

Date Fragments: A Hidden Bottleneck of Tokenisation for Temporal Reasoning

Gagan Bhatia¹ Maxime Peyrard² Wei Zhao¹

¹University of Aberdeen ²Université Grenoble Alpes & CNRS
{g.bhatia.24,wei.zhao}@abdn.ac.uk

Abstract

Modern BPE tokenisers often split calendar dates into meaningless fragments, e.g., “20250312” → “202”, “503”, “12”, inflating token counts and obscuring the inherent structure needed for robust temporal reasoning. In this work, we (1) introduce a simple yet interpretable metric, termed date fragmentation ratio, that measures how faithfully a tokeniser preserves multi-digit date components; (2) release DATEAUGBENCH, a suite of 6500 examples spanning three temporal reasoning tasks: context-based date resolution, format-invariance puzzles, and date arithmetic across historical, contemporary, and future time periods; and (3) through layer-wise probing and causal attention-hop analyses, uncover an emergent date-abstraction mechanism whereby large language models stitch together the fragments of month, day, and year components for temporal reasoning. Our experiments show that excessive fragmentation correlates with accuracy drops of up to 10 points on uncommon dates like historical and futuristic dates. Further, we find that the larger the model, the faster the emergent date abstraction heals date fragments. Lastly, we observe a reasoning path that LLMs follow to assemble date fragments, typically differing from human interpretation (year → month → day). Our datasets and code are made publicly available [here](#).

1 Introduction

Understanding and manipulating dates is a deceptively complex challenge for modern large language models (LLMs). Unlike ordinary words, dates combine numeric and lexical elements in rigidly defined patterns—ranging from compact eight-digit strings such as 20250314 to more verbose forms like “March 14, 2025” or locale-specific variants such as “14/03/2025.” Yet despite their structured nature, these date expressions often fall prey to subword tokenisers that fragment them into

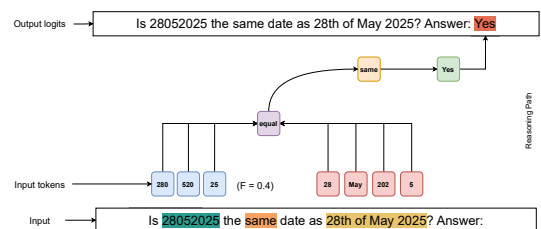


Figure 1: Internal processing of dates for temporal reasoning. Here $F=0.4$ shows the date fragmentation ratio.

semantically meaningless pieces. A tokeniser that splits “2025-03-14” into “20”, “25”, “-0”, “3”, “-1”, “4” not only inflates the token count but also severs the natural boundaries of year, month, and day. This fragmentation obscures temporal cues and introduces a hidden bottleneck: even state-of-the-art LLMs struggle to resolve, compare, or compute dates accurately when their internal representations have been so badly fragmented. This issue has a critical impact on real-world applications:

Mis-tokenised dates can undermine scheduling and planning workflows, leading to erroneous calendar invites or appointments (Vasileiou and Yeoh, 2024). They can skew forecasting models in domains ranging from time-series analysis (Tan et al., 2024; Chang et al., 2023) to temporal knowledge graph reasoning (Wang et al., 2024). In digital humanities and historical scholarship, incorrect splitting of date expressions may corrupt timelines and misguide interpretative analyses (Zeng, 2024). As LLMs are increasingly deployed in cross-temporal applications, such as climate projection (Wang and Karimi, 2024), economic forecasting (Carriero et al., 2024; Bhatia et al., 2024), and automated curriculum scheduling (Vasileiou and Yeoh, 2024), the brittleness introduced by subword fragmentation poses a risk of propagating temporal biases and inaccuracies into downstream scientific discoveries and decision-making systems (Tan et al., 2024).

In this work, we provide a pioneer outlook on

the impact of date tokenisation on downstream temporal reasoning. Figure 1 illustrates how dates are processed internally for temporal reasoning. Our contributions are summarized as follows:

- (i) We introduce DATEAUGBENCH, a benchmark dataset comprising 6,500 examples with 21 date formats. It is leveraged to evaluate a diverse array of LLMs from 8 model families in three temporal reasoning tasks.
- (ii) We present date fragmentation ratio, a metric that measures how fragmented the tokenisation outcome is compared to the actual year, month, and day components. We find that the fragmentation ratio generally correlates with temporal reasoning performance, namely that the more fragmented the tokenisation, the worse the reasoning performance.
- (iii) We analyse internal representations by tracing how LLMs “heal” a fragmented date embeddings in their layer stack—an emergent ability that we term *date abstraction*. We find that larger models can quickly compensate for date fragmentation at early layers to achieve high accuracy for date equivalence reasoning.
- (iv) We leverage causal analysis to interpret how LLMs stitch date fragments for temporal reasoning. Our results show that LLMs follow a reasoning path that is typically not aligned with human interpretation (year $\rightarrow$ month $\rightarrow$ day), but relies on subword fragments that statistically represent year, month, and day, and stitch them in a flexible order that is subject to date formats.

Our work fills the gap between tokenisation research (Goldman et al., 2024; Schmidt et al., 2024) and temporal reasoning (Su et al., 2024; Fatemi et al., 2024), and we suggest future work to consider date-aware vocabularies and adaptive tokenisers to ensure that date components remain intact.

2 Related Works

Tokenisation as an information bottleneck. Recent scholarship interrogates four complementary facets of sub-word segmentation: (i) *tokenisation fidelity*, i.e. how closely a tokeniser preserves semantic units: Large empirical studies show that higher compression fidelity predicts better downstream accuracy in symbol-heavy domains such as code, maths and dates (Goldman et al., 2024; Schmidt et al., 2024); (ii) *numeric segmentation strategies*

that decide between digit-level or multi-digit units: Previous work demonstrates that the choice of radix-single digits versus 1-3 digit chunks induces stereotyped arithmetic errors and can even alter the complexity class of the computations LLMs can realise (Singh and Strouse, 2024; Zhou et al., 2024); (iii) *probabilistic or learnable tokenisers* whose segmentations are optimised jointly with the language model: Theory frames tokenisation as a stochastic map whose invertibility controls whether maximum-likelihood estimators over tokens are consistent with the underlying word distribution (Gastaldi et al., 2024; Rajaraman et al., 2024) and (iv) *pre-/post-tokenisation adaptations* that retrofit a model with a new vocabulary: Zheng et al. (2024) introduce an *adaptive tokeniser* that co-evolves with the language model, while Liu et al. (2025) push beyond the “sub-word” dogma with *SuperBPE*, a curriculum that first learns subwords and then merges them into cross-whitespace “super-words”, cutting average sequence length by 27 %. Complementary studies expose and correct systematic biases introduced by segmentation (Phan et al., 2024) and propose *trans-tokenisation* to transfer vocabularies across languages without re-training the model from scratch (Remy et al., 2024). Our work builds on these insights but zooms in on calendar dates—a hybrid of digits and lexical delimiters whose multi-digit fields are routinely shredded by standard BPE, obscuring cross-field regularities crucial for temporal reasoning.

Temporal reasoning in large language models.

Despite rapid progress on chain-of-thought and process-supervised reasoning, temporal cognition remains a conspicuous weakness of current LLMs. Benchmarks such as TIMEBENCH (Chu et al., 2024), TEMPReason (Tan et al., 2023), TEST-OF-TIME (Fatemi et al., 2024), MENATQA (Wei et al., 2023) and TIMEQA (Chen et al., 2021) reveal large gaps between model and human performance across ordering, arithmetic and co-temporal inference. Recent modelling efforts attack the problem from multiple angles: temporal-graph abstractions (Xiong et al., 2024), instruction-tuned specialists such as TIMO (Su et al., 2024), pseudo-instruction augmentation for multi-hop QA (Tan et al., 2023), and alignment techniques that re-ground pretrained models to specific calendar years (Zhao et al., 2024). Yet these approaches assume a faithful internal representation of the input dates themselves. By introducing the notion of *date frag-*

mentation and demonstrating that heavier fragmentation predicts up to ten-point accuracy drops on DATEAUGBENCH, we uncover a failure mode that is *orthogonal* to reasoning algorithms or supervision: errors arise before the first transformer layer, at the level of subword segmentation. Addressing this front-end bottleneck complements existing efforts to further improve LLMs for temporal reasoning.

3 DateAugBench

We introduce DATEAUGBENCH, benchmark designed to isolate the impact of date tokenisation on temporal reasoning in LLMs. DATEAUGBENCH comprises 6,500 augmented examples drawn from two established sources, TIMEQA (Chen et al., 2021) and TIMEBENCH (Chu et al., 2024), distributed across three tasks splits (see Table 1). Across all the splits, our chosen date formats cover a spectrum of common regional conventions (numeric with slashes, dashes, or dots; concatenated strings; two-digit versus four-digit years) and deliberately introduce fragmentation for atypical historical (e.g. “1799”) and future (e.g. “2121”) dates. This design enables controlled measurement of how tokenisation compression ratios and subsequent embedding recovery influence temporal reasoning performance.

