



Evaluating Robustness of Large Language Models with Neologisms

Jonathan Zheng, Alan Ritter, Wei Xu

College of Computing

Georgia Institute of Technology

jzheng324@gatech.edu; {wei.xu, alan.ritter}@cc.gatech.edu

Abstract

The performance of Large Language Models (LLMs) degrades from the temporal drift between data used for model training and newer text seen during inference. One understudied avenue of language change causing data drift is the emergence of neologisms – new word forms – over time. We create a diverse resource of recent English neologisms by using several popular collection methods. We analyze temporal drift using neologisms by comparing sentences containing new words with near-identical sentences that replace neologisms with existing substitute words. Model performance is nearly halved in machine translation when a single neologism is introduced in a sentence. Motivated by these results, we construct a benchmark to evaluate LLMs’ ability to generalize to neologisms with various natural language understanding tasks and model perplexity. Models with later knowledge cutoff dates yield lower perplexities and perform better in downstream tasks. LLMs are also affected differently based on the linguistic origins of words, indicating that neologisms are complex for static LLMs to address. We will release our benchmark and code for reproducing our experiments.

1 Introduction

Neologisms – recent word forms representing a new meaning, sense, or connotation (Cartier, 2017) – consistently surface as language changes. Neologisms emerge to describe the ever-changing state of the world, such as new terms created during the COVID-19 pandemic. While humans easily adapt to language change, large language models (LLMs) struggle with the misalignment of training data and new test data distributions (Luu et al., 2022).

Prior work on temporal language change (Lazari-dou et al., 2021; Onoe et al., 2022; Luu et al., 2022) observed model degradation when finetuning on older text and evaluating on newer data and named entities (Rijhwani and Preotiuc-Pietro, 2020; Agar-

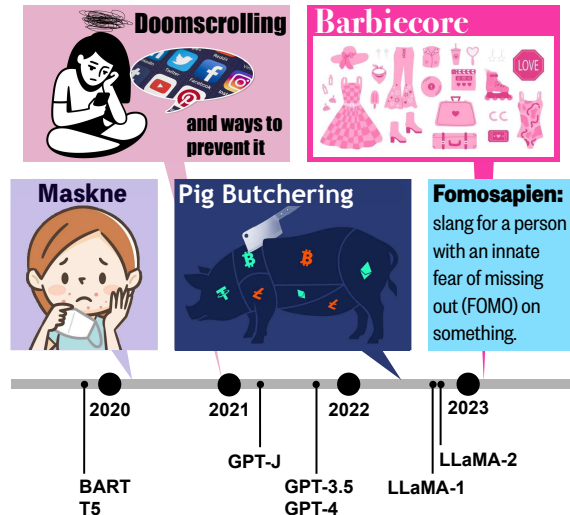


Figure 1: NEO-BENCH collects neologisms from 2020-2023 for LLM evaluation. “Pig Butchering” originated as a Mandarin expression (杀猪盘).

wal and Nenkova, 2022; Liu and Ritter, 2023). However, as far as we are aware there has not been prior work that analyzes the robustness of LLMs on handling neologisms. We show that adding a neologism to text decreases machine translation quality by an average of 43% in a human evaluation (§2), even for popular words emerging before 2020.

In this paper, we present NEO-BENCH, a new benchmark designed to test the ability of LLMs to understand and process neologisms. We combine multiple methods and online text corpora to collect a diverse set of 2,505 neologisms based on the linguistic taxonomy devised by Pinter et al. (2020): (i) **lexical neologisms** – words representing new concepts, e.g., “long covid”; (ii) **morphological neologisms** – blends of existing subwords, e.g., “doomscrolling”; and (iii) **semantic neologisms** – existing words that convey a new meaning or sense, e.g., “ice” (a term that refers to petrol- or diesel-powered cars taking electric car charging spots). We estimate word prevalence over time with Google Trends to obtain trending neologisms. We also create 4 benchmark tasks to evaluate the

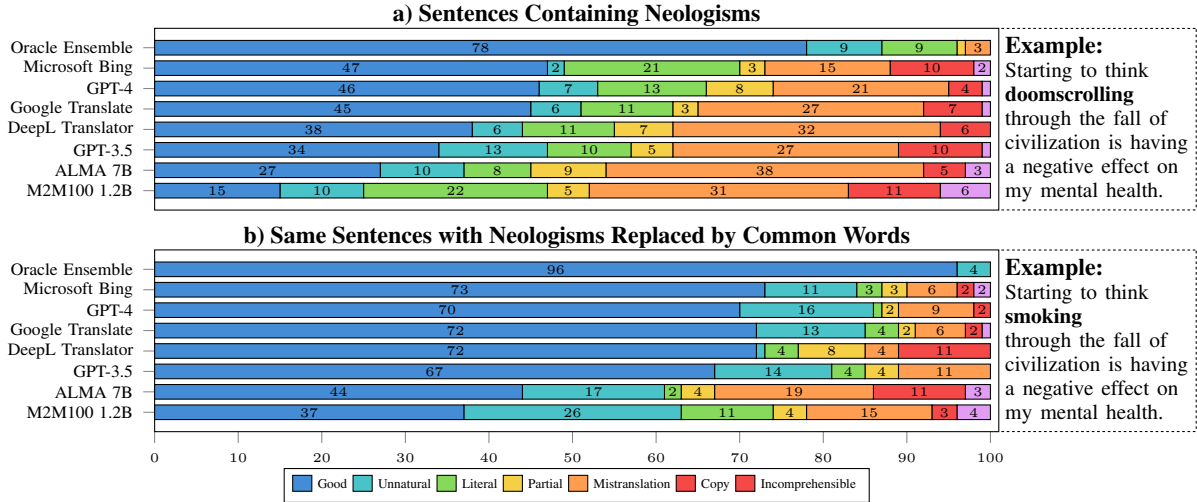


Figure 2: A single neologism can dramatically affect model output, as shown by human evaluation of Machine Translation models on sentences containing neologisms and the same sentences with neologisms replaced by carefully chosen words that also fit in the context. Oracle ensemble selects the best translation from all models.

impact of neologisms on LLMs with Perplexity, Cloze Question Answering, Definition Generation, and Machine Translation.

We show that lower neologism perplexities correlate with higher downstream task performance. Older LLMs – BART, T5, GPT-J, and Flan-T5 – perform much worse with an average of 32.20% and 12.27% accuracy in question answering and definition generation, respectively. We also find that automatic metrics do not accurately measure the quality of translated sentences containing neologisms, evidenced by Spearman’s ρ rank correlation between COMETKiwi (a state-of-the-art metric) and human judgment, which is 0.491. This is lower than the average ρ of 0.629 for COMETKiwi across 5 language pairs reported in the WMT23 Quality Estimation task (Blain et al., 2023). LLM performance in NEO-BENCH also differs based on a word’s linguistic type, as lexical neologisms without derivations yield the highest perplexities and the most fragmented subword tokenization, while semantic neologisms that repurpose existing words result in literal definitions and translations.

NEO-BENCH evaluates a diverse set of LLM capabilities on handling neologisms in various tasks. Models must also understand compositionality for morphological neologisms, differentiate between word senses for semantic neologisms, and handle different contexts for lexical neologisms.

2 Motivation

We start by using machine translation as an example to illustrate the significant challenge neologisms pose on state-of-the-art NLP systems. We

manually collect 100 neologism words with sentential context from social media, news articles, and dictionaries. GPT-3.5, GPT-4 and commercial translation systems, e.g., Google Translate,¹ Microsoft Bing,² and DeepL Translator,³ only managed to correctly translate about 34-47% of these 100 sentences that contain neologisms based on our manual inspection (Figure 2; from English to Chinese). In stark contrast, when replacing the neologism with a common word in these sentences, the percentage of correct translations rises substantially to 67-73%. We observe similar trends in open-source translation models, such as ALMA (Xu et al., 2023) and M2M100 (Fan et al., 2020).

One thing to note is that these replacement words are not exact synonyms, but words that have been carefully chosen to create a near-identical, semantically plausible sentence; because new words emerge in areas not occupied by existing words (Ryskina et al., 2020), true synonyms would often be verbose and incompatible with the sentence context. Because the original sentences containing neologisms were collected in the wild, one might assume they would be even more natural in comparison to their modified counterparts, but yet, **there is a large gap in translation quality between neologism and non-neologism words for all models.**

A closer look reveals that six typical types of errors are made in mistranslated model outputs, which include (ordered by severity):

- **Unnatural:** Imperfect translation of the sen-

¹<https://translate.google.com/>

²<https://www.bing.com/translator>

³<https://www.deepl.com/translator>

tence due to grammatical errors;

- **Literal:** Inaccurate output that literally translates the neologism or remaining sequence;
- **Partial:** Part of the sentence is untranslated and left out of the output;
- **Mistranslation:** Incorrectly translated sentence portion leads to a poor understanding of the overall sentence meaning;
- **Copy:** Part of the output is not translated and copied from the English input;
- **Incomprehensible:** Incoherent output that fails to capture any original sentence meaning;

Table 12 in the Appendix shows translations for each error type. The most common errors are mistranslations and literal translations with an average of 27.3% and 13.7% respectively. Model output for non-neologism sentences is more likely to have minor errors and be labeled unnatural by annotators.

Another interesting observation is that newer neologisms indeed show lower rates of good translations and often higher rates of mistranslations, as one may expect. Figure 3 shows the percentage of good translations and mistranslations over time for varied models. Compared to non-neologism sentences, models still yield lower rates of correct translations for neologisms that emerged before 2020. Many neologisms use existing words to convey meanings, such that the poor performance of models is not wholly explained by the absence of these word forms in training data. We propose a novel benchmark (§3) to systematically study the impact of neologisms on LLMs (§4).

3 NEO-BENCH: A Neologism Benchmark

We create NEO-BENCH, a benchmark that consists of 2,505 neologisms (both words and phrases) that newly emerged around 2020–2023 and 4 intrinsic/extrinsic tasks (Table 1) to evaluate LLMs’ abilities to generalize on neologisms. To facilitate continuous research on neologisms and language change, we intend to periodically update NEO-BENCH with neologisms emerging after 2023.

3.1 Neologism Collection

A neologism is a term that represents a new meaning or sense (Cartier, 2017). Previous datasets (McCrae, 2019; Ryskina et al., 2020; Zhu and Jurgens, 2021) only collected specific word types, ignored neologisms conveying new meanings with existing words, and did not utilize word prevalence trends

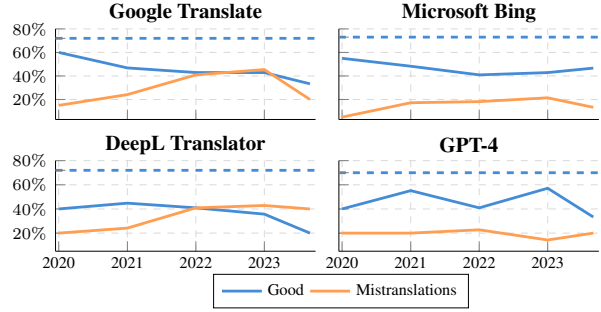


Figure 3: Percentage of good translations and mistranslations of neologism sentences over time. The dashed line represents the percentage of good translations achieved on non-neologism sentences.

Task	Dataset	Evaluation
Machine Translation	240 sentences containing neologisms	BLEU, COMET
Perplexity Ranking	422 Cloze passages with one-word answers	Word ranking
Cloze Questions	750 Cloze passages with multiple choice answers	Accuracy
Definition Generation	750 "What is [neologism]?" questions	Accuracy

Table 1: Summary of datasets in NEO-BENCH.

(more in Related Work §6). We design a more systematic collection process to quantify the effect of neologisms on a language’s data distribution.

Filtering Reddit Data based on Google Trends (Method 1). New words commonly propagate in online communities (Zhu and Jurgens, 2021), thus, we count word frequencies in monthly Reddit data to find single-word neologism candidates. We set a frequency cutoff between 50 and 100 per month to obtain uncommon words and remove misspellings and named entities using SpaCy (Montani et al., 2023), resulting in 74,542 candidates. We further obtain word search frequencies from 2010 to 2023 on Google Trends⁴ and automatically filter out 87.13% of neologism candidates based on these trend lines (see Figure 4 for examples) by a combination of curve fitting, argmax detection, and integrals over time. Appendix §B.1 provides more details about trend filtering. From the set of 9,590 remaining candidates, we find that 10.48% are prevalent neologisms by manual inspection. In total, we collected 1,005 neologism words from Reddit (310 lexical, 588 morphological, and 107 semantic neologisms).

