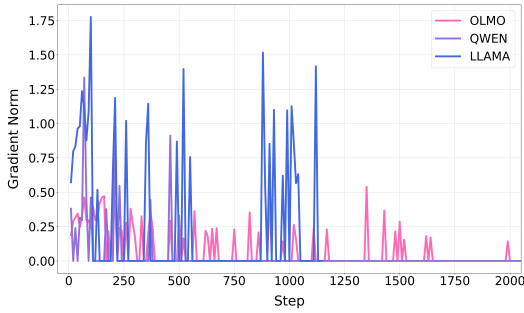


Figure 4: Policy model gradient norm during d-RLAIF using the reward signal R_0 generated by Selene-1-mini.

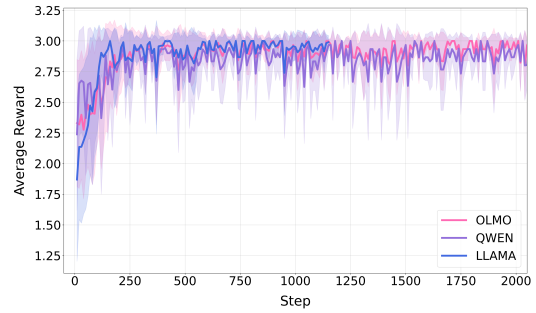


Figure 5: Mean and standard deviation of the reward signal R_0 generated by Selene-1-mini during d-RLAIF.

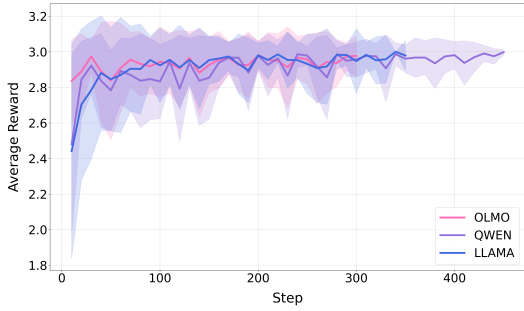


Figure 6: Mean and standard deviation of the reward signal R_0 generated by M-Prometheus during d-RLAIF.

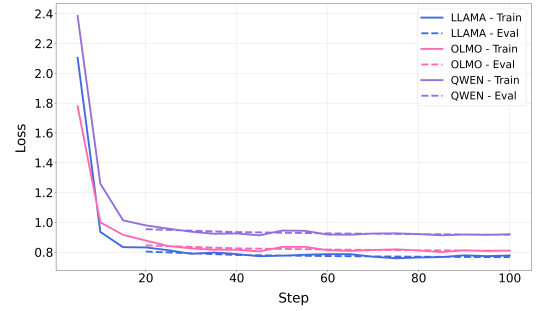


Figure 7: Training and evaluation loss during SFT.

Author	Logical	Rational	Complete _N	min _{LRC}	Narrativity	BLEU-4	ROUGE-L
LLAMA-3.1-8B-INSTRUCT							
BASE	2.735	2.385	2.528	2.080	4.347	0.340	0.502
few-shot	2.768	2.615	2.505	2.213	4.303	0.306	0.473
SFT	2.320	2.501	2.934	2.260	4.692	0.810	0.835
GRPO _{ROUGE-L}	2.007	2.212	2.938	1.952	4.628	0.855	0.863
RRR _O -Prometheus	2.801	2.825	2.594	2.400	4.534	0.179	0.386
RRR _{O5} -Prometheus	2.917	2.943	2.059	2.019	4.53	0.007	0.194
RRR _N -Prometheus	2.655	2.717	2.246	2.058	4.219	0.163	0.357
RRR _O -Selene	1.820	2.080	2.779	1.761	4.489	0.174	0.458
RRR _N -Selene	2.785	2.860	2.740	2.554	4.733	0.087	0.299

Table 12: Evaluation of post-trained LLMs using the test split (n=1871) from TimeTravel dataset using Gemini-3-Flash. Additional models trained on Llama-3.1.