Context-based task. In the *Context-based* split, we sample 500 question–context pairs from TIMEQA, each requiring resolution of a date mentioned in the passage (e.g. Which team did Omid Namazi play for in 06/10/1990?). Every date expression is systematically rendered in six canonical serialisations—including variants such as MM/DD/YYYY, DD-MM-YYYY, YYYY.MM.DD and concatenations without delimiters—yielding 3,000 examples that jointly probe tokenisation fragmentation and contextual grounding.

Simple Format Switching task. The *Simple Format Switching* set comprises 150 unique date pairs drawn from TIMEBENCH, posed as binary same-day recognition questions (e.g. “Are 20251403 and 14th March 2025 referring to the same date?”). Each pair is presented in ten different representations, spanning slash-, dash-, and dot-delimited formats, both zero-padded and minimally notated, to stress-test format invariance under maximal tokenisation drift. This produces 1,500 targeted examples of pure format robustness. We also have examples

where the dates are not equivalent, complicating the task.

Date Arithmetic task. The *Date Arithmetic* split uses 400 arithmetic instances from TIMEBENCH (e.g. What date is 10,000 days before 5/4/2025?). With the base date serialised in five distinct ways—from month-day-year and year-month-day with various delimiters to compact eight-digit forms. This results in 2,000 examples that examine the model’s ability to perform addition and subtraction of days, weeks, and months under various token fragmentation.

4 Experiment Design

4.1 Date Tokenisation

Tokenisers. For tokenisation analysis, we compare a deterministic, rule-based *baseline tokeniser* against model-specific tokenisers. The baseline splits each date into its semantic components—year, month, day or Julian day—while preserving original delimiters. For neural models, we invoke either the OpenAI Tiktok encodings (for gpt-4, gpt-3.5-turbo, gpt-4o, text-davinci-003) or Hugging Face tokenisers for open-source checkpoints. Every date string is processed to record the resulting sub-tokens, token count, and reconstructed substrings.

Distance metric. To capture divergence from the ideal, we define a distance metric θ between a model’s token distribution and the baseline’s:

$$\theta(\mathbf{t}, \mathbf{b}) = 1 - \frac{\mathbf{t} \cdot \mathbf{b}}{|\mathbf{t}|, |\mathbf{b}|}, \quad (1)$$

where $\mathbf{t}$ and $\mathbf{b}$ are vectors of sub-token counts for the model and baseline, respectively. A larger θ indicates greater sub-token divergence.

Date fragmentation ratio. Building on θ , we introduce the *date fragmentation ratio* F , which quantifies how fragmented a tokeniser’s output is relative to the baseline. We initialise $F = 0.0$ for a perfectly aligned segmentation and apply downward adjustments according to observed discrepancies: a 0.10 penalty if the actual year/month/day components are fragmented (i.e., $\mathbf{1}_{\text{split}} = 1$), a 0.10 penalty if original delimiters are lost (i.e., $\mathbf{1}_{\text{delimiter}} = 1$), a 0.05 penalty multiplied by the token count difference ($N - N_b$) between a tokeniser and the baseline, and a $0.30 \times \theta$ penalty for distributional divergence. The resulting $F \in [0, 1]$ provides an interpretable score: values close to 0

Dataset and Task	# Formats	# Raw	Size	Evaluation	
				Example	GT
Context based	6	500	3000	Which team did Omid Namazi play for in 06/10/1990?	Maryland Bays
2 Date Format Switching	10	150	1500	Are 20251403 and March 14th 2025 referring to the same date?	Yes
Date Arithmetic	5	400	2000	What date is 10,000 days before 5/4/2025?	18 November 1997; 17 December 1997
Total	21	1500	6500		

Table 1: Overview and examples of task splits in DATEAUGBENCH.

denote minimal fragmentation, and values near 1 indicate severe fragmentation.

$$F = 0.10 \times \mathbf{1}_{\text{split}} + 0.10 \times \mathbf{1}_{\text{delimiter}} + 0.05 \times (N - N_b) + 0.30 \times \theta \quad (2)$$

This date fragmentation ratio is pivotal because tokenisation inconsistencies directly impair a model’s ability to represent and reason over temporal inputs. When date strings are split non-intuitively, models encounter inflated token sequences and fragmented semantic cues, which can potentially lead to errors in tasks such as chronological comparison, date arithmetic, and context-based resolution.

Validation of Date Fragmentation Ratio. To ensure our custom metric is well-founded, we performed a two-part validation. First, we conducted a human evaluation study, in which we found that our F metric’s scores align strongly with human judgments of "fragmentation severity" (Spearman’s $\rho = 0.84$), significantly outperforming standard metrics like BLEU ($\rho = 0.52$). Second, we used a data-driven approach to learn the metric’s weights by training a model to predict the human severity scores from our fragmentation components. This process confirmed that our intuitively chosen weights accurately reflect the factors driving human perception of fragmentation. For a detailed breakdown of the human evaluation protocol and the data-driven validation, please see Appendix A.3.

4.2 Temporal Reasoning Evaluation

Models. We evaluate a spectrum of model ranging from 0.5 B to 14 B parameters: five open-source Qwen 2.5 models (0.5 B, 1.5 B, 3 B, 7 B, 14 B) (Yang et al., 2024), two Llama 3 models (3 B, 8 B) (Touvron et al., 2023), and two

OLMo (Groeneveld et al., 2024) models (1 B, 7 B). For comparison with state-of-the-art closed models, we also query the proprietary GPT-4o and GPT-4o-mini endpoints via the OpenAI API (OpenAI et al., 2024).

LLM-as-a-judge. To measure how date tokenisation affects downstream reasoning, we employ an LLM-as-judge framework using GPT-4o. For each test instance in DATEFRAGBENCH, we construct a JSONL record that includes the question text, the model’s predicted answer, and a set of acceptable gold targets to capture all semantically equivalent date variants (e.g., both “03/04/2025” and “April 3, 2025” can appear in the gold label set). This record is submitted to GPT-4o via the OpenAI API with a system prompt instructing it to classify the prediction as CORRECT, INCORRECT, or NOT ATTEMPTED. A prediction is deemed CORRECT if it fully contains any one of the gold target variants without contradiction; INCORRECT if it contains factual errors relative to all gold variants; and NOT ATTEMPTED if it omits the required information. We validate GPT-4o’s reliability by randomly sampling 50 judged instances across all splits and obtaining independent annotations from four student evaluators. In 97% of cases, GPT-4o’s judgments of model answers agree with the averaged human judgments across four student evaluators, with a Cohen’s κ of 0.89 as the inter-annotator agreement, affirming the reliability of our automatic evaluation setup.

4.3 Internal Representations

Layerwise probing. We use four Qwen2.5 (Yang et al., 2024) model checkpoints (0.5B, 1.5B, 3B, and 7B parameters) to trace how temporal information is processed internally across different layers. During inference, each question is prefixed with a fixed system prompt and a chain-of-thought cue,

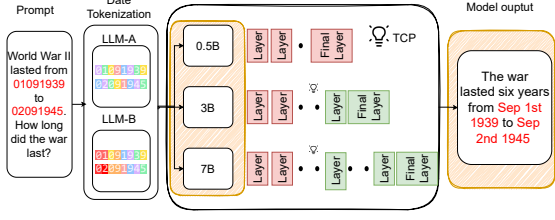


Figure 2: Illustration of how LLMs with various model sizes process dates. TCP means Tokenization Compensation Point, defined as the first layer at which LLMs achieve above-chance accuracy (see details in Sec. 6).

then passed through the model in evaluation mode. At each layer i , we extract the hidden-state vector corresponding to the final token position, yielding an embedding $h_i \in \mathbb{R}^d$ for that layer. Repeating over all examples produces a collection of layer-wise representations for positive and negative cases. We then quantify the emergence of temporal reasoning by training lightweight linear probes on these embeddings. For layer i , the probe is trained to distinguish “same-date” (positive) vs “different-date” (negative) examples. To explain when the model’s date understanding is achieved, we define the *tokenisation compensation point* as the layer at which the model’s representation correctly represents the date in the given prompt. We experiment with this idea across various model sizes, aiming to test our hypothesis: larger models would recover calendar-level semantics from fragmented tokens at earlier stages, i.e., tokenisation compensation is accomplished at early layers, as illustrated in Figure 2.

Causal attention-hop analysis. We introduce a framework intended to understand in which order date fragments are stitched together for LLMs to answer a temporal question. Figure 1 depicts the idea of our framework: given an input prompt requiring a date resolution (e.g., “Is 28052025 the same date as 28th of May 2025?”), we define two sets of tokens: (1) *concept tokens* corresponding to year, month, and day fragments, and (2) *decision tokens* corresponding to the model answer (“yes” or “no”). Our framework aims to identify a stitching path for temporal reasoning, or reasoning path for short. A reasoning path is defined as a sequence of tokens containing date fragments and the model answer¹. Given that there are multiple potential

¹The idea of reasoning paths was introduced by Lindsey et al. (2025), which we leverage to interpret how LLMs address date fragments for temporal reasoning.

paths, we score each path and select the highest-scoring one as the LLM’s reasoning path for the given prompt. To score a reasoning path, our idea is the following: we identify when a date fragment or model answer is activated, by which input token and at which layer, and then determine how important each input token is for the date fragment and model answer. Our idea is implemented by using two different approaches: (i) next token prediction (§A.2.1): how likely a date fragment and model answer follows a given input token and (ii) token importance (§A.2.2): how important an input token is to a date fragment and model answer (by replacing the input token with a random token). Lastly, we combine the results of the two approaches to yield the final score of a reasoning path (§A.2.3). This causal framework not only pinpoints *where* and *when* date fragments are activated, but also *in which order* they are stitched together to yield the model answer.