Retrieving News Articles about Neologisms (Method 2). As Method 1 is only good at find-

⁴<https://trends.google.com/trends/>

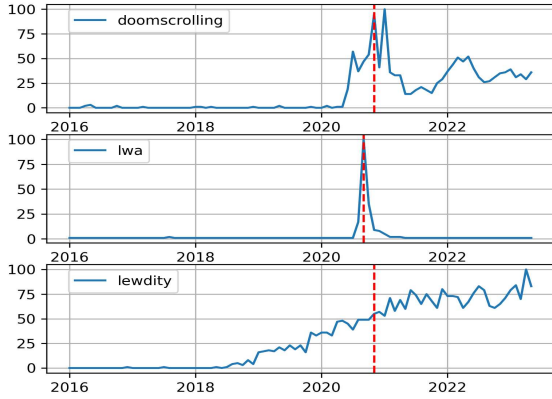


Figure 4: Example Google Trend lines measuring neologism prevalence. The dashed line estimates the date a neologism becomes popular while not yet conventional.

ing single-word neologisms, we turn to news articles that explain the meanings of neologisms to collect multi-word expressions. We first manually get 100 neologisms from news articles, recording news headlines of neologisms. Then, based on the shared text patterns of headlines, we created 16 headline **templates** (e.g., “___: What is it?”) to retrieve Google News articles from 2019 to 2023. Using SpaCy, we identify 60,671 noun and verb phrases with a Part-of-Speech tagger and remove duplicates and named entities. We used the same aforementioned filtering method for these phrases using Google Trends. From the remaining 8,039 candidates, we manually extracted 1,100 neologisms (778 lexical, 222 morphological, and 100 semantic neologisms), of which 713 are multiwords.

Sampling Existing Neologism Datasets (Method 3). To supplement our dataset with additional neologisms, we also sample from two existing open-source resources that contain a lot of rare words, many of which have no Google Trends data available. The NYT First Said Twitter bot (@NYT_first_said) tweets out words when they are used for the first time in New York Times articles by using exclusion lists. We retrieve 1,100 of its tweets from 2020 to collect 200 derived neologisms (192 morphological and 8 semantic). We also sample 1,400 entries from another noisy, automatically constructed, dataset of 80,071 new slang dictionary entries (Zhu and Jurgens, 2021). We manually filter the sample and collect 200 derived neologisms (4 lexical, 194 morphological, and 2 semantic).

Overall. We collected 2,505 neologism words. While semantic neologisms are infrequent in all sources, Google Trends data enables the collection of them, as these words change in baseline prevalence when a new sense is being popularized. Only

Neologism: <i>doomscrolling</i>	
<i>The silver lining of this website no longer functioning as an even vaguely reliable information source is that ___ has basically been completely undermined. It wouldn't even work now since everything is too geared to outrage clickbait and actual reporting has disappeared, so there is no point staying on the app.</i>	
a) misinformation	b) surfing
c) doomscrolling	d) lying
e) gaming	f) anti-productivity (distractor)

Table 2: Example passage in NEO-BENCH for multiple-choice Cloze Question Answering with correct neologism answers and partially correct distractor answers.

5.04% of words from Reddit, 1.12% of phrases from news articles, and 3.09% of entries from previous datasets overlap with candidates from the other two sources — highlighting the importance of using multiple diverse data sources and methods for neologism collection. We also verified that 44.23% of these 2,505 words actually appear in the Urban Dictionary, a crowdsourced English-language online dictionary for slang words and phrases.

3.2 Benchmark Tasks

NEO-BENCH consists of 4 tasks – 3 downstream and 1 intrinsic metric – to evaluate models’ knowledge of neologisms: (i) Machine Translation with human and automatic evaluation; (ii) Cloze Question Answering to evaluate models in context; (iii) Definition Generation to evaluate models in a context-free setting; and (iv) perplexity to compare single-word neologisms to commonly used words. We describe the setup and result tables/figures in this section, then discuss the key findings based on these results more in-depth in §4.

Machine Translation (Task 1). We sample from our collected neologisms (§3.1) and search for reference sentences containing these words on social media and Google. We construct 240 sentences, including the 100 used in §2. We work with in-house native speakers to create reference translations (English to Chinese) and evaluate system outputs in Table 5 with BLEU (Papineni et al., 2002), COMET (Rei et al., 2020), and MetricX-23 (XXL and XL) (Juraska et al., 2023). We also use the reference-free metrics COMETKiwi (Rei et al., 2022) and MetricX-23-QE (XXL and XL). Lower MetricX scores indicate higher performance. We report the correlations of metrics with human ratings for the MT models in §2 using Spearman’s ρ in Table 5.

Cloze Question Answering (Task 2). We sample 750 neologisms to create text passages, where one sentence contains a neologism and the remaining passage serves as preceding or following context

a) Definition Generation Output Examples	
Stablecoin	Reference Definition: A stablecoin is a type of cryptocurrency where the value of the digital asset is supposed to be pegged to a reference asset, which is either fiat money, exchange-traded commodities, or another cryptocurrency.
	Model Output (Correct): Stablecoins are cryptocurrencies designed to maintain a stable value, typically by pegging their value to a specific asset or basket of assets, such as the US dollar, gold, or a combination of assets.
Angel Shot	Reference Definition: An angel shot is a code to inform a bartender that a customer is not safe and needs assistance.
	Model Output (Incorrect): An angel shot is a cocktail made with whiskey and cream, served in a shot glass.
b) Machine Translation Output Examples	
Longcovid	Input: Each reinfection increases the risk of longcovid, hospitalization, & death.
	Model Output (Correct): 每次再感染都会增加长新冠病毒、住院和死亡的风险。 (Every reinfection increases the risk of long COVID, hospitalization, and death.)
	Human Translation: 每一次新冠感染都会提高出现后遗症、住院治疗，甚至死亡的风险。 (Each COVID-19 infection increases the risk of developing sequelae, hospitalization, and even death.)
Doomscrolling	Input: Starting to think doomscrolling through the fall of civilization is having a negative effect on my mental health.
	Model Output (Incorrect): 开始认为在文明的衰落中滚动的厄运对我的心理健康产生了负面影响。 (Start to think that the doom rolling in the decline of civilization is having a negative impact on my mental health.)
	Human Translation: 开始觉得，刷关于文明衰败的负能量新闻对我的心理健康产生了负面影响。 (Starting to feel that scrolling through negative news about the decline of civilization is having a negative impact on my mental health.)

Table 3: Example model definitions and translations for NEO-BENCH tasks. “Doomscrolling” is the act of spending an excessive amount of time reading negative news online. (English translations are shown for information only.)

(see example in Table 2). We mask out the neologism and provide four incorrect answers plus one distractor answer, which is a common word or phrase that is feasible in context. We evaluate BART-large (Lewis et al., 2019), T5-Large (Raffel et al., 2020), Flan-T5-Large (Chung et al., 2022), GPT-J 6B (Wang and Komatsuzaki, 2021), LLaMA-1 7B (Touvron et al., 2023a), Alpaca 7B (Taori et al., 2023), LLaMA-2, LLaMA-2-Chat (Touvron et al., 2023b), OLMo-7B (Groeneveld et al., 2024), OLMo-7B-Instruct, Mistral-7B (Jiang et al., 2023), Mistral-7B-Instruct, GPT 3.5 (Brown et al., 2020), and GPT-4 in multiple-choice Cloze Question Answering (QA). We experiment with 5-shot prompting and test three sizes of LLaMA-2 models. We show results in Figure 5 with the stratified and combined accuracies of selecting either the neologism or distractor answer.

Open-ended Definition Generation (Task 3). We evaluate the same models from Task 2 for their context-free knowledge of 750 neologisms with question prompts (i.e., “What is doomscrolling?”) to obtain neologism definitions. We construct human reference definitions and use GPT-4 to evaluate if model generations are semantically equivalent to the gold reference. We use 5-shot prompting and report results with accuracy in Figure 5. Table 3 shows example LLM-generated definitions.

Perplexity Rankings (Task 4). Using 422 Cloze passages that have both singular distractor and ne-

Label	Complete	Partial	Unknown
Good	53.13%	34.78%	30.77%
Unnatural	9.38%	4.35%	0.00%
Literal	10.94%	17.39%	15.38%
Partial	4.69%	13.04%	15.38%
Mistranslation	20.31%	21.74%	23.09%
Copy	1.55%	4.35%	15.38%
Incomprehensible	0.00%	4.35%	0.00%
Total	64	23	13

Table 4: GPT-4’s understanding of a neologism does not result in high machine translation performance. GPT-4 MT output is separated by its performance in Cloze QA, Definition Generation, and Definition Prompting. GPT-4 shows full, partial, and no knowledge of a neologism if zero, one, or multiple tasks are incorrect, respectively.

ologism answers, we use perplexity to evaluate GPT-J 6B, LLaMA-1 7B, Alpaca 7B, LLaMA-2 7B, LLaMA-2 Chat 7B, OLMo-7B, OLMo-7B-Instruct, Mistral-7B, and Mistral-7B-Instruct. For each passage, we use rank classification (Brown et al., 2020), where we fill in the mask with the neologism and measure the perplexity of the passage. We replace the mask with the distractor answer and the top 5000 singular words from Reddit by frequency and measure perplexities of all 5002 sequences. The mask-filling words are sorted by the corresponding sequence perplexity, and the average rankings of neologisms and distractors are reported in Figure 7. Lower neologism rankings represent lower relative perplexities and show that models are likely to complete the passage with a neologism.

Model (human rank)		Reference-Based Metrics				Reference-Free Metrics		
		BLEU $\uparrow$	COMET $\uparrow$	MX-23 _{XXL} $\downarrow$	MX-23 _{XL} $\downarrow$	COMET _{Kiwi} $\uparrow$	MX-QE _{XXL} $\downarrow$	MX-QE _{XL} $\downarrow$
Bing Translator	(1)	0.452 (2)	0.825 (5)	2.419 (6)	2.343 (6)	0.788 (5)	1.679 (5)	2.246 (5)
GPT-4	(2)	0.446 (3)	0.854 (1)	1.550 (1)	1.793 (1)	0.793 (3)	1.432 (3)	2.089 (3)
Google Translate	(3)	0.507 (1)	0.853 (2)	1.825 (4)	1.945 (4)	0.800 (2)	1.429 (2)	1.940 (1)
DeepL Translator	(4)	0.406 (4)	0.842 (3)	1.775 (3)	1.901 (3)	0.807 (1)	1.260 (1)	1.944 (2)
GPT-3.5	(5)	0.399 (5)	0.841 (4)	1.705 (2)	1.796 (2)	0.792 (4)	1.467 (4)	2.157 (4)
ALMA 7B (LLaMA-2)	(6)	0.285 (7)	0.801 (6)	2.382 (5)	2.251 (5)	0.746 (6)	2.038 (6)	2.462 (6)
M2M100 1.2B	(7)	0.337 (6)	0.776 (7)	3.454 (7)	3.142 (7)	0.745 (7)	2.821 (7)	2.976 (7)
Spearman’s ρ		0.244	0.445	0.457	0.380	0.491	0.451	0.445

Table 5: Machine Translation models evaluated on neologisms with BLEU, COMETKiwi, COMET, **MetricX-23**, and **MetricX-23-QE**. We use the XXL and XL sizes for MetricX. Rankings of models are provided for metrics and human evaluation for models used in §2. Spearman’s ρ between each metric and human evaluation is also reported.