I Significance Testing

Model	Method	Logical	Rational	Complete _N	min _{LRC}	Narrativity
LLAMA-3.1-8B-INST.	BASE	2.735*	2.385*	2.528*	2.080*	4.347*
	few-shot	2.768	2.615*	2.505*	2.213*	4.303*
	SFT	2.320*	2.501*	2.934*	2.260*	4.692
	RRR _O -Prometheus	2.801	2.825*	2.594*	2.400*	4.534*
	RRR _N -Selene	2.785 ^{Base}	2.860 ^{Base}	2.740 ^{Base}	2.554 ^{Base}	4.733 ^{Base}

Table 13: Evaluation of Llama-3.1-8B variants using the test split ($n = 1871$) from the TimeTravel dataset. Values marked with * indicate a statistically significant difference ($p < 0.05$ via Wilcoxon signed-rank test) compared to the **RRR_N-Selene** baseline. Bolded values represent the best performance for each metric.

Model	Method	Logical	Rational	Complete _N	min _{LRC}	Narrativity
QWEN-3-8B	BASE	2.429*	2.132*	2.759*	1.940*	4.551*
	few-shot	2.543*	2.219*	2.686	2.006*	4.559*
	SFT	2.402*	2.532*	2.925*	2.317*	4.693*
	RRR _O -Prometheus	2.717*	2.336*	2.566*	2.024*	4.356*
	RRR _N -Selene	2.781 ^{Base}	2.890 ^{Base}	2.705 ^{Base}	2.516 ^{Base}	4.894 ^{Base}

Table 14: Evaluation of Qwen-3-8B variants using the test split ($n = 1871$) from the TimeTravel dataset. Values marked with * indicate a statistically significant difference ($p < 0.05$ via Wilcoxon signed-rank test) compared to the **RRR_N-Selene** baseline. Bolded values represent the best performance for each metric within the model family.

Model	Method	Logical	Rational	Complete _N	min _{LRC}	Narrativity
OLMO-3-7B-INST.	BASE	2.762	2.812*	2.400*	2.222*	4.328*
	few-shot	2.745*	2.837*	2.377*	2.208*	4.415*
	SFT	2.295*	2.441*	2.910*	2.211*	4.647*
	RRR _O -Prometheus	2.770	2.857*	2.491*	2.325*	4.460*
	RRR _N -Selene	2.793 ^{Base}	2.893 ^{Base}	2.595 ^{Base}	2.434 ^{Base}	4.899 ^{Base}

Table 15: Evaluation of Olmo-3-7B variants using the test split ($n = 1871$) from the TimeTravel dataset. Values marked with * indicate a statistically significant difference ($p < 0.05$ via Wilcoxon signed-rank test) compared to the **RRR_N-Selene** baseline. Bolded values represent the best performance for each metric within the model family.

J Example Annotations (with Additional Reasoning/Rationale)

Narrativity (Equilibrium→Disruption→Recognition→Attempt→New Equilibrium)

A: (4) My nephew had a bad back one day. He took some ibuprofen and was fine. The pills were never even opened, and his back felt better by the next day.

B: (4) My nephew had a bad back one day. He took some ibuprofen and was fine. The pills were never even opened, and his back felt better by the next day.

C: (4) My nephew had a bad back one day. He took some ibuprofen was fine. The pills were never even opened, and his back felt better by the next day.

Logical

A: (1) It's illogical and inconceivable that the nephew took ibuprofen yet the pills were never opened.

B: (3) The edited text is logically coherent. Taking ibuprofen and leaving the prescription pills unopened is fully consistent.

C: (3) The story does not present inconceivable scenarios. Whilst at first there appears to be a logical inconsistency, "the pills" do not necessarily refer to ibuprofen.

Rational

A: (1) There's causal and contextual disconnection between taking ibuprofen and the pills being never opened.