5 Experiment Results

5.1 Date fragmentation

Model	Past	Near Past	Present	Future	Avg
Baseline	0.00	0.00	0.00	0.00	0.00
OLMo	0.15	0.14	0.07	0.25	0.15
GPT-3	0.17	0.14	0.06	0.25	0.16
Llama 3	0.29	0.28	0.27	0.30	0.29
GPT-4o	0.32	0.31	0.22	0.30	0.29
GPT-3.5	0.47	0.22	0.26	0.36	0.33
GPT-4	0.36	0.26	0.29	0.39	0.33
Qwen	0.58	0.55	0.49	0.58	0.55
Gemma	0.58	0.55	0.49	0.58	0.55
DeepSeek	0.58	0.55	0.49	0.58	0.55
Llama 2	0.63	0.63	0.63	0.63	0.63
Phi	0.63	0.63	0.63	0.63	0.63

Table 2: Date fragmentation ratio across models and data splits over time. In case a family of model variants (Qwen, Gemma, DeepSeek and Phi) uses the same tokeniser, only the family name is referenced.

Cross-temporal performance. Table 2 reports the mean date fragmentation ratio across four time periods—*Past* (pre-2000), *Near Past* (2000–2009), *Present* (2010–2025), and *Future* (post-2025)—for each evaluated model. A ratio of 0.00 signifies perfect alignment with our rule-based baseline tokeniser, whereas higher values indicate progressively greater fragmentation. The rule-based Baseline unsurprisingly attains the maximal ratio of 0.00 in all periods, serving as a lower bound.

Models	Context Rlt	Fmt Switch	Date Arth.	Avg.
GPT-4o-mini	53.20	95.66	56.67	68.51
OLMo-2-7B	32.13	97.24	64.72	64.70
Qwen2.5 14B	47.56	94.56	51.35	64.49
Qwen2.5 7B	39.56	91.24	40.56	57.12
Qwen2.5 3B	25.45	90.10	39.45	51.67
LLama3.1 8B	26.20	90.22	34.50	50.31
Qwen2.5 1.5B	21.32	89.65	32.34	47.77
Qwen2.5 0.5B	10.23	88.95	31.32	43.50
OLMo-2-1B	9.26	90.09	25.90	41.75
LLama3.2 3B	9.51	88.45	23.66	40.54

Table 3: Average accuracies per task. Context Rlt stands for context based resolution, Fmt Switch refers to format switching, and Date Arth. refers to date arithmetic.

Among neural architectures, OLMo (Groeneveld et al., 2024) demonstrates the highest robustness, with an average fragmentation ratio of 0.15, closely followed by GPT-3 at 0.16. Both maintain strong fidelity across temporal splits, although performance dips modestly in the Future category (0.25), reflecting novel token sequences not seen during pre-training.

Model	Tokenised output	Frag-ratio
Baseline	10 27 1606	0.00
OLMo	10 27 16 06	0.34
Llama 3	102 716 06	0.40
GPT-3	1027 16 06	0.40
GPT-4o	102 716 06	0.40
Gemma	1 0 2 7 1 6 0 6	0.55
DeepSeek	1 0 2 7 1 6 0 6	0.55
Cohere	1 0 2 7 1 6 0 6	0.55
Qwen	1 0 2 7 1 6 0 6	0.55
Phi 3.5	– 1 0 2 7 1 6 0 6	0.60
Llama 2	– 1 0 2 7 1 6 0 6	0.60

Table 4: Tokenisation of the MMDDYYYY string “10271606” across models.

Impact of subtoken granularity. A closer look, from Table 4, at sub-token granularity further explains these trends. Llama 3 (Touvron et al., 2023) and the GPT (OpenAI et al., 2023) families typically segment each date component into three-digit sub-tokens (e.g., “202”, “504”, “03”), thus preserving the semantic unit of “MMDDYYYY” as compact pieces. OLMo (Groeneveld et al., 2024) splits the date tokens into two digit tokens (e.g., “20”, “25”). By contrast, Qwen (Yang et al., 2024) and Gemma (Team et al., 2024) models break dates into single-digit tokens (e.g., “2”, “5”), whereas Phi (Abdin et al., 2024) divides it into single-digit tokens with an initial token (e.g. “_”, “2”, “0”, “2”, “5”), inflating the token count. Although single-digit tokenisation can enhance models’ ability to perform arbitrary numeric manipulations (by treat-

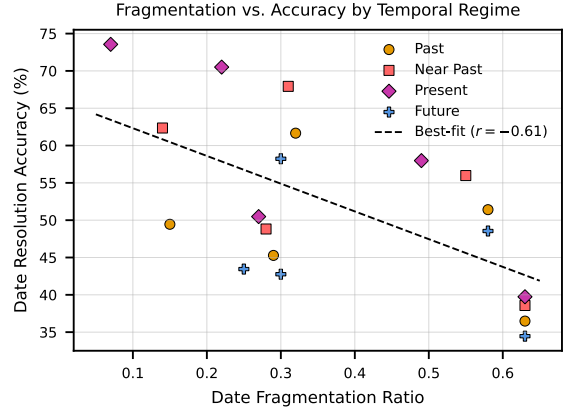


Figure 3: Date fragmentation ratio versus date resolution accuracy, stratified by four time periods and six LLMs: OLMo, Llama 3, GPT-4o, Qwen, Gemma, Phi.

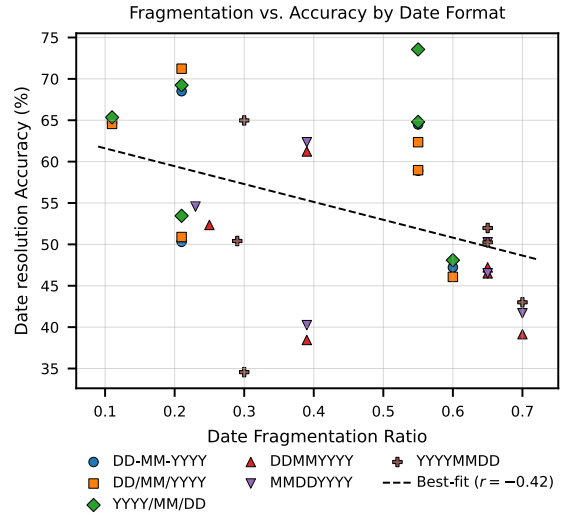


Figure 4: Date fragmentation ratio versus date resolution accuracy, stratified by six formats and six LLMs.

ing each digit as an independent unit), it comes at the expense of temporal abstraction: the tight coupling between day, month, and year is lost, inflating the compression penalty and increasing the θ divergence from the baseline.

5.2 DATEFRAGBENCH Evaluation

Performance on temporal reasoning tasks. We compare model accuracies in three tasks: Context-based Resolution, Format Switching, and Date Arithmetic (see Table 3). All models effectively solve Format Switching (e.g. 97.2% for OLMo-2-7B, 95.7% for GPT-4o-mini, 94.6% for Qwen2.5-14B, 90.2% for Llama3.1-8B). By contrast, Context Resolution and Arithmetic remain challenging: GPT-4o-mini scores 53.2% and 56.7%, Qwen2.5-

14B 47.6% and 51.4%, Llama3.1-8B 26.2% and 34.5%, and OLMo-2-7B 32.1% and 64.7%, respectively. The fact that arithmetic performance consistently exceeds resolution suggests that, given a correctly tokenised date, performing addition or subtraction is somewhat easier than resolving the date within free text—which requires encyclopedic knowledge.

Correlating date fragmentation with model accuracy over time. Figure 3 plots date fragmentation ratio against resolution accuracy, with 24 data points across six models and four temporal splits. Accuracy rises as we move from Past (1600–2000) to Near Past (2000–2009) and peaks in the Present (2010–2025), mirroring the negative correlation between fragmentation and accuracy (dashed line, Pearson correlation of -0.61). We note that the correlation is not particularly strong. This is because (i) for some models (e.g., Phi), their date fragmentation ratios remain unchanged across temporal data splits and (ii) models differ greatly by their sizes: a larger model could outperform a substantially smaller model in terms of temporal reasoning performance, even if the former has a much higher fragmentation ratio.

As seen from Table 8, GPT-4o-mini climbs from 61.7 % in Past to 67.9 % in Near Past, peaks at 70.5 % for Present, and falls to 58.2 % on Future dates. Qwen-2.5-14B and Llama-3.1-8B trace the same contour at lower absolute levels. OLMo-2-7B shows the steepest Near-Past jump ($49.5 \rightarrow 62.4$ %) and achieves the highest Present accuracy (73.6 %), consistent with its finer-grained tokenisation of “20XX” patterns. These results indicate that while finer date tokenisation (i.e., lower fragmentation ratios) boosts performance up to contemporary references, today’s models still generalise poorly to genuinely novel (post-2025) dates, highlighting an open challenge for robust temporal reasoning.