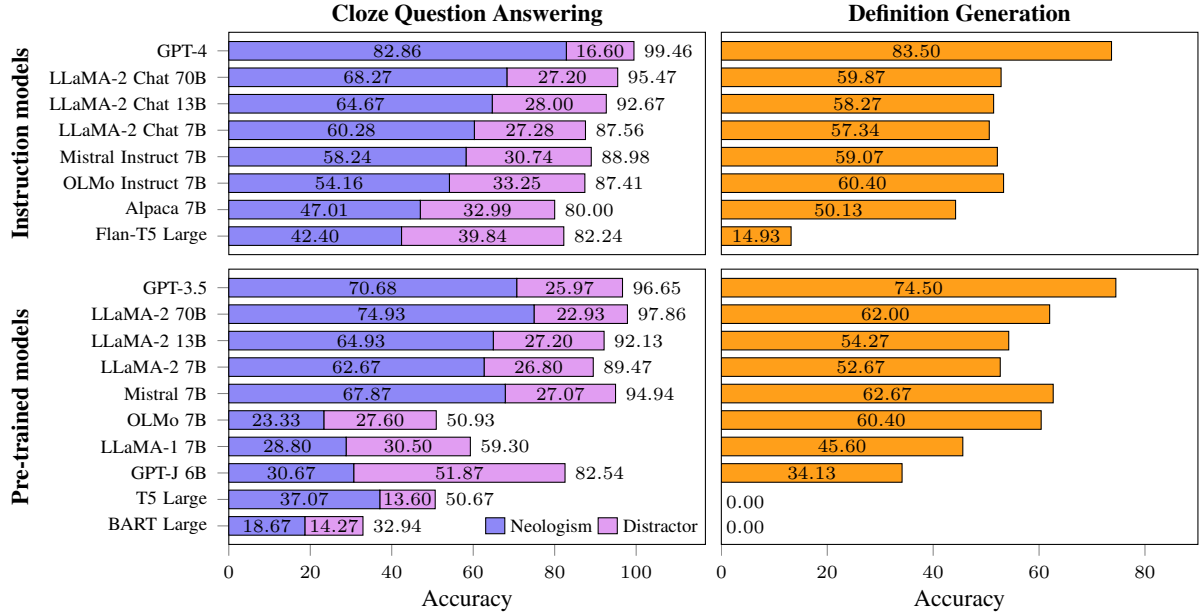


Figure 5: **Left:** Results of the Cloze Question Answering task reported by accuracy of selecting the neologism or distractor option. Combined accuracy for selecting either answer is provided. **Right:** Results of the Definition Generation task reported with accuracy of correct definitions. 5-shot prompting of models is used for both tasks.

4 Key Findings

We utilize NEO-BENCH tasks to evaluate the ability of various LLMs to adapt to neologisms. The following are our key findings:

Current automatic metrics cannot accurately evaluate MT models that struggle with neologisms. In §2, MT models decrease in performance by 43% when translating neologisms with Bing being the best model based on human evaluation. However, COMET and COMETKiwi scores are notably high and MetricX-23 error scores are low for all models in Table 5. The best models are Google Translate for BLEU (0.507); DeepL for COMETKiwi (0.807) and MetricX-23-QE (1.260); and GPT-4 for COMET (0.854) and MetricX-23 (1.550), highlighting that automatic metrics show poor system-level correlations with human judgments. For sentence-wise correlation between

MT metrics and human evaluations, the average Spearman’s ρ of BLEU, COMET, COMETKiwi, MetricX-23 and MetricX-23-QE (XXL) is 0.244, 0.445, 0.491, 0.457, 0.451, respectively. In contrast, COMETKiwi, our highest correlating metric, has an average ρ of 0.629 for five language pairs on the WMT23 Quality Estimation task for direct assessment (Blain et al., 2023). From our reference sentences, translating neologisms often requires paraphrasing, resulting in low ρ for BLEU.

GPT-4’s knowledge of neologisms is task specific. Table 4 presents the human annotations of GPT-4 translations of neologisms, separated by the corresponding performance of neologisms in Cloze QA, Definition Generation, and Definition Prompting, where we ask GPT-4 if the provided human reference definitions of neologisms are correct. GPT-4 shows complete knowledge of a neologism if all tasks are correct; partial knowledge if one task is in-

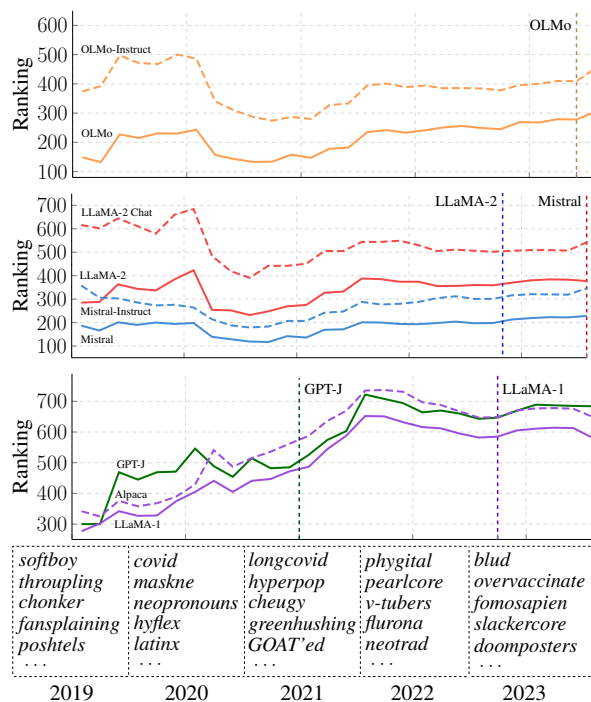


Figure 6: Rankings of neologisms over time compared to 5000 common words. Newer models are plotted separately. Dashed lines show model knowledge cutoffs⁵. Example neologisms from each year are provided, and neologisms without trendlines are reported at the end.

correct; and unknown if multiple tasks are incorrect. There are higher rates of correct translations for neologisms GPT-4 understands, as good translations drop by 20.3% if GPT-4 does not fully know a neologism’s meaning. However, the rate of mistranslations are constant regardless of GPT-4 performance in other tasks, and GPT-4 only correctly translates 53.13% of neologisms it has complete knowledge of. GPT-4’s knowledge of neologisms is compartmentalized and does not result in similarly high performance in machine translation compared to other Neo-Bench tasks, emphasizing the difficulty of translating neologisms.

Models perform worse on neologisms than pre-existing words. For Cloze questions in Figure 5, neologism answers are designed to be more natural as the original passages contained these neologisms, yet all models select a large portion (27.99% on average) of distractor answers. Neologisms also have an average perplexity rank of 463 compared to distractor rankings of 45 in Figure 7, indicating much lower perplexity for pre-existing words.

Older LLMs perform significantly worse. In Figure 5, the average performance of GPT-J, BART, T5, and Flan-T5 is 26.99% lower in Cloze QA and

⁵Mistral AI does not reveal training data details, so we provide our best estimation for the model’s knowledge cutoff.

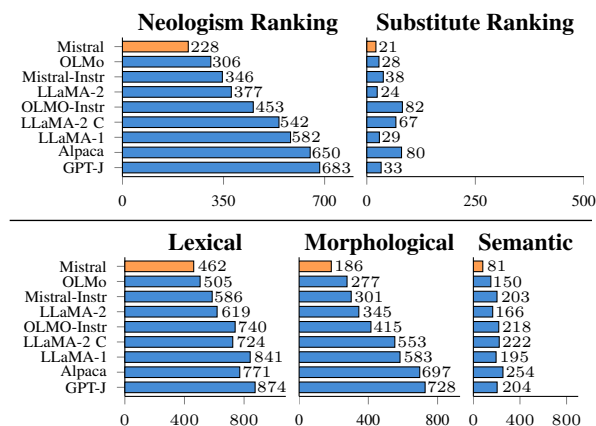


Figure 7: Average rankings of neologisms and pre-existing substitute terms compared to 5000 common words, sorted by model perplexities of texts filled in with each word. Neologisms are separated by linguistic type: lexical, morphological, and semantic.

47.78% lower in Definition Generation than other models. In Figure 7, GPT-J and LLaMA-1 models exhibit higher neologism rankings than newer open-source models, correlating with lower downstream performance. Newer models – GPT-3.5, GPT-4, LLaMA-2, OLMO, and Mistral – perform better as they are trained on data containing newer neologisms, generally have algorithmic improvements, and are trained with more resources than older models.

Perplexity rankings of older models increase drastically from 2019 until 2021. While neologisms are often gradually worked into a vocabulary (Zhu and Jurgens, 2021), we use trend lines to best estimate the date when a neologism becomes popular and report perplexity over time in Figure 6. Newer models dip in 2020 but increase afterward as 52% of neologisms from this period are now conventionalized terms related to COVID-19.

Larger models handle neologisms better. Increasing the sizes of LLaMA-2 and LLaMA-2 Chat leads to consistent improvements across both Cloze Question Answering and Definition Generation. On average, LaMA-2 70B and LLaMA-2 Chat 70B yield 10.13% higher accuracy in Cloze QA and 5.93% higher accuracy in Definition Generation than LLaMA-2 7B and LLaMA-2 Chat 7B.

Instruction-tuning results in high neologism perplexities. In Figure 7, LLaMA-1, LLaMA-2, OLMO, and Mistral models have, on average, 125 lower neologism rankings than their instruct-tuned counterparts. Instruct models are trained with dialogue (Wei et al., 2022), so uncommon generation is less desired.

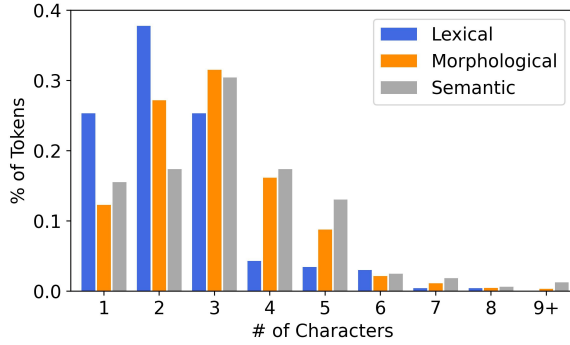


Figure 8: Distributions of characters per subword token of each neologism type, reported by the proportion of tokens that have a certain number of characters.

5 Linguistic Taxonomy Analysis

We separate NEO-BENCH task results by neologism linguistic structure: lexical, morphological, and semantic. Figure 7 presents perplexity rankings, Figure 9 reports human evaluation for MT, and Figure 10 shows the results for the best models on Cloze QA and Definition Generation.

Lexical neologisms produce the highest perplexities, but yield the best downstream results. On average, lexical neologisms have 226 higher rankings than other words, indicating higher relative model perplexities. Figure 8 shows the distribution of characters per token of neologisms using the LLaMA tokenizer, and the average number of characters per token for lexical, morphological, and semantic neologisms is 2.36, 2.98, and 3.24 respectively. Lexical neologisms have more fragmented tokenizations, as these words have the highest proportion of 1-2 character tokens. Lexical neologisms are less likely to be separated into long, common word roots or segments representing subword information, instead producing uncommon token sequences that result in higher neologism rankings and perplexity. In downstream tasks, however, lexical neologisms yield 0.6% higher Cloze accuracy, 8.8% more correct definitions, and 21.5% more good translations than other neologisms.

Morphological neologisms produce low perplexities but yield poor downstream performance. Compared to lexical neologisms, morphological neologisms are, on average, segmented into longer tokens and constructed with common subwords, resulting in lower perplexity rankings. However, they yield 4.2% lower Cloze accuracy, 9.1% more incorrect definitions, and 26.1% less good translations than lexical neologisms. 76.8% of neologisms without trend lines are morphological. Compared to lexical and semantic neologisms that require preva-

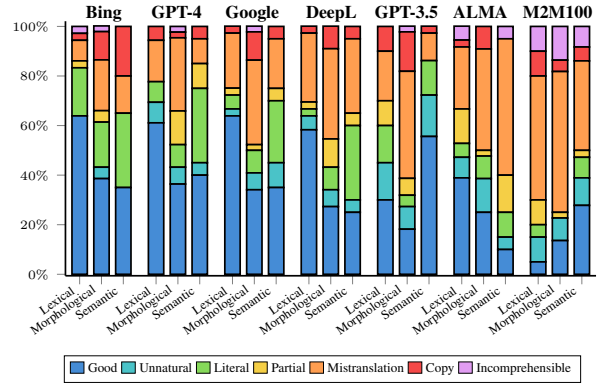


Figure 9: Results of the human-annotated MT models for each linguistic type of neologism.

lence to be differentiated from incoherent strings, morphological neologisms are created with polynomial combinations of common subwords. Many of these intelligible combinations are largely unused, resulting in lower downstream performance.

Semantic neologisms produce the lowest perplexities and the worst performance in generation tasks. Since these neologisms use existing word forms, they have an average of 266 lower perplexity rankings than other neologism words. While semantic neologisms yield high Cloze QA accuracy, they also achieve the lowest percentages of correct definitions and translations. Models produce popular definitions and literal translations based on a word’s most common meaning, as the new sense of semantic neologisms is often nuanced and difficult to capture.