B: (3) The edited text is rationally interpretable. The progression from pain to treatment to recovery is clear.

C: (1) The story cannot rationally be interpreted using Todorov's stages in its entirety because there is a causal disconnection; the most likely interpretation is that "the pills" refer to the ibuprofen in sentence two.

Completeness

A: (1) There are no recognisable disruptions. The edited text is not as much of a narrative as the original.

B: (3) The edited text is complete. It preserves the same broad Todorovian stage structure as the original.

C: (3) All stages of Todorov's theory presented in the original text are present in the rewritten story.

Table 16: An example of valid annotator disagreement when tagging and evaluating stories.

Narrativity (Equilibrium→Disruption→Recognition→Attempt→New Equilibrium)

A: (5) I started using coupons. I realized it was too much work, with no real savings. I still collected coupons from the Sunday paper, emails and snail mail, but I didn't put much effort into researching stores that double the value of coupons. In the end, I didn't clip, tear, or cut any coupons, and I didn't end up saving any money last week.

B: (5) I wanted to see if i could save some money so I started using coupons. I realized it was too much work, with no real savings. I still collected coupons from the Sunday paper, emails and snail mail, but I didn't put much effort into researching stores that double the value of coupons. In the end, I didn't clip, tear, or cut any coupons, and I didn't end up saving any money last week.

C: (4) I wanted to see if i could save some money so I started using coupons. I realized it was too much work, with no real savings. I still collected coupons from the Sunday paper, emails and snail mail, but I didn't put much effort into researching stores that double the value of coupons. In the end, I didn't clip, tear, or cut any coupons, and I didn't end up saving any money last week.

Logical

A: (1) It is illogical that the narrator would keep collecting coupons after realising that it was too much work with no real savings.

B: (3) The edited text is logically coherent. The speaker tries coupons, recognises the effort outweighs the savings, reduces that effort, and ends without saving money.

C: (3) The events of the story progress without inconceivable scenarios. A logical inconsistency could be interpreted between sentences three and four... However, this inconsistency could be explained if we viewed the time of narration as a flashback.

Rational

A: (1) There is no rational causal connection between realising that coupons are too much work and continuing to collect coupons.

B: (3) The edited text is rationally interpretable. The causal progression from initial goal to failed outcome is clear.

C: (3) The events progress rationally: the character wants to see if they can save money with coupons, they realise it is too much work, then stop attempting to do so.

Completeness

A: (3) It has a narrative structure of disruption, recognition, attempt and equilibrium. Even though it doesn't make sense, it's as much of a narrative as the original.

B: (3) The edited text is complete. It includes disruption, recognition, attempt, and a new equilibrium, and is at least as stage-complete as the original.

C: (1) The rewritten story is less complete than the original because it omits the crucial "attempt" stage. Even though the rewritten story includes as many Todorovian stages, it is not as narratively complete because character agency isn't represented as fully.

Table 17: An example of valid annotator disagreement when tagging and evaluating stories.

Narrativity (Equilibrium→Disruption→Recognition→Attempt→New Equilibrium)

A: (4) Evan was ready to buy himself a bicycle. He went to the store but forgot his money. Then he looked around some more, trying to find something else he could afford or a way to get home.

B: (4) Evan was ready to buy himself a bicycle. He went to the store but forgot his money. Then he looked around some more, trying to find something else he could afford or a way to get home.

C: (3) Evan was ready to buy himself a bicycle. He went to the store but forgot his money. Then he looked around some more, trying to find something else he could afford or a way to get home.

Logical

A: (3) I interpret 'but forgot his money' as Evan forgetting a large sum of money intended for buying a bicycle. Evan may still have the usual money he carries everyday, and is able to logically find something else he could afford.

B: (3) The edited text is logical. Forgetting his money and then looking for an alternative or a way home is internally consistent.

C: (1) There is a logical impossibility in the story because: 1. Evan forgets his money, but 2. proceeds to look around the store for something he can "afford," which would be nothing.