Correlating date fragmentation with model accuracy over formats. Figure 4 plots model accuracy against date fragmentation ratio across six date formats and six LLMs. A moderate negative trend emerges (dashed line, Pearson correlation of -0.42): formats that contain explicit separators (DD-MM-YYYY, DD/MM/YYYY, YYYY/MM/DD) are tokenised into more pieces and, in turn, resolved more accurately than compact, separator-free strings (DDMMYYYY, MMD-DYYYY, YYYYMMDD). As shown in Table 9, GPT-4o-mini tops every format and receives

a moderate performance drop from 71.2 % on DD/MM/YYYY to 61.2 % on DDMMYYYY, with the highest overall average (66.3 %). OLMo-2-7B and Qwen-2.5-14B both exceed 70 % on the highly fragmented YYYY/MM/DD form, but slip into the low 50s on MMDDYYYY and YYYYMMDD. Lower date fragmentation ratio models, such as Llama-3.1-8B and Phi-3.5, lag behind; their accuracy plunges below 40 %. Even so, all models score much better on separator-rich formats compared to the date formats without separators. In summary, model accuracy is correlated to how cleanly a model can tokenise the string into interpretable tokens: more visual structure (slashes or dashes) means lower fragmentation, which suggests more straightforward reasoning, and in turn, leads to better performance.

6 In which layer do LLMs compensate for date fragmentation?

Layerwise linear probing. To pinpoint in which layer a model learns to recognize two equivalent dates, we define the *tokenisation compensation point* (TCP) as the earliest layer at which a lightweight linear probe on the hidden state achieves above-chance accuracy, which is defined as 80%, on the date equivalence task. Figure 5a reports TCPs for the DATES_PAST benchmark (1600–2010): Qwen2.5-0.5B reaches TCP at layer 12 (50% depth), Qwen2.5-1.5B at layer 15 (53.6%), Qwen2.5-3B at layer 8 (22.2%), and Qwen2.5-7B at layer 4 (14.3%). The leftward shift of the 3B and 7B curves suggests how larger models recover calendar-level semantics from fragmented tokens more rapidly. Figure 5b shows the DATES_PRESENT benchmark (2010–2025), where only the 1.5B, 3B, and 7B models surpass TCP—at layers 16 (57.1%), 21 (58.3%), and 17 (60.7%), respectively—while the 0.5B model never does. The deeper TCPs here reflect extra layers needed to recombine the two-digit “20” prefix, which is fragmented unevenly by the tokeniser. In Figure 10, we evaluate DATES_FUTURE (2025–2599), where novel four-digit sequences exacerbate fragmentation. Remarkably, TCPs mirror the Past regime: layers 12, 15, 8, and 4 for the 0.5B, 1.5B, 3B, and 7B models, respectively. This parallelism indicates that model scale dictates how quickly LLMs can compensate for date fragmentation to achieve high accuracy, even when dates are novel.

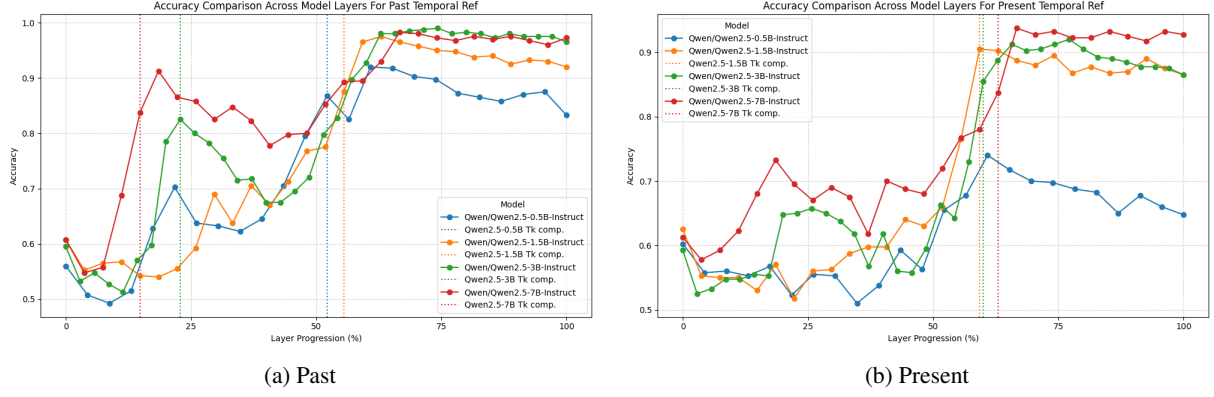


Figure 5: Layer-wise accuracies in the two time periods: Past and Present.

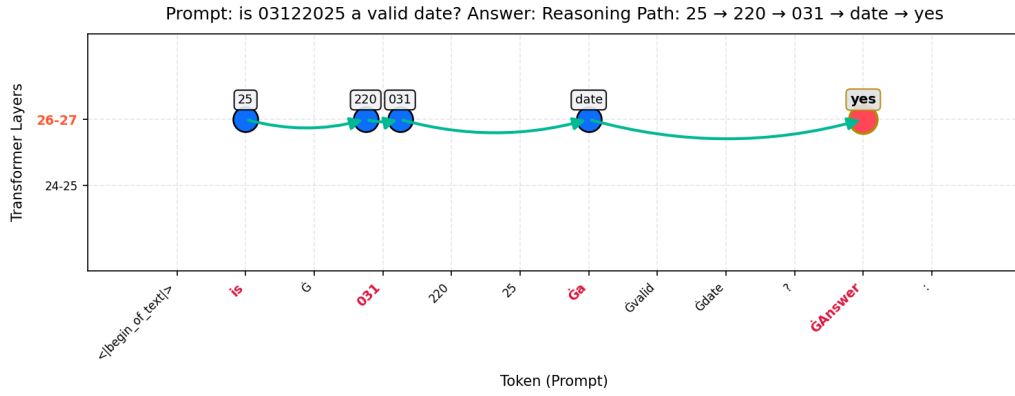


Figure 6: Reasoning path for the "03122025 is a valid date" prompt.

Tokenisation compensation point. Overall, we observe a sharp decline in TCP as model size increases: small models defer date reconstruction to middle layers, whereas the largest model does so within the first quarter of layers. Across all the three temporal benchmarks, TCP shifts steadily toward the first layers as model size grows.

7 How do LLMs stitch date fragments for temporal reasoning?

Causal path tracing. To investigate how LLMs like Llama 3 (Touvron et al., 2023) internally stitch date fragments to yield a model answer, we apply our causal framework to identify the model’s reasoning path over a specific prompt. Figure 6 plots model layers on the y axis against prompt tokens (e.g., Is 03122025 a valid date?) on the x axis. Green arrows mark the reasoning path with the highest score that is responsible for generating the answer “yes”. Date fragments “25”, “220”, “031”, and the model answer “yes” are activated in sequence at layer 26-27 by the input tokens “is”, “031”, “a” and “Answer” respectively. As such, the model performs a kind of discrete, step-by-step to-

ken aggregation, stitching together substrings of the input until a binary valid/invalid verdict emerges.

Misalignment between LLMs and human. In contrast, human readers parse dates by immediately mapping each component to a coherent temporal schema: “03” is March, “12” is day of month, “2025” is year, and then checking whether the day falls within the calendar bounds of that month. Humans bring rich world knowledge of calendars and leap-year rules to bear in parallel. However, LLMs exhibit no explicit calendar “module”; instead, they rely on learned statistical associations between digit-patterns and the training-time supervisory signal for “valid date”. The reasoning path in Figure 6 thus illustrates a fundamentally different mechanism of date comprehension in LLMs, based on date fragments re-routing rather than holistic semantic interpretation. We repeated causal tracing on 100 date strings in 6 different date formats to test whether the reasoning path difference between human and LLMs is consistent across date formats. In most of cases, we observe that model reasoning paths are not aligned with human inter-

pretation (year $\rightarrow$ month $\rightarrow$ day), rather rely on *sub-word fragments* that statistically represent year, month, and day, and stitch these date fragments in a flexible order that is subject to date formats (see examples in Figures 7-8). However, such a reasoning path becomes tricky when a date is greatly fragmented: given the date abstraction is learned from frequency rather than hard-coded rules, the abstraction is biased toward standard Western formats and contemporary years. As a result, a model often addresses popular dates (in the same format) with similar reasoning paths. However, the reasoning path becomes obscure on rare, historical, or locale-specific strings outside the distribution of pre-training data (see Figure 9).

8 Discussion

The moderate Pearson correlations of -0.61 (by temporal regime) and -0.42 (by date format) are a significant finding in themselves. They confirm that date fragmentation is a consistent and independent bottleneck, while also highlighting that it is not the sole factor influencing performance. The remaining variance is naturally explained by confounding factors such as model architecture, scale, and pre-training data exposure. For instance, a larger model may have a greater capacity to "heal" a poorly tokenised date, partially masking the negative impact of a high fragmentation ratio. Nonetheless, our results demonstrate that, all else being equal, higher fragmentation consistently predicts a drop in accuracy. This reveals a fundamental impediment that exists at the input level, before the model's core reasoning layers are even engaged.