6 Related Work

Temporal Drift in LLMs. Prior work has explored temporal data drift by creating temporal splits of training data (Loureiro et al., 2022; Luu et al., 2022; Röttger and Pierrehumbert, 2021; Jin et al., 2022; Luu et al., 2022; Lazaridou et al., 2021). New factual updates of concepts are studied with temporal splits of text corpora and QA datasets (Jang et al., 2022; Margatina et al., 2023; Zhao et al., 2022; Vu et al., 2023). Other work has observed model degradation from new named entities (Onoe et al., 2022; Rijhwani and Preotiuc-Pietro, 2020; Chen et al., 2021). Temporal degradation occurs during short-term crisis events where information changes quickly (Pramanick et al., 2022). Studies have consistently found model degradation with perplexity and downstream tasks. There are no studies on model degradation from language change of neologisms, so we create a benchmark to evaluate models on neologisms with similar tasks.

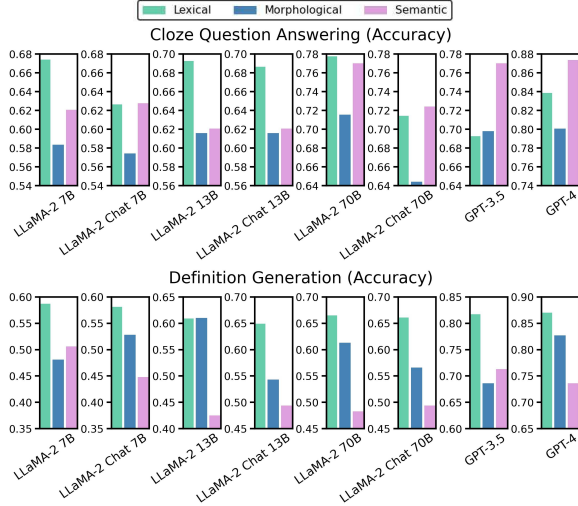


Figure 10: Results of Cloze QA and Definition Generation stratified by linguistic types of neologisms.

Neologism Collection. Using reference texts as exclusion lists to filter common words from target corpora is the most documented method of neologism detection. Target texts and exclusion lists include news articles (Pinter et al., 2020; Falk et al., 2014), dictionaries (Kerremans et al., 2018; Langemets et al., 2020; Liu et al., 2013; Dhuliawala et al., 2016), social media (Pyo, 2023; Zalmout et al., 2019; Megerdooian and Hadjarian, 2010) and other corpora (Cartier, 2017; Lejeune and Cartier, 2017). These texts are slow to curate, and semantic neologisms are filtered out. Moreover, no resource collects general semantic neologisms.

Some resources measure word prevalence with time-series data of search queries to collect single-word neologism candidates (Broad et al., 2018) or cybersecurity neologisms (Li et al., 2021). They are limited in scope by collecting only one type of neologism based on rising popularity. Other methods collect neologisms with new slang dictionary entries (Dhuliawala et al., 2016; Zhu and Jurgens, 2021). Dictionaries are slow to update, so new entries are often conventionalized words. There is no resource that uses time-series data to collect a variety of neologisms rising in prevalence.

Previous work has utilized search templates of explanation patterns to collect automatically neologisms (Breen et al., 2018). A few efforts have used neural methods to automatically detect one specific type of neologism, such as adjective-noun neologism pairs (McCrae, 2019), blend words (Megerdooian and Hadjarian, 2010), and grammatical neologisms (Janssen, 2012; Falk et al., 2014), which are existing words with new parts of speech. No automatic resource collects neologisms

from various topics and linguistic backgrounds. To address these limitations, NEO-BENCH uses multiple methods to semi-automatically collect a variety of neologisms, including multiword, semantic, and prevalent neologisms.

Unseen Words. Rare words or typos are typically unseen in data when training a model but may show up during inference. Prior work has measured model degradation from unseen words (Chirkova and Troshin, 2021; Nayak et al., 2020) and used contextual subword embeddings (Garneau et al., 2018; Chen et al., 2019; Hu et al., 2019; Wang et al., 2019; Araabi et al., 2022), similar surface forms of common words (Chen et al., 2022; Fukuda et al., 2020), and morphological structure (Lochter et al., 2020, 2022) to represent unseen words. Comparatively, the neologism lifecycle follows three stages: emergence, dissemination, and conventionalization (Cartier, 2017). New words often become prevalent and drastically shift a language’s distribution. Semantic neologisms also use existing word forms and are not classified as unseen words.

7 Conclusion

In this paper, we present NEO-BENCH, a new benchmark to test the ability of LLMs to generalize on neologisms. We use several methods to collect a variety of neologisms, including prevalent, multiword, and semantic neologisms. In our experiments, we find that models struggle with neologisms in both perplexity and downstream tasks. Machine Translation is especially difficult, as translating neologisms often requires paraphrasing the sentence. Current automatic metrics cannot measure translation quality, and human evaluation is still needed. Neologisms also affect models differently based on linguistic structure, indicating that this phenomenon is complex for LLMs to address.

Limitations

Most of our neologisms largely originated in US and UK English, as we collect textual data from news articles from this region. We do not restrict our locations for Reddit data, but the majority of English-speaking Reddit users are also from the same regions. Given our limited expertise in other English dialects, especially regions whose English variations are largely influenced by other languages, we do not collect many neologisms from English-speaking regions outside these regions. However,

our computational framework for collecting neologisms can be applied to any language or local variation. For temporal drift of multilingual language modeling, we leave multilingual neologism collection and temporal drift analysis up for future work. Additionally, NEO-BENCH is static as we collect neologisms from mostly 2020-2023, which will become outdated over time as newer models will be exposed to new language in context. However, the semi-automatic collection methods require minimal human supervision and can be dynamically updated to continuously obtain neologisms. These methods require time-series data of words and online text corpora without needing human-curated information like updated dictionaries to filter words. The time-independent filters can collect recent neologisms without needing the time-consuming, manual curation of temporal splits of textual data. Provided that the Google and Reddit Terms of Service and API access enable NEO-BENCH collection, we intend on periodically updating the set of neologisms in our dataset.

Ethical Considerations

We utilize Reddit monthly dumps to obtain uncommon words, which often include sensitive information such as account usernames. We take the appropriate measures to ensure that no personally identifiable information (PII) is included in our dataset. We use a named-entity recognition model via SpaCy to identify named entities that are potentially PII and largely remove this information automatically when filtering for neologism candidates. We also manually inspect all the candidates to ensure that no PII is included in our dataset. As we use natural references from Google to construct our model inputs, we also review our hand-crafted sentences to ensure that there is no PII contained in these sentences. Many neologism entries in our work emerge from slang, and some slang words have expletive or offensive meanings. The purpose of our dataset and benchmark is to obtain a representative sample of neologisms and comprehensively evaluate the impact of neologisms on Large Language Models. We present examples that do not contain such offensive information, but these offensive entries are nonetheless a consistent source of neologisms. For expletive neologisms, we strive to create input sentences that capture the meaning of the neologism while not perpetuating gender, racial, and other potential biases. We do

not collect any neologisms that directly reference stereotypes and demographic biases.

Acknowledgments

We would like to thank Piranava Abeyakaran, Kaushik Sriram, Vishnesh Jayanthi Ramanathan, Yao Dou, Chao Jiang, Cai Yang, and Xiaofeng Wu for their help in constructing model inputs and translating reference sentences; Qihang Yao and Sam Stevens for their scripts for processing Google Trends and Urban Dictionary data, respectively; Graham Neubig and Paul Michel for valuable discussions. We would also like to thank Azure’s Accelerate Foundation Models Research Program for graciously providing access to API-based models, such as GPT-4. This research is supported in part by the NSF CAREER awards IIS-2144493 and IIS-2052498, ODNI and IARPA via the HIATUS program (contract 2022-22072200004). The views and conclusions contained herein are those of the authors and should not be interpreted as necessarily representing the official policies, either expressed or implied, of NSF, ODNI, IARPA, or the U.S. Government. The U.S. Government is authorized to reproduce and distribute reprints for governmental purposes notwithstanding any copyright annotation therein.

References

- Oshin Agarwal and Ani Nenkova. 2022. Temporal effects on pre-trained models for language processing tasks. *Transactions of the Association for Computational Linguistics*.
- Ali Araabi, Christof Monz, and Vlad Niculae. 2022. [How effective is byte pair encoding for out-of-vocabulary words in neural machine translation?](#) In *Proceedings of the 15th biennial conference of the Association for Machine Translation in the Americas (Volume 1: Research Track)*, pages 117–130, Orlando, USA. Association for Machine Translation in the Americas.
- Steven Bird, Ewan Klein, and Edward Loper. 2009. *Natural language processing with Python: analyzing text with the natural language toolkit*. " O’Reilly Media, Inc."
- Frederic Blain, Chrysoula Zerva, Ricardo Ribeiro, Nuno M. Guerreiro, Diptesh Kanojia, José G. C. de Souza, Beatriz Silva, Tânia Vaz, Yan Jingxuan, Fatemeh Azadi, Constantin Orasan, and André Martins. 2023. [Findings of the WMT 2023 shared task on quality estimation](#). In *Proceedings of the Eighth Conference on Machine Translation*, pages 629–653,

- Singapore. Association for Computational Linguistics.
- James Breen, Timothy Baldwin, and Francis Bond. 2018. [The company they keep: Extracting Japanese neologisms using language patterns](#). In *Proceedings of the 9th Global Wordnet Conference*, pages 163–171, Nanyang Technological University (NTU), Singapore. Global Wordnet Association.
- Claire Broad, Helen Langone, and David Guy Brizan. 2018. [Candidate ranking for maintenance of an online dictionary](#). In *Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)*, Miyazaki, Japan. European Language Resources Association (ELRA).
- Tom B Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *arXiv preprint arXiv:2005.14165*.
- Emmanuel Cartier. 2017. [Neoville, a web platform for neologism tracking](#). In *Proceedings of the Software Demonstrations of the 15th Conference of the European Chapter of the Association for Computational Linguistics*, pages 95–98, Valencia, Spain. Association for Computational Linguistics.
- Lihu Chen, Gaël Varoquaux, and Fabian M. Suchanek. 2022. [Imputing out-of-vocabulary embeddings with love makes language models robust with little cost](#).
- Mingqing Chen, Rajiv Mathews, Tom Ouyang, and Françoise Beaufays. 2019. [Federated learning of out-of-vocabulary words](#). *CoRR*, abs/1903.10635.
- Shuguang Chen, Leonardo Neves, and Tamar Solorio. 2021. [Mitigating temporal-drift: A simple approach to keep ner models crisp](#).
- Nadezhda Chirkova and Sergey Troshin. 2021. [A simple approach for handling out-of-vocabulary identifiers in deep learning for source code](#).
- Hyung Won Chung, Le Hou, Shayne Longpre, Barret Zoph, Yi Tay, William Fedus, Yunxuan Li, Xuezhi Wang, Mostafa Dehghani, Siddhartha Brahma, Albert Webson, Shixiang Shane Gu, Zhuyun Dai, Mirac Suzgun, Xinyun Chen, Aakanksha Chowdhery, Alex Castro-Ros, Marie Pellat, Kevin Robinson, Dasha Valter, Sharan Narang, Gaurav Mishra, Adams Yu, Vincent Zhao, Yanping Huang, Andrew Dai, Hongkun Yu, Slav Petrov, Ed H. Chi, Jeff Dean, Jacob Devlin, Adam Roberts, Denny Zhou, Quoc V. Le, and Jason Wei. 2022. [Scaling instruction-finetuned language models](#).
- Shehzaad Dhuliawala, Diptesh Kanojia, and Pushpak Bhattacharyya. 2016. [SlangNet: A WordNet like resource for English slang](#). In *Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC’16)*, pages 4329–4332, Portorož, Slovenia. European Language Resources Association (ELRA).
- Ingrid Falk, Delphine Bernhard, and Christophe Gérard. 2014. [From non word to new word: Automatically identifying neologisms in French newspapers](#). In *Proceedings of the Ninth International Conference on Language Resources and Evaluation (LREC’14)*, pages 4337–4344, Reykjavik, Iceland. European Language Resources Association (ELRA).
- Angela Fan, Shruti Bhosale, Holger Schwenk, Zhiyi Ma, Ahmed El-Kishky, Siddharth Goyal, Man-deep Baines, Onur Celebi, Guillaume Wenzek, Vishrav Chaudhary, Naman Goyal, Tom Birch, Vitaliy Liptchinsky, Sergey Edunov, Edouard Grave, Michael Auli, and Armand Joulin. 2020. [Beyond english-centric multilingual machine translation](#). *CoRR*, abs/2010.11125.
- Markus Freitag, George Foster, David Grangier, Viresh Ratnakar, Qijun Tan, and Wolfgang Macherey. 2021. [Experts, errors, and context: A large-scale study of human evaluation for machine translation](#). *Transactions of the Association for Computational Linguistics*, 9:1460–1474.
- Nobukazu Fukuda, Naoki Yoshinaga, and Masaru Kit-suregawa. 2020. [Robust Backed-off Estimation of Out-of-Vocabulary Embeddings](#). In *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 4827–4838, Online. Association for Computational Linguistics.
- Nicolas Garneau, Jean-Samuel Leboeuf, and Luc Lamontagne. 2018. [Predicting and interpreting embeddings for out of vocabulary words in downstream tasks](#). In *Proceedings of the 2018 EMNLP Workshop BlackboxNLP: Analyzing and Interpreting Neural Networks for NLP*, pages 331–333, Brussels, Belgium. Association for Computational Linguistics.
- Dirk Groeneveld, Iz Beltagy, Pete Walsh, Akshita Bhagia, Rodney Kinney, Oyvind Taffjord, 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, Tushar Khot, William Merrill, Jacob Morrison, Niklas Muen-nighoff, Aakanksha Naik, Crystal Nam, Matthew E. Peters, Valentina Pyatkin, Abhilasha Ravichander, Dustin Schwenk, Saurabh Shah, Will Smith, Emma Strubell, Nishant Subramani, Mitchell Wortsman, Pradeep Dasigi, Nathan Lambert, Kyle Richardson, Luke Zettlemoyer, Jesse Dodge, Kyle Lo, Luca Soldaini, Noah A. Smith, and Hannaneh Hajishirzi. 2024. [Olmo: Accelerating the science of language models](#).
- David Heineman, Yao Dou, and Wei Xu. 2023. [Thresh: A unified, customizable and deployable platform for fine-grained text evaluation](#).
- Matthew Honnibal and Ines Montani. 2017. spaCy 2: Natural language understanding with Bloom embeddings, convolutional neural networks and incremental parsing. To appear.