Rational

A: (3) Looking for an alternative after forgetting the money to buy the bicycle is causally and contextually connected.

B: (1) The edited text is irrational in Todorovian terms. The attempt is present, but the narrative does not clearly resolve the disruption or establish a justified new equilibrium.

C: (1) The events of the story are causally disconnected.

Complete_N

A: (1) The edited text does not reach a new equilibrium whereas the original text does: "Then he rode home!" so it is not as complete of a narrative as the original.

B: (3) The edited text is complete. It includes disruption, attempt, and a form of new equilibrium, which matches the broad stage structure of the original.

C: (1) The original text presented a disruption in Evan's desire to buy a bicycle; an 'attempt at doing so in looking around the store; and a 'new equilibrium' after he purchases the bicycle he initially desired.

Table 18: An example of valid annotator disagreement when tagging and evaluating stories.

Narrativity (Equilibrium→Disruption→Recognition→Attempt→New Equilibrium)

A: (1) I had to go to Boston for a work trip. It was a boring trip filled with conferences and no fun. We attended several conferences and had no memorable experiences, but I learned a lot for my career that I will carry with me.

B: (4) I had to go to Boston for a work trip. It was a boring trip filled with conferences and no fun. We attended several conferences and had no memorable experiences, but I learned a lot for my career that I will carry with me.

C: (4) I had to go to Boston for a work trip. It was a boring trip filled with conferences and no fun. We attended several conferences and had no memorable experiences, but I learned a lot for my career that I will carry with me

Logical

A: (3) The scenario is conceivable and breaks no logical rules. "boring/unmemorable" and "learn a lot" are not mutually exclusive.

B: (1) The edited text contains a weak internal contradiction: the trip is described as having "no memorable experiences," yet it also results in a significant and lasting professional benefit. Because this shift is not sufficiently explained, the narrative's internal logic is weakened.

C: (3) At first, it seems there is a logical impossibility because the narrator, after having an uneventful trip "filled with [...] no fun" and "no memorable experiences," expresses in the last line that they learnt a lot. However, the conjunction "but" suggests that even though they had a bad time, the experience nevertheless taught them something. Bad or uneventful experiences can also be informative.

Rational

A: (3) Looking for an alternative after forgetting the money to buy the bicycle is causally and contextually connected.

B: (1) Although the text signals disruption and ends in a new equilibrium, it does not rationally develop the causal steps between them. The transition from a boring, uneventful trip to a meaningful career lesson lacks narrative grounding.

C: (3) Again, there appears to be a causal disconnection within the last sentence (from recognition to equilibrium). However, the word "but" makes the story rational. If the conjunction were "and", the story would be causally disconnected.

Complete_N

A: (1) There is no disruption in the edited text. It is not as complete as the original narrative, which has disruption and new equilibrium.

B: (1) The edited text is not fully complete because it excludes the Attempt stage. Without that stage, the sequence of Todorovian development is reduced, making the edited version less narratively complete than the original.

C: (3) Even though the fourth stage of Todorov's theory seems more important to create a complete narrative, the rewritten story feels just as complete because of the final clause "but I learned a lot for my career that I will carry with me."

Table 19: An example of valid annotator disagreement when tagging and evaluating stories.

Narrativity (Equilibrium→Disruption→Recognition→Attempt→New Equilibrium)

A: (5) In the last mile of a marathon only two runners were in contention. One runner got cramps on the ground and tripped the other. They both rested on the ground. before they drop out.

B: (5) In the last mile of a marathon only two runners were in contention. They both got cramps and had to drop out of the race. One runner got cramps on the ground and tripped the other. They both rested on the ground. They rested for a few second before they drop out.

C: (3) In the last mile of a marathon only two runners were in contention. They both got cramps and had to drop out of the race. One runner got cramps on the ground and tripped the other. They both rested on the ground. They rested for a few second before they drop out.