Our findings also shed light on the debate between memorisation and actual logical reasoning in LLMs. The performance disparity between "Present" dates and "Past" or "Future" dates (Figure 3) suggests that models heavily rely on statistical patterns and memorised facts from their pre-training data. For common contemporary dates, strong learned associations allow models to effectively parse and reason about them, even if tokenisation is suboptimal. However, for less frequent historical or novel future dates, this reliance on memorisation becomes a liability. The "date abstraction" mechanism struggles, and models must generalise from sparser data, leading to the observed accuracy drops. This contrasts with human date parsing, which leverages explicit, rule-based calendar knowledge rather than frequency-based

recall.

9 Conclusion

In this paper, we identified date tokenisation as a critical yet overlooked bottleneck in temporal reasoning with LLMs. We demonstrated a correlation between date fragmentation and task performance in temporal reasoning, i.e., the more fragmented the tokenisation, the worse the reasoning performance. Our layerwise and causal analyses in LLMs further revealed an emergent "date abstraction" mechanism that explains when and how LLMs understand and interpret dates. Our results showed that larger models can compensate for date fragmentation at early layers by stitching fragments for temporal reasoning, while the stitching process appears to follow a reasoning path that connects date fragments in a flexible order, differing from human interpretation from year to month to day.

Limitations

While our work demonstrates the impact of date tokenisation on LLMs for temporal reasoning, there are several limitations. First, DATEAUGBENCH focuses on a finite set of canonical date serialisations and does not capture the full diversity of natural-language expressions (e.g., "the first Monday of May 2025") or noisy real-world inputs like OCR outputs. Second, our experiments evaluate a representative but limited pool of tokenisers and model checkpoints (up to 14B parameters); therefore, the generalizability of date fragmentation ratio and our probing and causal analyses to very large models with 15B+ parameters remains unknown. Third, while the fragmentation ratio measures front-end segmentation fidelity, it does not account for deeper world-knowledge factors such as leap-year rules, timezone conversions, and culturally grounded calendar systems, all of which may influence temporal interpretation; further, the fragmentation ratio metric, though straightforward and interpretable, is not rigorously evaluated. Our work and the DATEAUGBENCH benchmark deliberately focus on a specific, foundational challenge: the tokenisation of explicit, multi-digit date strings in Anglo-centric Gregorian formats. This narrow scope allows us to isolate the impact of subword fragmentation on core temporal reasoning. However, this focus means we do not address the full complexity of temporal understanding, such as dates expressed in natural language (e.g., "the first Monday of May 2025"), dates

with missing components, non-Gregorian calendars (e.g., Hijri, Hebrew), or dates represented with non-Latin numeral systems. We consider the extension to these diverse and important cases as critical future work that can build upon our foundational analysis of fragmentation. Lastly, the core idea of our causal framework is inspired by [Lindsey et al. \(2025\)](#); however, our extension to temporal reasoning is not evaluated. Future work should extend to more diverse date expressions, broader model and tokeniser families, equipping tokenisers with external calendar-wise knowledge to improve further robust temporal reasoning, and conducting rigorous evaluation of the fragmentation ratio metric and the causal framework.

Ethical Considerations

DATEAUGBENCH is derived solely from the public, research-licensed TIMEQA and TIMEBENCH corpora that do not contain sensitive data; our augmentation pipeline rewrites only date strings. However, our dataset focuses on 21 Anglo-centric Gregorian formats. Therefore, our data potentially reinforce a Western default and overlook calendars or numeral systems used in many other cultures, and our date fragmentation metric may over-penalise tokenisers optimised for non-Latin digits.

Acknowledgements

We thank the anonymous reviewers for their thoughtful comments that greatly improved the work. We gratefully thank Madiha Kazi, Cristina Mahanta, and MingZe Tang for their support of conducting human evaluation for LLMs-as-judge. We also thank Ahmad Isa Muhammad for participating in early discussions.

References

- Marah Abdin, Jyoti Aneja, Hany Awadalla, Ahmed Awadallah, Ammar Ahmad Awan, Nguyen Bach, Amit Bahree, Arash Bakhtiari, Jianmin Bao, Harkirat Behl, Alon Benhaim, Misha Bilenko, Johan Bjorck, Sébastien Bubeck, Martin Cai, Qin Cai, Vishrav Chaudhary, Dong Chen, Dongdong Chen, and 110 others. 2024. [Phi-3 technical report: A highly capable language model locally on your phone](#). 2404.14219v4.
- Gagan Bhatia, El Moatez Billah Nagoudi, Hasan Cavusoglu, and Muhammad Abdul-Mageed. 2024. [Fintral: A family of gpt-4 level multimodal financial large language models](#). *Preprint*, arXiv:2402.10986.
- Andrea Carriero, Davide Pettenuzzo, and Shubhran-shu Shekhar. 2024. [Macroeconomic forecasting with large language models](#). *arXiv preprint arXiv:2407.00890*.
- Ching Chang, Wei-Yao Wang, Wen-Chih Peng, and Tien-Fu Chen. 2023. [Llm4ts: Aligning pre-trained llms as data-efficient time-series forecasters](#). *arXiv preprint arXiv:2308.08469*.
- Wenhu Chen, Xinyi Wang, and William Yang Wang. 2021. [A dataset for answering time-sensitive questions](#). *Preprint*, arXiv:2108.06314.
- Zheng Chu, Jingchang Chen, Qianglong Chen, Weijiang Yu, Haotian Wang, Ming Liu, and Bing Qin. 2024. [Timebench: A comprehensive evaluation of temporal reasoning abilities in large language models](#). *Preprint*, arXiv:2311.17667.
- Bahare Fatemi, Mehran Kazemi, Anton Tsitsulin, Karishma Malkan, Jinyeong Yim, John Palowitch, Sungyong Seo, Jonathan Halcrow, and Bryan Perozzi. 2024. [Test of time: A benchmark for evaluating llms on temporal reasoning](#). 2406.09170v1.
- Juan Luis Gastaldi, John Terilla, Luca Malagutti, Brian DuSell, Tim Vieira, and Ryan Cotterell. 2024. [The foundations of tokenization: Statistical and computational concerns](#). 2407.11606v3.
- Omer Goldman, Avi Caciularu, Matan Eyal, Kris Cao, Idan Szpektor, and Reut Tsarfaty. 2024. [Unpacking tokenization: Evaluating text compression and its correlation with model performance](#). 2403.06265v2.
- Dirk Groeneveld, Iz Beltagy, Pete Walsh, Akshita Bhagia, Rodney Kinney, Oyvind Tafjord, Ananya Harsh Jha, Hamish Ivison, Ian Magnusson, Yizhong Wang, Shane Arora, David Atkinson, Russell Authur, Khyathi Raghavi Chandu, Arman Cohan, Jennifer Dumas, Yanai Elazar, Yuling Gu, Jack Hessel, and 24 others. 2024. [Olmo: Accelerating the science of language models](#). 2402.00838v4.
- Jack Lindsey, Wes Gurnee, Emmanuel Ameisen, Brian Chen, Adam Pearce, Nicholas L. Turner, Craig Citro, David Abrahams, Shan Carter, Basil Hosmer, Jonathan Marcus, Michael Sklar, Adly Templeton, Trenton Bricken, Callum McDougall, Hoagy Cunningham, Thomas Henighan, Adam Jermy, Andy Jones, and 8 others. 2025. [On the biology of a large language model](#). *Transformer Circuits Thread*.
- Alisa Liu, Jonathan Hayase, Valentin Hofmann, Se-wong Oh, Noah A. Smith, and Yejin Choi. 2025. [Superbpe: Space travel for language models](#). *arXiv preprint arXiv:2503.13423*.
- OpenAI, :, Aaron Hurst, Adam Lerer, Adam P. Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, Aleksander Mądry, Alex Baker-Whitcomb, Alex Beutel, Alex Borzunov, Alex Carney, Alex Chow, Alex Kirillov, and 401 others. 2024. [Gpt-4o system card](#). 2410.21276v1.