- Ziniu Hu, Ting Chen, Kai-Wei Chang, and Yizhou Sun. 2019. [Few-shot representation learning for out-of-vocabulary words](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 4102–4112, Florence, Italy. Association for Computational Linguistics.
- Joel Jang, Seonghyeon Ye, Changho Lee, Sohee Yang, Joongbo Shin, Janghoon Han, Gyeonghun Kim, and Minjoon Seo. 2022. [TemporalWiki: A lifelong benchmark for training and evaluating ever-evolving language models](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 6237–6250, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Maarten Janssen. 2012. [NeoTag: a POS tagger for grammatical neologism detection](#). In *Proceedings of the Eighth International Conference on Language Resources and Evaluation (LREC’12)*, pages 2118–2124, Istanbul, Turkey. European Language Resources Association (ELRA).
- Albert Q. Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, Léo Renard Lavaud, Marie-Anne Lachaux, Pierre Stock, Teven Le Scao, Thibaut Lavril, Thomas Wang, Timothée Lacroix, and William El Sayed. 2023. [Mistral 7b](#).
- Xisen Jin, Dejiao Zhang, Henghui Zhu, Wei Xiao, Shang-Wen Li, Xiaokai Wei, Andrew Arnold, and Xiang Ren. 2022. [Lifelong pretraining: Continually adapting language models to emerging corpora](#). In *Proceedings of BigScience Episode #5 – Workshop on Challenges & Perspectives in Creating Large Language Models*, pages 1–16, virtual+Dublin. Association for Computational Linguistics.
- Juraj Juraska, Mara Finkelstein, Daniel Deutsch, Aditya Siddhant, Mehdi Mirzazadeh, and Markus Freitag. 2023. [MetricX-23: The Google submission to the WMT 2023 metrics shared task](#). In *Proceedings of the Eighth Conference on Machine Translation*, pages 756–767, Singapore. Association for Computational Linguistics.
- Daphné Kerremans, Jelena Prokić, Quirin Würschinger, and Hans-Jörg Schmid. 2018. [Using data-mining to identify and study patterns in lexical innovation on the web](#). *Pragmatics and Cognition*, 25(1):174–200.
- Margit Langemets, Jelena Kallas, Kaisa Norak, and Indrek Hein. 2020. [New estonian words and senses: Detection and description](#). *Dictionaries: Journal of the Dictionary Society of North America*, 41(1):69–82.
- Angeliki Lazaridou, Adhiguna Kuncoro, Elena Gribovskaya, Devang Agrawal, Adam Liska, Tayfun Terzi, Mai Gimenez, Cyprien de Masson d’Autume, Tomas Kocisky, Sebastian Ruder, Dani Yogatama, Kris Cao, Susannah Young, and Phil Blunsom. 2021. [Mind the gap: Assessing temporal generalization in neural language models](#).
- Gaël Lejeune and Emmanuel Cartier. 2017. [Character based pattern mining for neology detection](#). In *Proceedings of the First Workshop on Subword and Character Level Models in NLP*, pages 25–30, Copenhagen, Denmark. Association for Computational Linguistics.
- Mike Lewis, Yinhan Liu, Naman Goyal, Marjan Ghazvininejad, Abdelrahman Mohamed, Omer Levy, Ves Stoyanov, and Luke Zettlemoyer. 2019. [Bart: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension](#).
- Ying Li, Jiaxing Cheng, Cheng Huang, Zhouguo Chen, and Weina Niu. 2021. [NEDetector: Automatically extracting cybersecurity neologisms from hacker forums](#). *Journal of Information Security and Applications*, 58:102784.
- Shuheng Liu and Alan Ritter. 2023. [Do CoNLL-2003 named entity taggers still work well in 2023?](#) In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Toronto, Canada. Association for Computational Linguistics.
- Tsun-Jui Liu, Shu-Kai Hsieh, and Laurent Prévot. 2013. [Observing features of PTT neologisms: A corpus-driven study with n-gram model](#). In *Proceedings of the 25th Conference on Computational Linguistics and Speech Processing, ROCLING 2015, National Sun Yat-sen University, Kaohsiung, Taiwan, October 4-5, 2013*. Association for Computational Linguistics and Chinese Language Processing (ACLCLP), Taiwan.
- Johannes V. Lochter, Renato M. Silva, and Tiago A. Almeida. 2022. [Multi-level out-of-vocabulary words handling approach](#). *Knowledge-Based Systems*, 251:108911.
- Johannes V. Lochter, Renato Moraes Silva, and Tiago A. Almeida. 2020. [Deep learning models for representing out-of-vocabulary words](#). *CoRR*, abs/2007.07318.
- Daniel Loureiro, Francesco Barbieri, Leonardo Neves, Luis Espinosa Anke, and Jose Camacho-Collados. 2022. [Timelms: Diachronic language models from twitter](#).
- Kelvin Luu, Daniel Khashabi, Suchin Gururangan, Karishma Mandyam, and Noah A. Smith. 2022. [Time waits for no one! analysis and challenges of temporal misalignment](#). In *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 5944–5958, Seattle, United States. Association for Computational Linguistics.
- Katerina Margatina, Shuai Wang, Yogarshi Vyas, Neha Anna John, Yassine Benajiba, and Miguel Ballesteros.

2023. [Dynamic benchmarking of masked language models on temporal concept drift with multiple views](#). In *Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics*, pages 2881–2898, Dubrovnik, Croatia. Association for Computational Linguistics.
- John Philip McCrae. 2019. [Identification of adjective-noun neologisms using pretrained language models](#). In *Proceedings of the Joint Workshop on Multiword Expressions and WordNet (MWE-WN 2019)*, pages 135–141, Florence, Italy. Association for Computational Linguistics.
- Karine Megerdooimian and Ali Hadjarian. 2010. Mining and classification of neologisms in persian blogs. In *HLT-NAACL 2010*.
- Ines Montani, Matthew Honnibal, Matthew Honnibal, Adriane Boyd, Sofie Van Landeghem, and Henning Peters. 2023. [explosion/spaCy: v3.7.2: Fixes for APIs and requirements](#).
- Anmol Nayak, Hariprasad Timmapathini, Karthikeyan Ponnalagu, and Vijendran Gopalan Venkoparao. 2020. [Domain adaptation challenges of BERT in tokenization and sub-word representations of out-of-vocabulary words](#). In *Proceedings of the First Workshop on Insights from Negative Results in NLP*, pages 1–5, Online. Association for Computational Linguistics.
- Yasumasa Onoe, Michael Zhang, Eunsol Choi, and Greg Durrett. 2022. [Entity cloze by date: What LMs know about unseen entities](#). In *Findings of the Association for Computational Linguistics: NAACL 2022*, pages 693–702, Seattle, United States. Association for Computational Linguistics.
- 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.
- Yuval Pinter, Cassandra L. Jacobs, and Max Bittker. 2020. [NYTWIT: A dataset of novel words in the New York Times](#). In *Proceedings of the 28th International Conference on Computational Linguistics*, pages 6509–6515, Barcelona, Spain (Online). International Committee on Computational Linguistics.
- Aniket Pramanick, Tilman Beck, Kevin Stowe, and Iryna Gurevych. 2022. [The challenges of temporal alignment on Twitter during crises](#). In *Findings of the Association for Computational Linguistics: EMNLP 2022*, pages 2658–2672, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Jiyeon Pyo. 2023. Detection and replacement of neologisms for translation. Master’s thesis, The Cooper Union for the Advancement of Science and Art.
- Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J. Liu. 2020. [Exploring the limits of transfer learning with a unified text-to-text transformer](#). *Journal of Machine Learning Research*, 21(140):1–67.
- Ricardo Rei, Craig Stewart, Ana C Farinha, and Alon Lavie. 2020. [Comet: A neural framework for mt evaluation](#).
- Ricardo Rei, Marcos Treviso, Nuno M. Guerreiro, Chrysoula Zerva, Ana C Farinha, Christine Maroti, José G. C. de Souza, Taisiya Glushkova, Duarte Alves, Luisa Coheur, Alon Lavie, and André F. T. Martins. 2022. [CometKiwi: IST-unbabel 2022 submission for the quality estimation shared task](#). In *Proceedings of the Seventh Conference on Machine Translation (WMT)*, pages 634–645, Abu Dhabi, United Arab Emirates (Hybrid). Association for Computational Linguistics.
- Shruti Rijhwani and Daniel Preotiuc-Pietro. 2020. Temporally-informed analysis of named entity recognition. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, Online. Association for Computational Linguistics.
- Paul Röttger and Janet Pierrehumbert. 2021. [Temporal adaptation of BERT and performance on downstream document classification: Insights from social media](#). In *Findings of the Association for Computational Linguistics: EMNLP 2021*, pages 2400–2412, Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Maria Ryskina, Ella Rabinovich, Taylor Berg-Kirkpatrick, David Mortensen, and Yulia Tsvetkov. 2020. [Where new words are born: Distributional semantic analysis of neologisms and their semantic neighborhoods](#). In *Proceedings of the Society for Computation in Linguistics 2020*, pages 367–376, New York, New York. Association for Computational Linguistics.
- Rohan Taori, Ishaan Gulrajani, Tianyi Zhang, Yann Dubois, Xuechen Li, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. 2023. Stanford alpaca: An instruction-following llama model. https://github.com/tatsu-lab/stanford_alpaca.
- Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurelien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. 2023a. [Llama: Open and efficient foundation language models](#).
- 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, Brian Fuller, Cynthia Gao, Vedanuj Goswami, Naman Goyal, Anthony Hartshorn, Saghar Hosseini, Rui Hou, Hakan

Inan, Marcin Kardas, Viktor Kerkez, Madian Khabza, Isabel Kloumann, Artem Korenev, Punit Singh Koura, Marie-Anne Lachaux, Thibaut Lavril, Jenya Lee, Diana Liskovich, Yinghai Lu, Yuning Mao, Xavier Martinet, Todor Mihaylov, Pushkar Mishra, Igor Molybog, Yixin Nie, Andrew Poulton, Jeremy Reizenstein, Rashi Rungta, Kalyan Saladi, Alan Schelten, Ruan Silva, Eric Michael Smith, Ranjan Subramanian, Xiaoqing Ellen Tan, Binh Tang, Ross Taylor, Adina Williams, Jian Xiang Kuan, Puxin Xu, Zheng Yan, Iliyan Zarov, Yuchen Zhang, Angela Fan, Melanie Kambadur, Sharan Narang, Aurelien Rodriguez, Robert Stojnic, Sergey Edunov, and Thomas Scialom. 2023b. [Llama 2: Open foundation and fine-tuned chat models](#).