Logical

A: (3) Conceivable. The events are not narrated in chronological order. “They both got cramps and had to drop out” is a summary of the explanations that follow.

B: (1) The edited text is not logical. The claim that one runner “got cramps on the ground and tripped the other” is internally unclear and weakens the event sequence.

C: (3) The events of the story progress logically, free from internal contradictions; the two runners both get cramps, drop to the floor, and will eventually drop out.

Rational

A: (3) Everything is contextually and causally connected even if it’s not in chronological order.

B: (1) The edited text is not rational. The causal progression from cramping to tripping to resting to dropping out is not presented in a fully coherent way.

C: (3) All events can be rationally interpreted using the stages. The last two sentences are an elaboration of the disruption that I tagged in sentence two.

Complete_N

A: (3) It is just as complete of a narrative as the original. Cramps are the disruption, resting is the attempt to resolve, and dropping out is the new equilibrium.

B: (3) The edited text is complete in stage coverage. It still presents equilibrium, disruption, recognition, attempt, and a final outcome.

C: (1) The rewritten story omits an attempt at new equilibrium because there is no activity to overcome the obstacle of falling to the ground with cramps.

Table 20: An example of valid annotator disagreement when tagging and evaluating stories.

K Example Outputs

Source	Text	L	R	C	M	N
Original Story						
Premise	Max and May were at a tavern.					
Initial	They saw the dartboard and challenged each other to a game.					
Original Ending	Max played very well. But May thru even better than him. May won the game.					
Counterfactual Retelling						
Premise	Max and May were at a tavern.					
Counterfactual	They didn't play darts there because they were not competitive.					
Edited Ending	???					
Human Ground Truth						
Crowd Worker	Max played pool very well. But May thought she was even better than him. May won the game, but neither really cared who won.	3	3	3	3	5
Olmo						
Base	Max and May decided not to play since they weren't competitive.	3	3	1	1	1
Few-Shot	Max and May talked and laughed instead of playing. They enjoyed the atmosphere together.	3	3	3	3	1
SFT	Max played very well. But May threw even better than him. May won the game.	1	1	3	1	1
RRR _O	They left the tavern satisfied with their decision not to compete.	3	3	3	3	1
RRR _N	They decided to sing a song instead. The tavern erupted in laughter and applause—both agreed it was far more fun than competition.	3	3	3	3	5
Llama						
Base	Max and May were at a tavern. They saw the dartboard but decided not to play because they were not competitive.	3	1	1	1	1
Few-Shot	They decided to skip the dartboard and instead ordered a couple of beers to relax and enjoy the evening.	3	3	1	1	1
SFT	Max played very well at cards. But May threw even better than him. May won the game.	1	3	3	1	5
RRR _O	They ended up chatting over a pint instead, enjoying each other's company without the competitive edge that darts had brought out in them.	3	3	1	1	5
RRR _N	They ordered a round and struck up a conversation. Enjoying each other's company, they decided to share a story. Max won over May with his tale.	3	3	3	3	5
Qwen						
Base	Max played very well. But May threw even better than him. May won the game.	1	1	3	1	1
Few-Shot	They didn't play darts there because they were not competitive. Instead, they talked about their favorite books and shared stories about their travels.	3	1	1	1	1
SFT	Max played very well at darts. But May threw even better than him. May won the game.	1	1	3	1	5
RRR _O	They decided to play a different game instead, one that allowed them to enjoy each other's company without competition.	3	3	1	1	4
RRR _N	They ordered drinks, then debated chess. Max won, but May vowed to train. They vowed to rematch.	3	3	3	3	5
Codex						
Code-Prompt	Max played very well. But May thru even better than him. They had a good time.	1	1	1	1	1

Table 21: Example outputs from the test set. Ratings are given by LLM-as-judge (Gemini-3-Flash): L=Logical, R=Rational, C=Complete_N, M=min_{LR}, N=Narrativity.