- OpenAI, Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, Red Avila, Igor Babuschkin, Suchir Balaji, Valerie Balcom, Paul Baltescu, Haiming Bao, Mohammad Bavarian, Jeff Belgum, and 262 others. 2023. [Gpt-4 technical report](#). 2303.08774v6.
- Buu Phan, Marton Havasi, Matthew Muckley, and Karen Ullrich. 2024. Understanding and mitigating tokenization bias in language models. *arXiv preprint arXiv:2406.16829*.
- Nived Rajaraman, Jiantao Jiao, and Kannan Ramchandran. 2024. [Toward a theory of tokenization in llms](#). 2404.08335v1.
- François Remy, Pieter Delobelle, Hayastan Avetisyan, Alfiya Khabibullina, Miryam de Lhoneux, and Thomas Demeester. 2024. Trans-tokenization and cross-lingual vocabulary transfers: Language adaptation of llms for low-resource nlp. *arXiv preprint arXiv:2408.04303*.
- Craig W. Schmidt, Varshini Reddy, Haoran Zhang, Alec Alameddine, Omri Uzan, Yuval Pinter, and Chris Tanner. 2024. [Tokenization is more than compression](#). 2402.18376v2.
- Aaditya K. Singh and DJ Strouse. 2024. [Tokenization counts: the impact of tokenization on arithmetic in frontier llms](#). 2402.14903v1.
- Zhaochen Su, Jun Zhang, Tong Zhu, Xiaoye Qu, Juntao Li, Min Zhang, and Yu Cheng. 2024. [Timo: Towards better temporal reasoning for language models](#). 2406.14192v2.
- Mingtian Tan, Mike A. Merrill, Vinayak Gupta, Tim Althoff, and Thomas Hartvigsen. 2024. Are language models actually useful for time series forecasting? In *Advances in Neural Information Processing Systems*.
- Qingyu Tan, Hwee Tou Ng, and Lidong Bing. 2023. [Towards robust temporal reasoning of large language models via a multi-hop qa dataset and pseudo-instruction tuning](#). 2311.09821v2.
- Gemma Team, Morgane Riviere, Shreya Pathak, Pier Giuseppe Sessa, Cassidy Hardin, Surya Bhupatiraju, Léonard Hussenot, Thomas Mesnard, Bobak Shahriari, Alexandre Ramé, Johan Ferret, Peter Liu, Pouya Tafti, Abe Friesen, Michelle Casbon, Sabela Ramos, Ravin Kumar, Charline Le Lan, Sammy Jerome, and 179 others. 2024. [Gemma 2: Improving open language models at a practical size](#). *Preprint*, arXiv:2408.00118.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, Dan Bikel, Lukas Blecher, Cristian Canton Ferrer, Moya Chen, Guillem Cucurull, David Esiobu, Jude Fernandes, Jeremy Fu, Wenyin Fu, and 49 others. 2023. [Llama 2: Open foundation and fine-tuned chat models](#). 2307.09288v2.
- Stylianios Loukas Vasileiou and William Yeoh. 2024. Trace-cs: A synergistic approach to explainable course scheduling using llms and logic. *arXiv preprint arXiv:2409.03671*.
- Jiapu Wang, Kai Sun, Linhao Luo, Wei Wei, Yongli Hu, Alan Wee-Chung Liew, Shirui Pan, and Baocai Yin. 2024. Large language models-guided dynamic adaptation for temporal knowledge graph reasoning. *arXiv preprint arXiv:2405.14170*.
- Yang Wang and Hassan A Karimi. 2024. Exploring large language models for climate forecasting. *arXiv preprint arXiv:2411.13724*.
- Jason Wei, Nguyen Karina, Hyung Won Chung, Yunxin Joy Jiao, Spencer Papay, Amelia Glaese, John Schulman, and William Fedus. 2024. [Measuring short-form factuality in large language models](#). *Preprint*, arXiv:2411.04368.
- Yifan Wei, Yisong Su, Huanhuan Ma, Xiaoyan Yu, Fangyu Lei, Yuanzhe Zhang, Jun Zhao, and Kang Liu. 2023. [Menatqa: A new dataset for testing the temporal comprehension and reasoning abilities of large language models](#). 2310.05157v1.
- Siheng Xiong, Ali Payani, Ramana Kompella, and Faramarz Fekri. 2024. [Large language models can learn temporal reasoning](#). 2401.06853v6.
- An Yang, Baosong Yang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Zhou, Chengpeng Li, Chengyuan Li, Dayiheng Liu, Fei Huang, Guanting Dong, Haoran Wei, Huan Lin, Jialong Tang, Jialin Wang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Ma, and 43 others. 2024. [Qwen2 technical report](#). 2407.10671v4.
- Yifan Zeng. 2024. Histolens: An llm-powered framework for multi-layered analysis of historical texts – a case application of yantie lun. *arXiv preprint arXiv:2411.09978*.
- Bowen Zhao, Zander Brumbaugh, Yizhong Wang, Hananeh Hajishirzi, and Noah A. Smith. 2024. [Set the clock: Temporal alignment of pretrained language models](#). 2402.16797v2.
- Mengyu Zheng, Hanting Chen, Tianyu Guo, Chong Zhu, Binfan Zheng, Chang Xu, and Yunhe Wang. 2024. Enhancing large language models through adaptive tokenizers. In *Proc. NeurIPS*.
- Zhejiang Zhou, Jiayu Wang, Dahua Lin, and Kai Chen. 2024. [Scaling behavior for large language models regarding numeral systems: An example using pythia](#). 2409.17391v2.

A Appendix

A.1 Experiment Design

Implementation details of evaluation. The evaluation pipeline is implemented in Python and supports asynchronous API requests with retry logic,

as well as multiprocessing to handle thousands of examples efficiently. After collecting GPT-4o’s label for each instance, we map CORRECT/INCORRECT NOT ATTEMPTED to categorical scores A, B, and C. We then compute three core metrics: overall accuracy (proportion of A scores), given-attempted accuracy (A over A+B), and the F1 score, defined as the harmonic mean of overall and given-attempted accuracy. Results are reported both globally and stratified by task split (Context-based, Format Switching, Date Arithmetic) and by temporal category (Past, Near Past, Present, Future). We adopt the sample prompts introduced in SimpleQA (Wei et al., 2024) as our LLM-as-judge queries, ensuring consistent scoring instructions across all evaluations. Our specific prompt used for evaluation can be found in Table 10. We have presented our examples of LLM as a judge and human evaluation in Table 11.

Date ambiguities. We explicitly enumerate all valid variants in the gold label set for each example to handle multiple correct answers arising from date-format ambiguities. This ensures that any prediction matching one of these variants is marked correct, avoiding penalisation for format differences.

Synthetic benchmark construction for linear probing. We construct a suite of synthetic true-false benchmarks to isolate temporal reasoning across different reference frames. For the DATES_PAST, DATES_PRESENT, and DATES_FUTURE datasets, we sample 1,000 date-date pairs each, drawing calendar dates uniformly from the appropriate range and rendering them in two randomly chosen, distinct formatting patterns (Ymd vs d/m/Y). Exactly half of each set are “YES” examples (identical dates under different formats), which are our positive examples, and half are “NO” (different dates), which are our negative examples. All three datasets are balanced, shuffled, and split into equal positive and negative subsets to ensure fair probing.

A.2 Causal Attention–Hop Analysis

A.2.1 Next Token Prediction

We treat each token in the prompt as a candidate “concept” to follow. After the model processes the input, it produces a hidden vector $h_{\ell,p}$ per token at position p and layer ℓ . To see how likely a concept c (e.g., a date fragment and model answer) follows

each input token, we project $h_{\ell,p}$ through W_U to yield the “probability” distribution of vocabulary tokens, and denote $s_{\ell,p}^c$ as the “probability” of the concept being the next token.

$$z_{\ell,p} = W_U h_{\ell,p}, \quad s_{\ell,p}^c = z_{\ell,p}[t_c], \quad (1)$$

where t_c is the index of concept c in the vocabulary.

A.2.2 Token Importance

To measure how important an input token is to a concept (e.g., a date fragment and model answer), we replace the token with an unrelated one (e.g., “Dallas” $\rightarrow$ “Chicago”) and compute the probability drop of the concept incurred by the replacement, denoted as $I_{c,p}$ (which we compute only at the last layer):

$$I_{c,p} = \sigma(z_p[t_c]) - \sigma(\tilde{z}_p[t_c]), \quad (2)$$

where σ is a softmax function. The bigger the $I_{c,p}$, the more important the original token at position p for the concept c .

A.2.3 Path scoring

A *reasoning path* $\mathcal{P} = (c_1, \dots, c_k)$ is a sequence of tokens, indicating in which order date fragments are stitched together for LLMs to answer a temporal question. We score each potential path by blending five components (ordering, activation strength, causal strength, gap penalty, and confidence in the final concept), into a single score:

$$S(\mathcal{P}) = \alpha \times S_{\text{order}} + \beta \times S_{\text{act}} + \gamma \times S_{\text{causal}} - \eta \times S_{\text{gap}} + \kappa \times S_{\text{final}} \quad (1)$$

Each term is designed to reward a different desirable property:

- **Ordering:** we give points if the concepts appear in roughly left-to-right order in the prompt, and secondarily in increasing layer order:

$$S_{\text{order}} = 0.7 \times \mathbf{1}[p_1 \leq \dots \leq p_k] + 0.3 \times \mathbf{1}[\ell_1 \leq \dots \leq \ell_k], \quad (2)$$

where $\mathbf{1}$ is an indicator function, $p_i = \max_{\ell,p} s_{\ell,p}^{c_i}$ indicating the position of the most important input token for a concept c_i at the last layer. Similarly, ℓ_i is the layer at which an input token pays the most attention to the concept c_i .

- **Activation:** we compute the average position of the most important input token for a concept from 1 to k , and normalize by a threshold $\tau = 0.2$, and clip to 1:

$$S_{\text{act}} = \min\left(\frac{1}{k} \sum_{i=1}^k p_i / \tau, 1\right), \quad (3)$$

- **Causal strength:** we use the token importance score, denoted as $d_i = |I_{c_{i+1}, p_i}|$ between two adjacent concepts c_{i+1} and c_i , up-weight latter scores, and downweight missing links by a coverage term ρ , which is defined as the fraction of actual causal connections observed between consecutive concepts out of the total possible consecutive pairs in the path. The combined score then multiplies the weighted average of the d_i by $\frac{1}{2} + \frac{1}{2}\rho$, giving:

$$S_{\text{causal}} = \left(\frac{\sum_i w_i d_i}{\sum_i w_i}\right) (0.5 + 0.5\rho), \quad (4)$$

where $w_i = 0.5 + 0.5 \frac{i-1}{k-2}$.