Tu Vu, Mohit Iyyer, Xuezhi Wang, Noah Constant, Jerry Wei, Jason Wei, Chris Tar, Yun-Hsuan Sung, Denny Zhou, Quoc Le, and Thang Luong. 2023. [Freshllms: Refreshing large language models with search engine augmentation](#).

Ben Wang and Aran Komatsuzaki. 2021. GPT-J-6B: A 6 Billion Parameter Autoregressive Language Model. <https://github.com/kingoflolz/mesh-transformer-jax>.

Changhan Wang, Kyunghyun Cho, and Jiatao Gu. 2019. [Neural machine translation with byte-level subwords](#). *CoRR*, abs/1909.03341.

Jason Wei, Maarten Bosma, Vincent Y. Zhao, Kelvin Guu, Adams Wei Yu, Brian Lester, Nan Du, Andrew M. Dai, and Quoc V. Le. 2022. [Finetuned language models are zero-shot learners](#).

Haoran Xu, Young Jin Kim, Amr Sharaf, and Hany Hassan Awadalla. 2023. [A paradigm shift in machine translation: Boosting translation performance of large language models](#).

Nasser Zalmout, Kapil Thadani, and Aasish Pappu. 2019. [Unsupervised neologism normalization using embedding space mapping](#). In *Proceedings of the 5th Workshop on Noisy User-generated Text (W-NUT 2019)*, pages 425–430, Hong Kong, China. Association for Computational Linguistics.

Zhixue Zhao, George Chrysostomou, Kalina Bontcheva, and Nikolaos Aletras. 2022. [On the impact of temporal concept drift on model explanations](#). In *Findings of the Association for Computational Linguistics: EMNLP 2022*, pages 4039–4054, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.

Jian Zhu and David Jurgens. 2021. [The structure of online social networks modulates the rate of lexical change](#). *CoRR*, abs/2104.05010.

A Related Work

Table 6 provides an overview and comparison of English neologism resources, including the types

of neologisms and collection methods used in each dataset. NEO-BENCH uses more collection methods and covers more types of neologisms, including multiword and semantic neologisms.

B Data Collection

We base the categories of neologisms on the linguistic taxonomy used in previous literature. Based on our empirical studies, we observe that all neologisms fall under three broad categories: lexical, morphological, and semantic. We additionally label neologisms based on subcategories of these broad linguistic classes (e.g., word blends, derivations, acronyms, and novel phrases).

For Reddit neologism candidates, we collected 500 million utterances in December 2021 and 200 million utterances from January to May 2022. We tokenize the utterances with the NLTK package (Bird et al., 2009) to get individual word counts and update a generic word counter. Neologism candidates are selected by filtering out typos and extremely rare words with less than a frequency of 50. We further filter out named entities by utilizing a SpaCy named entity recognition (NER) model (Honnibal and Montani, 2017) (en_core_web_sm) to detect proper nouns and update a named entity counter. We compare the counts of words from the general counter and the named entity counter and filter out the word if the proportion that a general word is in a named entity is greater than 0.5. The remaining words with the lowest frequencies are the list of uncommon words that we treat as neologism candidates for a given month. In total, we collect 74,542 neologism candidates.

For news articles, we use a script to collect 11,412 headlines from Google News from 2019–2023. In total, we get 60,671 noun and verb phrases with a Part-of-Speech Tagger via SpaCy (en_core_web_sm) that we treat as neologism candidates. We use an old dataset of 80,071 neologisms obtained from two slang dictionaries (Zhu and Jurgens, 2021) and sample 200 neologisms with interesting or no trend lines. Table 7 provides the breakdown of method overlap between each method pair in NEO-BENCH. Instead of the sample of 1,100 data points, we compare a total of all 6,908 words tweeted out by the NYT First Said bot from 2020 to 2023 with the other methods. Figure 11 provides the breakdown of NEO-BENCH by collection method and linguistic type.

Dataset	Neologism Type				Collection Method			
	Emerging?	Multiword?	Semantic?	Generalized?	Exclusion Lists	Dictionaries	Time-Series	Templates
(Pinter et al., 2020)	✗	✗	✗	✓	✓	✗	✗	✗
(Kerremans et al., 2018)	✗	✗	✗	✓	✓	✗	✗	✗
(Zalmout et al., 2019)	✗	✗	✗	✓	✓	✗	✗	✗
(Janssen, 2012)	✗	✗	✓	✗	✓	✗	✗	✗
(Dhuliawala et al., 2016)	✗	✓	✗	✗	✗	✓	✗	✗
(Zhu and Jurgens, 2021)	✗	✓	✗	✓	✗	✓	✗	✗
(McCrae, 2019)	✗	✓	✗	✓	✗	✓	✗	✗
(Broad et al., 2018)	✓	✗	✗	✓	✗	✗	✓	✗
(Li et al., 2021)	✗	✓	✗	✗	✗	✗	✓	✗
NEO-BENCH (this work)	✓	✓	✓	✓	✓	✓	✓	✓

Table 6: Comparison of English Neologism resources by the types of neologisms collected and the collection method used. NEO-BENCH covers more types of neologisms by using more methods than prior neologism datasets.

Source	Reddit	News	NYTimes	Dictionary
# Candidates	74542	60671	6908	80071
% Reddit	-	0.93%	0.94%	3.91%
% News	0.76%	-	0.26%	0.12%
% NYTimes	0.09%	0.03%	-	0.15%
% Dictionary	4.19%	0.16%	1.80%	-
% Total	5.04%	1.12%	3.00%	4.18 %

Table 7: Number of shared neologism candidates for each method pair. The overlap is reported as a percent of the total number of candidates for each method.

B.1 Google Trends Filtering

We collect Google Trends monthly data from January 2010 to July 2023. While Google Trends provides data from 2004, there are inconsistencies in word usage frequencies until 2010. To compare word prevalence between neologisms, we make a pairwise comparison of a neologism candidate with the misspelling ‘dangrous’, which provides a consistent baseline comparison for word usage data. We then use this normalized trend line for neologism candidate filtering.

In total, we create five differing methods that use a combination of filtering criteria, including curve-fitting, argmax detection, integral, line of best-fit, and maximum trend data values, to evaluate words as neologism candidates. We select the best combination based on which method yields both high precision and estimated recall in collecting neologisms. Using 20,000 words collected in February 2022, we filter this set through all five methods which filter out almost 90% of words. We sample each method for 100 candidates and manually annotate the samples for neologism classification, obtaining a sampled precision of each method. We combine all the neologisms from the manually-annotated samples to obtain a computationally derived neologism set. We evaluate each method for its estimated recall based on the proportion of words from the computational neologism

sample that is not filtered out. The sample precision is particularly low given the sparsity of neologisms that appear at a specific point in time, so we select the method with the highest precision of 0.2 and an estimated recall of 0.625 to reduce the amount of manual annotation required.

B.2 Dataset Analysis

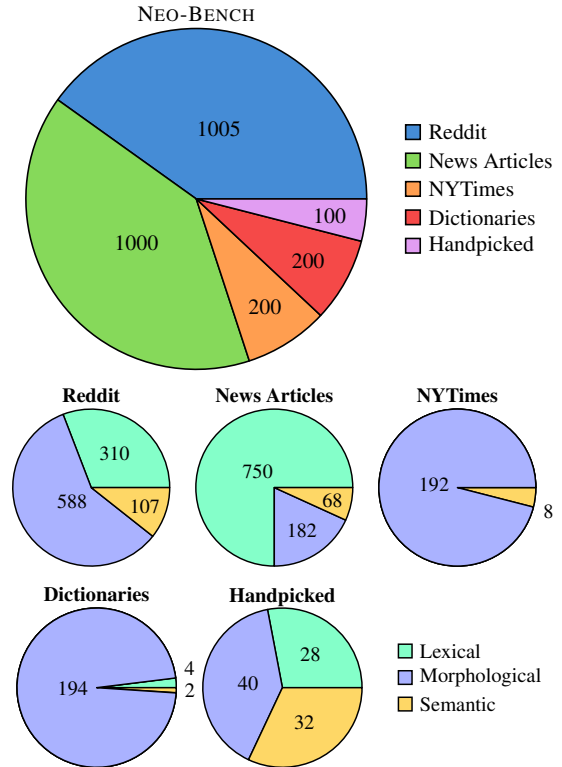


Figure 11: Breakdown of NEO-BENCH by collection method. Each method is further stratified by the linguistic type of neologisms.

With the best-performing filter method, we also estimate its recall with a set of 100 handpicked neologisms in our dataset. The estimated recall of our method is 0.55. This estimate is slightly lower than the estimated recall we used with the

neologism candidates we computationally gathered, but our collection methods remain consistent.

Analyzing the overlap of words in our dataset with Urban Dictionary by collection method, we find that 44.37% of Reddit neologisms, 38.4% of News neologisms, and 65.25% of neologisms obtained from dictionaries and exclusion lists overlap.

B.3 Emergence Date Labeling for Perplexity-Based Rankings

Using Google Trends to record the month and year where word usage spikes, we estimate the date for when a neologism enters the dissemination stage of its lifecycle. We find that 68% of neologisms emerged during 2020-2023, and 17% of all words have no dissemination date or trend line. The remaining words were prevalent before 2020, but a new connotation or usage has recently emerged. These dates are potentially inaccurate as a trend line is collapsed into a single date. Compared to using the entire graph to evaluate long-term word trends when filtering for neologisms, a single date may not perfectly capture neologism growth for words that exhibit a steady rise in growth.

B.4 NEO-BENCH Tasks

We sample 750 neologisms from our set of 2505 total neologisms. We stratify the sample based on the 5 collection methodologies used: handpicked, Reddit, News articles, NYT First Said Bot, and Slang dictionaries. We collect 500 neologisms from Reddit and News Articles, 100 from NYT First Said Bot, 100 from slang dictionaries, and 50 from our handpicked set.

We work with 3 in-house native English speakers to construct inputs for Cloze question answering, definition generation, and human-evaluated machine translation. Using a script, annotators are provided with Google Search Descriptions containing neologisms. With these reference sentences, annotators were instructed to create 750 multi-sentence Cloze passages and to create 4 multiple-choice options based on topical relevance to the neologism and 1 distractor answer that is a feasible substitute in the passage. Annotators also constructed 750 sentences containing neologisms and 750 questions asking for the definition of neologisms. From a subsample of 100 neologisms, annotators constructed minimal pair sentences using Google Search references for human-evaluated machine translation. Annotators were instructed to create any feasible substitute word, regardless of topic, to replace the

neologism in the minimal pair sentence. For perplexity rankings, we use 422 Cloze passages that have both a single-word neologism and distractor answer.

We additionally subsample 240 sentences containing neologisms for automatic machine translation. We work with 3 in-house native Chinese speakers to construct 240 reference translations for automatic MT evaluation. Annotators were provided the input sentence, the neologism, and the definition and were instructed to create fluent and accurate translations. If no associated term in Chinese exists, annotators were instructed to paraphrase sentences or use substitute terms with minimal information loss. Annotators were then instructed to retranslate the Chinese sentence back into English to enable further study on accurate translation techniques of neologisms, including paraphrasing the sentence. All 6 annotators we work with have a college-level education.

C Experimental Details

All models are evaluated on two NVIDIA A40 GPUs for a single run since models are not fine-tuned for NEO-BENCH tasks.

C.1 Machine Translation

For human-evaluated MT, our error categorization is partially adapted from the widely-used MQM framework (Freitag et al., 2021), which is adjusted based on the pilot studies we conducted on translating sentences containing neologism words. Translation errors are ordered by severity in affecting the understanding of the sentence, and annotators label sentences based on the most severe translation error. For instance, if there are grammatical mistakes and English words in the translation output, the output will be labeled “Copy” since that is the most severe translation error. Table 12 provides example outputs for each translation category.