Source	Text	L	R	C	M	N
Original Story						
Premise	Kimmy was on a trip to Paris.					
Initial	Her plane had just landed and she and her classmates her excited.					
Original Ending	They couldn't wait to start seeing museums and churches! They spent lots of time shopping and seeing historical sites. Kimmy was very excited.					
Counterfactual Retelling						
Premise	Kimmy was on a trip to Paris.					
Counterfactual	Kimmy's plane crashed over the ocean and killed everyone on board.					
Edited Ending	???					
Human Ground Truth						
Crowd Worker	Kimmy had just wanted to see the museums and churches! Her parents spent a lot of time planning her funeral. Kimmy was dead.	3	3	3	3	5
Olmo						
Base	Given the tragic counterfactual, there is no living person to experience excitement or see museums and churches, so the premise cannot be fulfilled in the same way as the original ending.	3	3	1	1	3
Few-Shot	There is no edited ending, as the counterfactual scenario is a tragic accident with no possible positive or alternative ending.	3	3	1	1	3
SFT	They couldn't wait to start seeing museums and churches! They spent lots of time shopping and seeing historical sites. Kimmy was very excited.	1	1	3	1	1
RRR _O	Due to the tragic accident, the trip never happened, and Kimmy and her classmates never experienced Paris or its attractions. There were no museums, churches, shopping, or historical sites to see. Kimmy could never have been excited in the same way.	3	3	3	3	4
RRR _N	Grief-stricken, the world mourned. No legacy remained—only memories in the waves. The tragedy ended hope; Paris dreams faded. The story closed with silent sorrow.	3	3	3	3	5
Llama						
Base	They never made it to Paris. Kimmy's plane crashed over the ocean and there were no survivors.	3	1	3	1	3
Few-Shot	Kimmy's plane crashed over the ocean and killed her.	3	3	1	1	3
SFT	They couldn't wait to start seeing museums and churches! They spent lots of time shopping and seeing historical sites. Kimmy was very excited.	1	1	3	1	1
RRR _O	The plane wreckage was found at the bottom of the ocean, a tragic reminder of the devastating crash that took Kimmy and her classmates' lives.	3	3	3	3	4
RRR _N	Panicked, Kimmy clung to debris, the only survivor. She was rescued and returned home, re-evaluating life. Kimmy was forever changed.	1	3	3	1	5
Qwen						
Base	There were no survivors, and Kimmy's adventure in Paris was cut short tragically. The classmates and everyone on board lost their lives in the crash.	3	3	3	3	3
Few-Shot	Kimmy's plane crashed over the ocean and killed everyone on board. There was no one left to see museums and churches, and no one to share the excitement with. Kimmy's trip ended tragically before it could begin.	3	1	3	1	5
SFT	They couldn't wait to start seeing museums and churches. They spent lots of time shopping and seeing historical sites. Kimmy was very excited about the trip.	1	1	3	1	3
RRR _O	Unfortunately, Kimmy and her classmates were never able to explore Paris. The plane crash claimed their lives before they could experience the city's wonders.	3	3	3	3	4
RRR _N	Miraculously, Kimmy survived and was rescued. She mourned her classmates but vowed to honor them. She traveled solo, explored Paris, and wrote a heartfelt tribute. Her story inspired many.	1	3	3	1	5
Codex						
Code-Prompt	Kimmy was very sad. She was very sad. She was very sad.	1	1	3	1	3

Table 22: Example outputs from the test set. Ratings are given by LLM-as-judge (Gemini-3-Flash): L=Logical, R=Rational, C=Complete_N, M=min_{LRC}, N=Narrativity.