- **Gap penalty:** to discourage large jumps in position, we compute the mean gap $\bar{g}$ and apply a small multiplier $\lambda = 0.1$:

$$S_{\text{gap}} = 1 - \lambda \bar{g}, \quad S_{\text{gap}} \leq 1. \quad (6)$$

This is done to encourage model paths to think step by step instead of directly jumping to the conclusion (yes/no).

- **Final confidence:** We compute the position of the most important input token for the last concept c_k :

$$S_{\text{final}} = \max_{\ell, p} s_{\ell, p}^{c_k}. \quad (7)$$

The reasoning path with the highest total score $S(\mathcal{P})$ is chosen as the model’s reasoning path over a specific prompt. We note that *Ordering*, *Activation*, *Gap penalty* and *Final confidence* components are built upon next token prediction signals $s_{\ell, p}^c$, whereas the *Causal strength* component is derived solely from token importance score I_{c_{i+1}, π_i} , i.e. the drop in the softmax probability for concept c_{i+1} when the token at position p_i is replaced.

A.3 Detailed Validation of the Date Fragmentation Ratio

This appendix provides a detailed account of the two-part validation process for our custom date fragmentation ratio (F), demonstrating its alignment with human intuition and its empirical soundness.

A.3.1 Human Evaluation of Fragmentation Severity

This study was designed to confirm that our F metric captures what humans perceive as semantic disruption in tokenized dates more effectively than general-purpose text similarity metrics.

Methodology. We recruited five computer science graduate students, who were familiar with NLP but blind to our hypotheses, to serve as annotators. We created a stimulus set of 100 tokenised date strings, stratified to represent a wide range of models, date formats, and fragmentation levels from our experiments. For each item, annotators were shown the original date and the list of sub-tokens, and asked to rate the “**fragmentation severity**” on a 5-point Likert scale, according to the following rubric with examples:

- **1 (No Fragmentation):** Tokens perfectly preserve the semantic components. *Example:* ‘10-27-1606’ $\rightarrow$ [‘10’, ‘-’, ‘27’, ‘-’, ‘1606’]
- **2 (Minor Fragmentation):** Mostly preserved, with minor, non-ideal splits. *Example:* ‘1606’ $\rightarrow$ [‘16’, ‘06’]
- **3 (Moderate Fragmentation):** Core components are broken, making the structure harder to discern. Delimiters might be lost or numbers oddly grouped. *Example:* ‘10271606’ $\rightarrow$ [‘102’, ‘716’, ‘06’]
- **4 (High Fragmentation):** Date split into many small pieces (e.g., single digits), though the original characters are easily reassembled. *Example:* ‘1606’ $\rightarrow$ [‘1’, ‘6’, ‘0’, ‘6’]
- **5 (Severe Fragmentation):** Tokenization completely obscures the date’s structure, often by adding non-numeric tokens or creating highly unintuitive groupings. *Example:* ‘1606’ $\rightarrow$ [‘_’, ‘1’, ‘6’, ‘0’, ‘6’]

The human judgments were highly reliable, with a Krippendorff’s Alpha for inter-annotator agreement of $\alpha = 0.81$.

Results. We computed the Spearman’s rank correlation coefficient (ρ) between the average human rating for each item and the scores from our F metric, BLEU, and character-level Edit Distance. As shown in Table 6, our F metric demonstrated a

Model	Tokenised output	Frag-ratio	Avg. Human Severity Rating (1-5)
Baseline	10 27 1606	0.00	1.0
OLMo	10 27 16 06	0.34	2.0
Llama 3	102 716 06	0.40	3.4
GPT-4o	102 716 06	0.40	3.4
Cohere	1 0 2 7 1 6 0 6	0.55	4.6
Phi 3.5	_ 1 0 2 7 1 6 0 6	0.60	5.0
Llama 2	_ 1 0 2 7 1 6 0 6	0.60	5.0

Table 5: An illustrative example showing the strong correlation between the calculated fragmentation ratio (Frag-ratio) and the average human-perceived severity rating for the tokenisation of the MMDDYYYY string "10271606". Higher scores in both metrics indicate greater fragmentation.

Metric	Correlation with Human Ratings (ρ)
Date Fragmentation Ratio (F)	0.84
BLEU Score	0.52
Character-Level Edit Distance	0.21

Table 6: Spearman Correlation (ρ) of Metrics with Human Judgments of Fragmentation Severity.

strong correlation with human ratings, far exceeding the general-purpose metrics. In Table 5, we present examples from the human validation study.

This result confirms that our specialised metric is necessary and effective, as it successfully quantifies the semantic disruption that humans perceive but that generic text metrics fail to capture.

A.3.2 Data-Driven Validation of Metric Coefficients

This analysis provides an empirical grounding for the weights used in our F metric’s formula. By learning the weights from data, we can validate that our initial, intuitive design aligns with a more formal, data-driven approach.

Formal Problem Formulation. To directly tune our metric to align with human perception, we frame the task as a linear regression problem where the goal is to predict the average human severity rating. This setup is more straightforward for validating the metric’s components against human judgment.

- The target variable is the **Average Human Severity Rating**, a continuous score from 1 to 5, as described in Appendix A.3.1.
- The feature vector $\mathbf{x} \in \mathbb{R}^4$ consists of the four fragmentation components: $\mathbf{x} = [1_{\text{split}}, 1_{\text{delimiter}}, (N - N_b), \theta]$.
- We aim to learn a weight vector $\mathbf{w}$ such that: Avg. Human Severity Rating $\approx \mathbf{w}^T \mathbf{x} + \text{intercept}$.

We used a non-negative linear regression model, as each fragmentation component is hypothesised to increase, not decrease, the perceived severity. Features were standardised before training to ensure the learned coefficients were comparable.

Results and Confirmation. After fitting the model to our human evaluation data, we obtained a set of empirically derived coefficients. We normalised these weights to sum to 1 to compare them with the relative importance implied by our original formula. As Table 7 shows, the weights learned by predicting human ratings are remarkably similar to the normalised version of our original, intuitively set weights.

This result provides strong empirical validation of our F metric’s design from an alternative perspective. It demonstrates that our initial weights, chosen based on semantic principles, accurately reflect not only the impact on model performance but also the factors that drive human perception of the severity of fragmentation. The relative importance of the components remains consistent: distributional divergence (θ) is the most significant factor, followed by major structural breaks (splits and delimiter loss), and finally by token count inflation. This confirms the robustness and validity of our metric’s formulation.

Fragmentation Component	Original Intuitive Weight (Normalized)	Empirically Learned Weight (from Human Ratings)
l_{split} (Component Split)	$0.10/0.55 \approx \mathbf{0.1818}$	0.2015
$l_{\text{delimiter}}$ (Delimiter Loss)	$0.10/0.55 \approx \mathbf{0.1818}$	0.1932
$N - N_b$ (Token Difference)	$0.05/0.55 \approx \mathbf{0.0909}$	0.1053
θ (Distributional Divergence)	$0.30/0.55 \approx \mathbf{0.5455}$	0.5000

Table 7: Comparison of Original (Normalised) and Empirically Learned Weights for the F Metric, using human ratings as the target variable.

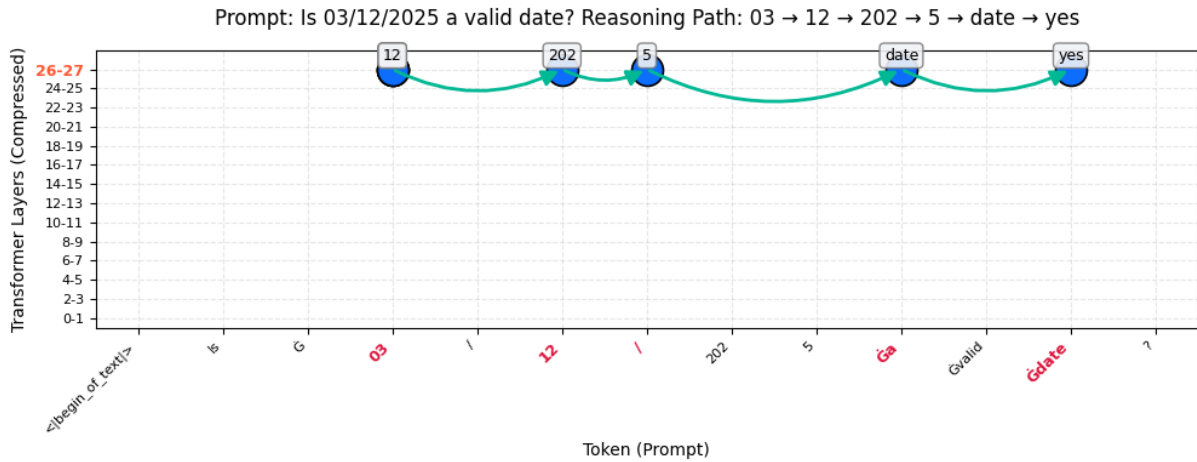


Figure 7: Reasoning path for the “03/12/2025 is a valid date” prompt.