We crowdsource annotations by screening for 5 fluent Mandarin speakers on Prolific⁶, a crowdsourcing website. All 5 annotators are fluent in English and Mandarin and reside in the United States and United Kingdom. All of the annotators have a college education and are informed about the nature of the study. Each annotator was given the same set of neologism sentences across all 5 models evaluated to ensure a standard comparison between models. Annotators were provided with

⁶<https://www.prolific.com>

neologism definitions and were instructed to select the label corresponding to the worst translation error in the sentence and highlight the corresponding spans in the input sentence and MT output. Annotators were also instructed to label the error severity and the confidence in their selection of translation categories on 3-point Likert scales. Finally, annotators marked if the translation error occurred from the neologism or from another portion of the sentence. The average time to annotate 20 minimal pairs was 80 minutes, and each annotator was paid \$12.00 an hour, which is on the high end for standard pay on Prolific. We use the Thresh (Heineman et al., 2023) interface for annotating translation sentences, and Figure 13 provides a screenshot of the interface.

Based on our human reference translations of neologisms, we find that translation is difficult. For neologisms representing new concepts, there may be information loss if there is not an associated word in the target language (e.g. boyflux $\rightarrow$ non-binary). There are often no exact words associated with neologisms in the target language, so translation requires providing the neologism definition and rephrasing the entire sentence (e.g. fossilflation $\rightarrow$ fossil fuel price increases). However, since neologisms are created with the same linguistic origins as common words, the novelty of neologisms results in lower MT performance.

C.2 Rank Classification with Perplexity

We also tested T5 and Flan-T5 for perplexity ranking and find that Flan-T5 exhibits higher neologism rankings than T5. However, when sorting words by lowest perplexity and filling in the mask, we find that these models produced entirely incoherent sequences, so we do not report these models. Given the computational intensity of evaluating 5,002 sequence perplexities, we only evaluate the base size of models.

C.3 Cloze Question Answering

Rank classification is used to select the lowest perplexity answer for BART, T5, and GPT-J. For other models, we shuffle the order of answers and conduct experiments with 5-shot prompting with the following format:

Fill in the blank with the options below:
 Question: [EXAMPLE CLOZE PASSAGE]
 a) [EXAMPLE INCORRECT ANSWER]
 b) [EXAMPLE DISTRACTOR ANSWER]

c) [EXAMPLE INCORRECT ANSWER]
 d) [EXAMPLE INCORRECT ANSWER]
 e) [EXAMPLE NEOLOGISM ANSWER]
 f) [EXAMPLE INCORRECT ANSWER]
 Answer: e) [EXAMPLE NEOLOGISM ANSWER]

...

Fill in the blank with the options below:

Question: [TEST CLOZE PASSAGE]

a) [TEST INCORRECT ANSWER]
 b) [TEST INCORRECT ANSWER]
 c) [TEST NEOLOGISM ANSWER]
 d) [TEST DISTRACTOR ANSWER]
 e) [EXAMPLE INCORRECT ANSWER]
 f) [EXAMPLE INCORRECT ANSWER]

Answer:

C.4 Definition Generation

We conduct experiments with 5-shot prompting with the following prompt:

Answer the question.

Question: [EXAMPLE DEFINITION QUESTION]

Answer: [EXAMPLE DEFINITION ANSWER]

...

Answer the question.

Question: [TEST CLOZE PASSAGE]

Question: [TEST DEFINITION QUESTION]

Answer:

One of the paper’s authors manually annotates 100 outputs. We measure the Cohen’s Kappa between automatic GPT-4 and human evaluation, and we obtain an average Cohen’s Kappa of 0.744, indicating high agreement between human judgment and GPT-4.

For automatic evaluation, we additionally use GPT-4 to determine if a correct model definition is better or worse than the reference definition provided by human input. For incorrect answers, we separate between incorrect and omitted generations, which are model outputs that are either left blank or, for GPT models, outputs where the model acknowledges that it does not recognize the neologism.

Table 13 provides the full results of the open-domain question-answering experiments, including the average length of definitions, manual evaluation and model-wise Cohen’s Kappa. While LLaMA-2 70B outperforms GPT-3.5 in Cloze QA, GPT 3.5 produces more correct definitions than LLaMA-2-70B. Instruction-tuned models produce a higher

Model	Pre-trained		
	Lexical	Morphological	Semantic
BART-Large	19.88	18.48	14.94
T5-Large	33.54	39.88	39.08
GPTJ 6B	34.78	25.22	36.78
LLaMA-1 7B	29.50	27.86	29.89
OLMo 7B	26.40	28.45	28.74
Mistral 7B	69.57	65.69	70.11
LLaMA-2 7B	67.39	58.36	62.07
LLaMA-2 13B	69.25	61.58	62.07
LLaMA-2 70B	77.95	71.55	<u>77.01</u>
GPT 3.5	69.25	69.79	<u>77.01</u>

Model	Instruction-Tuned		
	Lexical	Morphological	Semantic
Flan-T5 Large	41.99	42.99	41.61
Alpaca 7B	50.37	44.87	42.99
OLMo Instruct 7B	60.45	47.99	55.07
Mistral Instruct 7B	61.18	57.16	51.62
LLaMA-2 Chat 7B	62.64	57.42	62.76
LLaMA-2 Chat 13B	68.63	61.58	62.07
LLaMA-2 Chat 70B	71.43	64.22	72.41
GPT-4	83.85	80.06	87.36

Table 8: Neologism accuracies of models for the Cloze Question Answering task separated by linguistic type: lexical (322), morphological (341), and semantic (87). Best performing accuracy is presented in **bold**, while highest shared accuracy is reported in underline.

proportion of correct answers that are deemed better than the human reference sentence. We report the average length of generations for each model and conclude that GPT-4 prefers instruction model outputs because the human reference sentences are on average 19.20 words long, which is more concise compared to the more elaborative responses of instruct-tuned models. We find that 81.07% of preferred answers across all models are longer than the alternative correct answer. Table 14 provides examples of instruct-tuned model responses that are evaluated as better than the human reference and examples of GPT-omitted responses. Even when prompted with shortened answers, instruction-tuned models produce longer-form responses. We did not test for restraining the length of the model output as the complexity and reference definition length for each neologism varies extensively. Instruction-tuned models are less likely to produce omitted answers, and only GPT-3.5 and GPT-4 have generations that acknowledge that a term is unrecognized.

D Linguistic Taxonomy

Table 8 provides the linguistic breakdown of neologism accuracies for each model in Cloze QA. Table 11 provides the stratified results of definition generation by linguistic type of neologism. Tables 9 and 10 provides the breakdown of automatic machine translation evaluation, and Figure 12 provides the

stratified results of manually labeling translations by linguistic type. We report each category as a proportion to the total amount of neologisms for a certain linguistic type. Model performance discrepancy between the open source models and commercial systems is highest for lexical neologisms. For automatic metrics, lexical neologisms do not yield higher BLEU scores but yield higher COMET and COMETKiwi scores. Morphological neologisms yield 5.3 lower BLEU scores, 3.4% lower COMET scores, and 4.4% lower COMETKiwi scores than lexical neologisms. BLEU scores for semantic neologisms are high as there is high token overlap between different senses of the same word form. However, semantic neologisms often yield similarly low COMET and COMETKiwi scores as Morphological neologisms.

Model	MetricX-23 _{XXL}			MetricX-23-QE _{XXL}			MetricX-23 _{XL}			MetricX-23-QE _{XL}		
	Lex.	Morph.	Sem.	Lex.	Morph.	Sem.	Lex.	Morph.	Sem.	Lex.	Morph.	Sem.
Google Translate	1.494	2.087	1.864	1.126	1.762	1.180	1.642	2.145	2.099	1.742	2.129	1.867
Bing Translator	1.889	2.725	2.821	1.265	2.011	1.711	1.999	2.501	2.722	1.987	2.450	2.278
DeepL Translator	1.507	1.892	2.089	1.042	1.455	1.220	1.815	1.846	2.282	1.648	2.195	1.924
GPT-4	1.428	1.612	1.667	1.275	1.532	1.520	1.739	1.776	1.978	1.804	2.278	2.297
GPT-3.5	1.197	1.995	2.095	1.035	1.779	1.606	1.528	1.915	2.109	1.656	2.481	2.430
ALMA-7 B	2.025	2.553	2.758	1.710	2.246	2.233	1.950	2.489	2.284	2.172	2.719	2.413
M2M100 1.2B	2.764	3.917	3.779	2.204	3.370	2.704	2.647	3.553	3.138	2.487	3.373	3.004

Table 9: MetricX-23 and MetricX-23-QE scores of Machine Translation models evaluated on neologisms, separated by linguistic type of neologisms and aggregate score.

Model	COMET			COMETKiwi			BLEU		
	Lex.	Morph.	Sem.	Lex.	Morph.	Sem.	Lex.	Morph.	Sem.
Google Translate	0.870	0.842	0.849	0.820	0.782	0.805	0.530	0.487	0.507
Bing Translator	0.852	0.806	0.812	0.812	0.769	0.786	0.484	0.418	0.467
DeepL Translator	0.856	0.833	0.833	0.823	0.792	0.814	0.434	0.373	0.429
GPT-4	0.866	0.845	0.852	0.814	0.782	0.776	0.466	0.414	0.491
GPT-3.5	0.859	0.833	0.820	0.824	0.773	0.769	0.425	0.365	0.425
ALMA-7 B	0.814	0.795	0.786	0.770	0.731	0.730	0.303	0.262	0.287
M2M100 1.2B	0.816	0.743	0.774	0.786	0.715	0.730	0.357	0.310	0.361

Table 10: COMET and BLEU scores of Machine Translation models evaluated on neologisms, separated by linguistic type of neologisms and aggregate score.

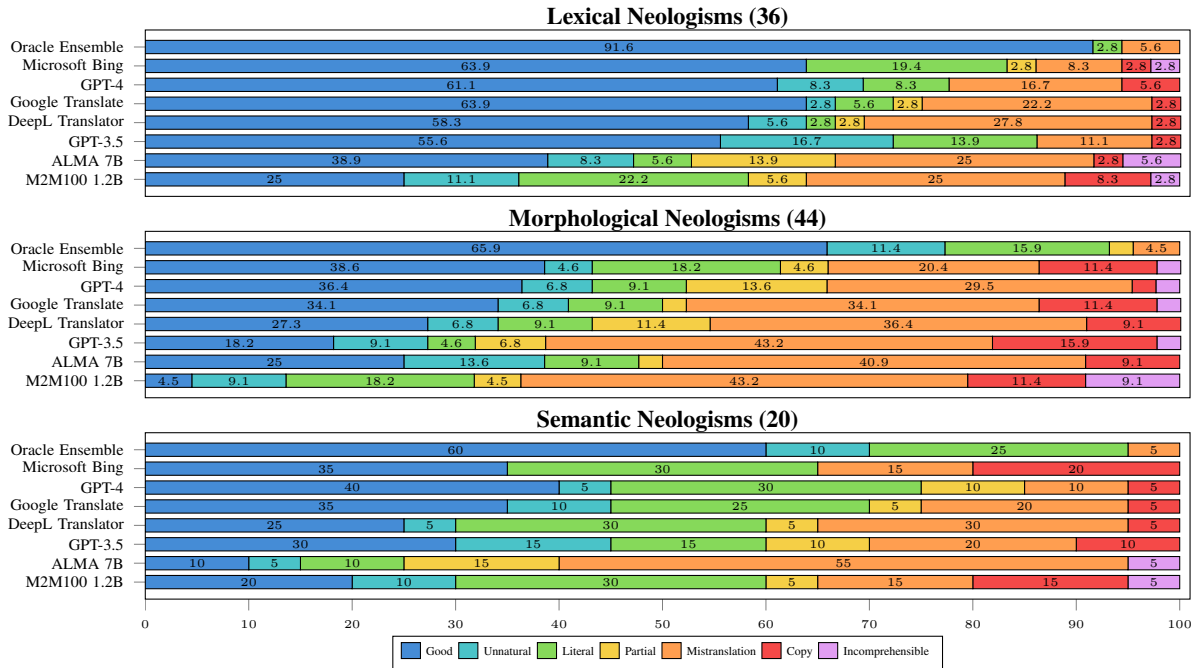


Figure 12: Results of the Machine Translation task with human-annotated labels for each linguistic type of neologism. Results are reported as percentages of the total number of neologisms of each linguistic category (provided in the titles). A Human Oracle Ensemble selecting the best model translation for each sentence is provided.