Source	Text	L	R	C	M	N
Original Story						
Premise	Amy’s class went on a field trip to see A Mid Summer Night’s Dream.					
Initial	They arrived at the theater and found their seats.					
Original Ending	Amy was in awe of the actors’ discipline and talent. After the play she was able to speak with the cast. Amy was thrilled.					
Counterfactual Retelling						
Premise	Amy’s class went on a field trip to see A Mid Summer Night’s Dream.					
Counterfactual	They forgot their tickets in the classroom.					
Edited Ending	???					
Human Ground Truth						
Crowd Worker	Amy was in awe at their stupidity. She saw the play the next week and then she was able to speak with the cast. Amy was thrilled.	3	3	3	3	5
Llama						
Base	They had to go back to the classroom to get the tickets, but were able to find their seats just in time for the start of the play. Amy was in awe of the actors’ discipline and talent. After the play, they were unable to speak with the cast because they had to rush back to the classroom to pick up their forgotten items, but Amy was still thrilled to have seen the performance.	3	3	3	3	5
Few-Shot	Amy was frustrated that they couldn’t get in to see the play. They had to go back to the classroom to get the tickets. After finally getting in, Amy was in awe of the actors’ discipline and talent. However, she was disappointed that they missed the beginning of the play.	3	3	3	3	5
SFT	Amy was in awe of the actors discipline and talent. After the play she was able to speak with the cast. Amy was thrilled.	2	1	3	3	3
RRR _O	Panicked, they rushed back to the classroom to retrieve their tickets, but unfortunately, the class had already left for the theater. Disappointed, they decided to watch from outside the theater, still managing to catch a glimpse of the performance and hearing the sound of the actors’ voices carrying through the windows.	1	3	1	1	5
RRR _N	Panicked, they rushed back to school. Luckily, they found the tickets. They arrived just in time, and Amy was amazed by the performance. She met the cast backstage, beaming with joy.	3	3	3	3	5
Qwen						
Base	Amy was in awe of the actors’ discipline and talent. After the play, she realised she had forgotten her tickets in the classroom and couldn’t speak with the cast. Amy was disappointed.	1	1	3	1	5
Few-Shot	They forgot their tickets in the classroom. Amy was disappointed and worried they would miss the play. After a frantic search, they found the tickets just in time to enter the theater.	3	1	3	1	5
SFT	Amy was in awe of the actors discipline and talent. After the play she was able to speak with the cast. Amy was thrilled.	2	1	3	1	5
RRR _O	Amy was frustrated when they realized they had forgotten their tickets in the classroom. She tried to find a solution but ultimately had to leave the theater without seeing the play.	3	3	3	3	5
RRR _N	They raced back, found tickets, and rushed in. The play was unforgettable—Amy cheered, then met the cast proudly.	3	3	3	3	5
Olmo						
Base	Due to forgetting their tickets, Amy and her classmates couldn’t enter the theater. Everyone was disappointed but decided to talk to the teachers about the situation to see if they could get in later. The experience left them eager to try again.	3	3	1	1	5
Few-Shot	Amy was frustrated that they couldn’t go into the show. She wished they had remembered the tickets so they could have enjoyed the performance and maybe even spoken with the cast afterward.	3	3	1	1	4
SFT	Amy was in awe of the actors discipline and talent. After the play she was able to speak with the cast. Amy was thrilled.	2	1	3	1	3
RRR _O	Amy was frustrated by not being able to enter the theater and had to explain the situation to staff before being redirected to a nearby event. Despite disappointment, she appreciated seeing live theatre from a different angle.	3	3	3	3	5
RRR _N	Realising the mix-up, the group rushed back to class, scrambled for tickets, and finally boarded the bus—reassured and racing to the theater. The play began, applause echoed, and relief washed over them.	3	3	3	3	5
Codex						
Code-Prompt	Amy was in awe of the actors discipline and talent. After the play she was able to speak with the cast. Amy was thrilled.	1	1	3	1	3

Table 23: Example outputs from the test set. Ratings are given by LLM-as-judge (Gemini-3-Flash): L=Logical, R=Rational, C=Complete_N, M=min_{LRC}, N=Narrativity.