Models	Past	Near Past	Present	Future
GPT-4o-mini	61.66	67.93	70.51	58.23
OLMo-2-7B	49.45	62.35	73.56	43.45
Qwen2.5 14B	58.97	64.80	67.22	55.69
Qwen2.5 7B	51.41	55.98	57.98	48.55
Qwen2.5 3B	46.50	50.25	51.98	43.91
LLama3.1 8B	45.28	48.82	50.48	42.76
Qwen2.5 1.5B	42.99	46.16	47.69	40.60
Qwen2.5 0.5B	39.15	41.68	43.00	36.98
OLMo-2-1B	36.07	38.09	40.49	34.07
LLama3.2 3B	36.48	38.57	39.74	34.46

Table 8: Model accuracy on context-based resolution across four data splits over time.

Model	DD-MM-YYYY	DD/MM/YYYY	YYYY/MM/DD	DDMMYYYY	MMDDYYYY	YYYYMMDD	Avg.
OLMo	64.70	64.56	65.35	52.35	54.56	50.41	58.65
Llama 3	50.31	50.89	53.45	38.45	40.24	34.56	44.65
GPT-4o	68.51	71.23	69.24	61.23	62.34	64.98	66.25
Qwen	64.49	62.35	73.56	46.50	50.25	51.98	58.19
Gemma	58.90	58.97	64.80	47.22	46.50	50.25	54.44
Phi	47.23	46.07	48.09	39.15	41.68	43.00	44.20

Table 9: Model accuracy on context-based resolution across date formats.

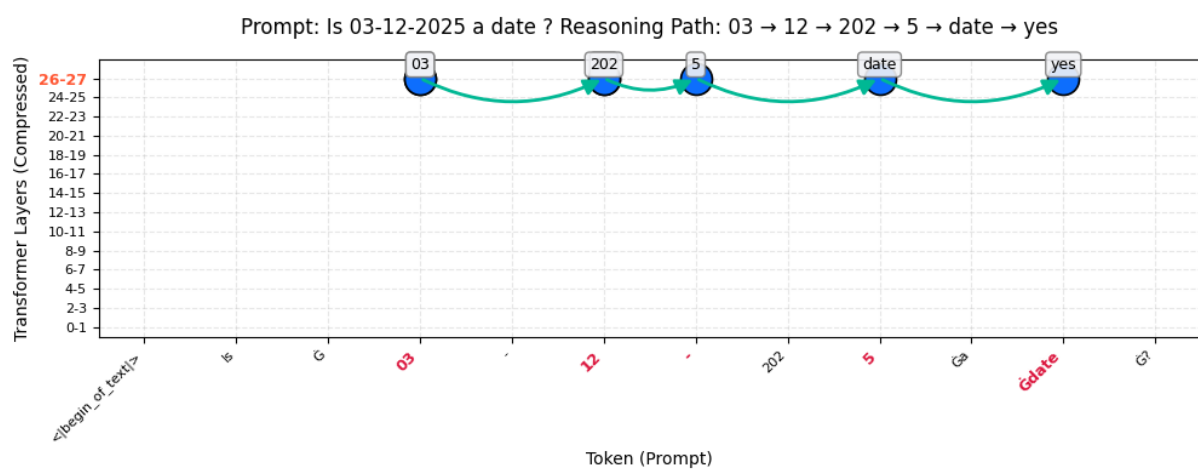


Figure 8: Reasoning path of the “03-12-2025 is a valid date” prompt.

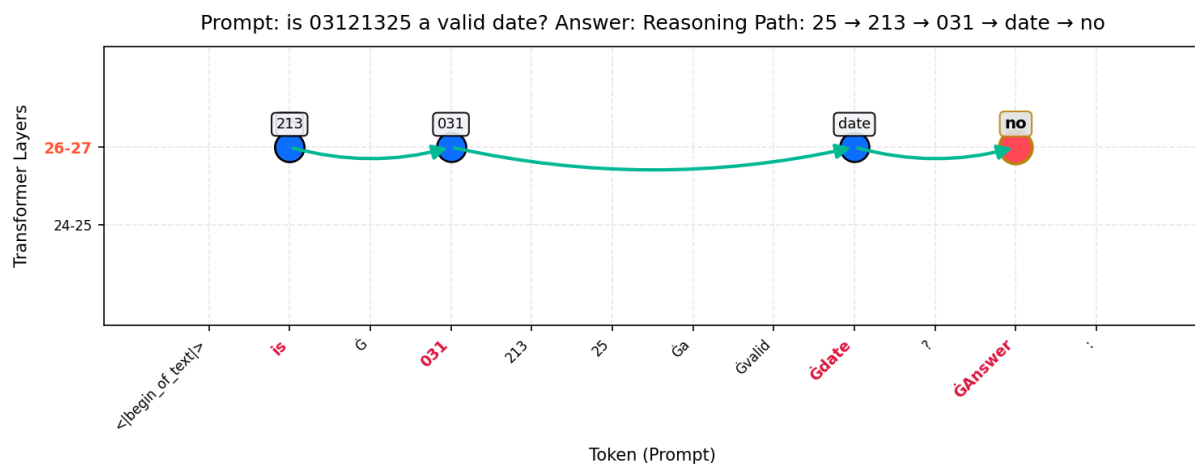


Figure 9: Reasoning path of the “03121325 is a valid date” prompt, where year = 1325.

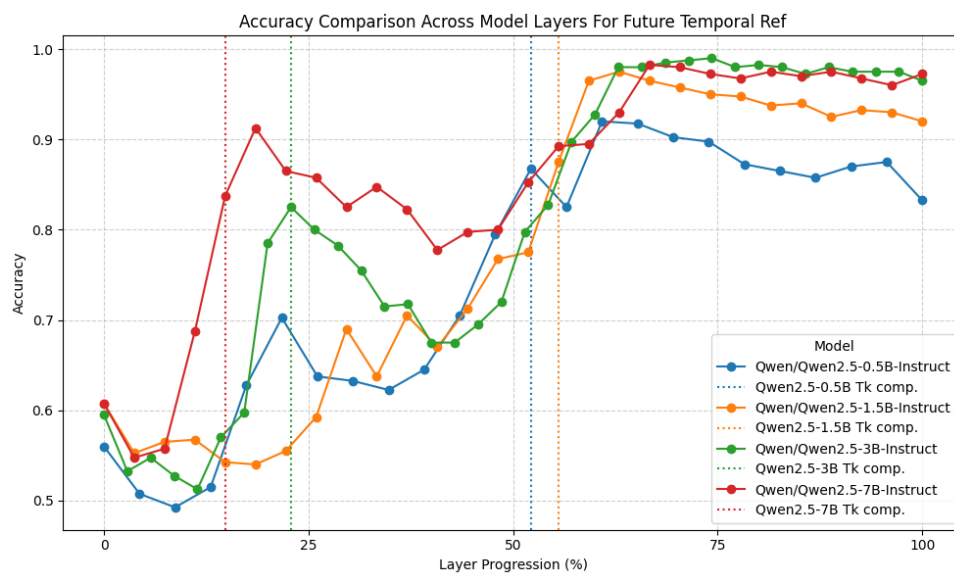


Figure 10: Layer-wise accuracies in the Future period.

LLM-as-Judge Evaluation Prompt

Your task: Evaluate one prediction at a time. You receive:

- **Question** – the task prompt shown to the model
- **Gold target** – *all* answers that are considered correct
- **Predicted answer** – the model’s response

Return **one letter only**:

A	CORRECT	prediction fully matches <i>one</i> gold variant
B	INCORRECT	prediction contradicts or misses required info
C	NOT_ATTEMPTED	prediction refuses, guesses, or answers irrelevantly

General rules:

1. Match semantics, ignore capitalisation, punctuation, order.
2. If any statement contradicts the gold target, grade **B**.
3. Hedging ("I think...") is fine if the correct info is present and no incorrect info is added.
4. Partial answers are **B**. Typos that preserve meaning are allowed.

DateAugBench specifics:

- **Date format ambiguity:** gold lists every valid interpretation; accept any.
- **Date arithmetic:** prediction must match *day, month, year* of a listed variant, any textual format allowed.
- **Format-switch questions:** answer with any synonym of Yes/True or No/False.
- **Numeric answers** – must match the gold number to the last shown significant digit.

Output format
Return exactly one capital letter:

A or B or C

No additional text or punctuation.

Example template

Question: {question}
Gold target: {target}
Predicted answer: {predicted_answer}

Now grade:

A or B or C

Table 10: LLM-as-Judge prompt used for comparing model and gold answers in the three DateAugBench tasks.

Human Evaluation	
Context-based resolution Prompt: Who was the chair of Allgemeiner Deutscher Fahrrad-Club in 17/10/2016? Gold Answer: Ulrich Syberg Model Prediction: As of October 17, 2016, the Federal Chairman was Ulrich Syberg Human Annotator Rating: A LLM-as-Judge Rating: A	
Date arithmetic Prompt: What date is 60 days after 05/01/1225? Gold Answer: March 6, 1225 , June 29, 1225 Model Prediction: July 30, 1225 Human Annotator Rating: B LLM-as-Judge Rating: B	

Table 11: Human evaluation of LLM-as-judge.