Pre-trained Models	Lexical (322)			Morphological (341)			Semantic (87)		
	Correct	% Worse	% Better	Correct	% Worse	% Better	Correct	% Worse	% Better
GPTJ 6B	0.367	78.7%	21.3%	0.323	68.1%	31.9%	0.322	78.6%	21.4%
LLaMA-1 7B	0.506	72.9%	27.1%	0.437	68.4%	31.6%	0.345	63.2%	36.8%
OLMo 7B	0.674	52.1%	47.9%	0.557	49.5%	50.5%	0.529	58.7%	41.3%
Mistral 7B	0.661	59.6%	40.4%	0.628	62.6%	37.4%	0.494	69.8%	30.2%
LLaMA-2 7B	0.587	65.1%	34.9%	0.481	58.0%	42.0%	0.506	72.7%	27.3%
LLaMA-2 13B	0.609	62.7%	37.3%	0.510	62.2%	37.8%	0.425	64.9%	35.1%
LLaMA-2 70B	0.665	57.4%	42.6%	0.613	54.0%	46.0%	0.483	57.1%	42.9%
GPT 3.5	0.817	6.1%	93.9%	0.686	9.8%	90.2%	0.713	11.4%	88.6%
Instruct-tuned Models	Lexical (322)			Morphological (341)			Semantic (87)		
	Correct	% Worse	% Better	Correct	% Worse	% Better	Correct	% Worse	% Better
Flan-T5 Large	0.158	96.2%	3.8%	0.158	83.5%	16.5%	0.080	86.3%	13.7%
Alpaca 7B	0.565	71.0%	29.0%	0.463	70.8%	29.2%	0.414	63.8%	36.2%
OLMo Instruct 7B	0.804	41.7%	58.3%	0.581	34.3%	65.7%	0.471	31.7%	68.3%
Mistral Instruct 7B	0.767	47.8%	52.5%	0.581	39.4%	60.6%	0.471	29.3%	70.7%
LLaMA-2 Chat 7B	0.581	59.4%	40.6%	0.528	63.3%	36.7%	0.448	66.7%	33.3%
LLaMA-2 Chat 13B	0.649	40.7%	59.3%	0.543	41.1%	58.9%	0.494	39.5%	60.5%
LLaMA-2 Chat 70B	0.661	46.0%	54.0%	0.566	44.5%	55.5%	0.494	37.2%	62.8%
GPT-4	0.870	8.3%	91.7%	0.827	11.4%	88.6%	0.736	11.0%	89.0%

Table 11: Results of the definition generation task when separated by linguistic type. Accuracy is reported as a proportion of correct answers compared to the total number of neologisms of each linguistic type. The percentages of correct answers that are labeled as ‘worse’ and ‘better’ than the human reference sentence by GPT-4 are provided.

< Hit 1 / 40 >

Instructions

Source:

Definition: Compersion is vicarious joy associated with seeing one's partner have a joyful romantic or sexual relation with another.

In a romantic relationship, the opposite of jealousy is compersion, the feeling of happiness that arises when you see your partner's happiness.

Target:

在浪漫关系中，嫉妒的反面是共享喜悦，当你看到伴侣的幸福时产生的快乐的感觉。

ADDING AN EDIT +

Select the Edit Category..

Good Translation

Incomprehensible

Mistranslation

Unnatural

Literal Translation

Partial Translation

No Translation

CANCEL X

SAVE ✓

ADDING AN EDIT ✎

Mistranslation: compersion with 共享喜悦

Rate the impact severity that this has on understanding the translation and being an accurate translation:

1 - minor

2 - somewhat

3 - a lot

What's your confidence on your choice?

1 - minor

2 - somewhat

3 - a lot

Which part of the sentence is mainly mistranslated?

Neologism

Rest of Sentence

CANCEL X

SAVE ✓

Figure 13: Thresh Interface used to crowdsource human annotations of Machine Translation Output.

Translation Output Examples	
Fire Weather	Input: They will not issue official warnings until fire weather is forecast to occur. <i>(Fire weather is the use of meteorological parameters such as relative humidity, wind speed, mixing heights, and soil moisture to determine whether conditions are favorable for fire growth and smoke dispersion.)</i>
	Model Output (Good): 在预测发生火灾天气之前，他们不会发布官方警告。 <i>(Before predicting the occurrence of fire weather, they will not issue an official warning.)</i>
	Human Translation: 他们直到预测到火灾天气才会发布官方警告。 <i>(They will only issue official warnings when a fire weather is forecasted.)</i>
Dupe	Input: Discover new affordable dupes for luxury expensive makeup products. <i>(A dupe is an abbreviation of the word "duplicate".)</i>
	Model Output (Unnatural): 为奢华昂贵的彩妆产品发现新的负担得起的复制品。 <i>(Find new affordable replicas for luxurious and expensive makeup products.)</i>
	Human Translation: 探索奢侈昂贵化妆品的平价替代新品。 <i>(Explore new affordable alternatives for luxurious and expensive makeup products.)</i>
Snowvember	Input: When sleet started falling during Thanksgiving it was officially snowvember . <i>(Snowvember refers to a particular November that experiences a lot of snowfall.)</i>
	Model Output (Literal): 当感恩节期间开始下雨夹雪时，正式进入了十一月雪。 <i>(When it starts to rain and snow during Thanksgiving, it officially enters the November snow.)</i>
	Human Translation: 感恩节开始下的雨夹雪标志着‘雪月’的正式开始。 <i>(The start of sleet during Thanksgiving marks the official beginning of 'Snowvember'.)</i>
Trollbaiting	Input: Trollbaiting has caused my growth on social media this past month to be super high. <i>(Trollbaiting describes when an internet user knowingly invites the hatred of a highly reactionary group of trolls.)</i>
	Model Output (Partial): 在过去的一个月里，我在社交媒体上的增长速度超快。 <i>(In the past month, my growth on social media has been extremely fast.)</i>
	Human Translation: 过去一个月里，通过挑衅网络喷子，我在社交媒体上的增长极为迅速。 <i>(In the past month, by provoking internet trolls, my growth on social media has been extremely rapid.)</i>
Forever Chemicals	Input: The environment cannot break down forever chemicals , and they will remain in our bodies for years if ingested. <i>(Forever chemicals are used to make products grease-proof, water-proof, stick-proof, and stain-resistant and are toxic to humans and nearly indestructible.)</i>
	Model Output (Mistranslation): 环境不能永远分解化学物质，如果摄入它们，它们会在我们体内停留数年。 <i>(The environment cannot always break down chemicals, and if ingested, they can stay in our bodies for years.)</i>
	Human Translation: 环境无法分解永久化学物质，一旦摄入，这些物质将在我们的身体中残留多年。 <i>(The environment cannot break down 'forever chemicals', and once ingested, these substances will remain in our bodies for many years.)</i>
Blud	Input: What is blud talking about I can't understand. <i>(Blud is slang that is used to address men and means bro.)</i>
	Model Output (Copy): blud 在说什么我听不懂。 <i>(blud, I don't understand what you are saying.)</i>
	Human Translation: 我听不懂这哥们儿在说什么。 <i>(I can't understand what this guy is saying.)</i>
Noctor	Input: Is the narcissist in your life a noctor and diagnosing you? <i>(A noctor is a health professional (usually nurse) who takes on some traditional roles performed by the doctor.)</i>
	Model Output (Incomprehensible): 你生活中的自恋者是夜幕降临，诊断着你吗？ <i>(Is the narcissist in your life like nightfall, diagnosing you?)</i>
	Human Translation: 那个在你生活中的自恋者是不是冒充医生给你做诊断？ <i>(Is the narcissist in your life pretending to be a doctor and diagnosing you?)</i>

Table 12: Example model outputs for all possible translation categories. For each neologism example, the English input and Chinese output is reported. A gold reference definition of the neologism is provided. (Neologism definitions and English translations are shown for information only.)

	Model	GPT-4 Eval. (750)				Avg. Length	Acc. (↑)	Human Eval. (100)	
		Incorrect	Omitted	Worse	Better			Acc. (↑)	Cohen's κ (↑)
Pre-trained	GPT-J 6B	494	0	190	66	19.04	0.341	0.38	0.711
	LLaMA-1 7B	384	24	240	102	16.69	0.456	0.48	0.697
	OLMo 7B	297	0	234	219	19.24	0.604	0.64	0.729
	Mistral 7B	280	0	291	179	17.35	0.627	0.64	0.768
	LLaMA-2 7B	311	42	251	144	19.49	0.529	0.61	0.681
	LLaMA-2 13B	262	81	255	152	18.76	0.544	0.56	0.698
	LLaMA-2 70B	191	94	260	205	17.29	0.620	0.67	0.827
	GPT 3.5	95	96	46	513	41.21	0.745	0.72	0.828
Instruct	Flan-T5 Large	638	0	100	12	15.42	0.149	0.17	0.670
	Alpaca 7B	374	0	112	264	34.18	0.501	0.56	0.761
	OLMo Instruct 7B	297	0	144	309	27.37	0.604	0.68	0.737
	Mistral Instruct 7B	307	0	165	278	22.15	0.591	0.65	0.643
	LLaMA-2 Chat 7B	320	0	155	275	24.02	0.573	0.58	0.758
	LLaMA-2 Chat 13B	313	0	178	259	24.14	0.583	0.66	0.731
	LLaMA-2 Chat 70B	300	1	200	249	24.22	0.599	0.64	0.771
	GPT 4	106	18	62	564	38.09	0.835	0.85	0.891

Table 13: Full results of the Definition Generation task showing the number of **correct** and **incorrect** answers per model. A sample of 100 neologisms are manually evaluated, and Cohen’s Kappa is calculated to determine annotator agreement between GPT-4 and human evaluation. Model accuracy is reported for both manual and GPT-4 evaluation based on the number of definitions evaluated. Correct outputs are divided between being worse or better than the human reference sentences. Wrong definitions are divided between incorrect or omitted output where models do not provide a definition when prompted. Average length of definitions is reported with number of words.

a) Better Model Output Examples	
Inclusive Education	Reference Definition: Inclusive education refers to all students being able to access and gain equal opportunities to education and learning.
	Flan – T5 Large Output: Inclusive education is a philosophy that emphasizes the inclusion of all students, regardless of their disability, gender, race, religion, or sexuality.
Greenwashing	Reference Definition: Greenwashing is the act or practice of making a product, policy, activity, etc. appear to be more environmentally friendly or less environmentally damaging than it really is.
	Alpaca 7B Output: Greenwashing is the practice of misleading consumers about the environmental practices of a company or product. It involves claiming to be environmentally friendly without having any genuine commitment to sustainability.
TCG	Reference Definition: TCG is an acronym for "trading card game."
	LLaMA – 2 Chat 7B Output: TCG stands for Trading Card Game, a type of game where players use cards to represent characters, spells, and other elements to battle each other.
Mermaidcore	Reference Definition: Mermaidcore is an aesthetic focusing on imagery and fashion related to mermaids.
	LLaMA – 2 Chat 13B Output: Mermaidcore is a type of aesthetic that features pastel colors, sparkly accessories, and flowing fabrics, inspired by the mythical creatures known as mermaids.
Wokeness	Reference Definition: Wokeness is the quality of being alert and concerned about social injustice and discrimination.
	LLaMA – 2 Chat 70B Output: It is a term used to describe the state of being aware of and attuned to social injustices, particularly those related to racism, sexism, and economic inequality.
Longcovid	Reference Definition: Longcovid is the signs and conditions that continue or develop after acute COVID-19 infection.
	GPT – 4 Output: Long COVID, also known as post-acute sequelae of SARS-CoV-2 infection (PASC), refers to a range of symptoms that can continue for weeks or months after the acute phase of a COVID-19 infection has resolved. Symptoms may include fatigue, breathlessness, brain fog, depression and anxiety, pain, and other long-term complications.
b) Omitted Model Output Examples	
Falstalgia	Reference Definition: Falstalgia is a misguided feeling of nostalgia, a longing for a past that really wasn't that great and you weren't around for.
	GPT – 4 Output: I'm sorry, "falstalgia" does not appear to be a recognized term in English language or any other languages. It could be a typo. If you meant "nostalgia," it is a sentimental longing or wistful affection for the past, typically for a period or place with happy personal associations.
Flurona	Reference Definition: Flurona is a coinfection with influenza and COVID-19.
	GPT – 3.5 Output: I'm sorry, but I couldn't find any information on the term "flurona." It is possible that it is a newly coined term or a term that is not widely recognized.

Table 14: Example model outputs of the definition generation task. **a)** Instruction-tuned model outputs evaluated as better than the reference definition by GPT-4 and **b)** GPT model outputs that omit definitions are provided.