

Multilingual Structured Sentiment Analysis for Environmental Sustainability

Anonymous ACL submission

Abstract

To effectively address global environmental challenges, we must have tools that allow us to carefully monitor how citizens, policy makers and other stakeholders debate sustainability. However, there are currently very few NLP resources and tools specialized for this topic. This paper presents ENVIS, a multilingual corpus (Italian, English, and Indonesian) for investigating the debate on environmental sustainability in social media using Structured Sentiment Analysis. We introduce a framework for the automatic aggregation of span-level annotations that preserves the annotators' perspective and avoids manual intervention by safeguarding the quality of the annotations. We performed a series of experiments with four open-source instruction-based Large Language Models in zero-shot and few-shot settings, where we have measures the impact of the order and number of shots. The results show further confirm the ineffectiveness of LLMs in extracting fine-grained sentiment information, being outperformed by a supervised state-of-the-art neural method trained on very few data. This questions the suitability of LLMs for knowledge/information extraction tasks. Our error analysis indicates that LLMs mostly struggle in identifying the sentiment term or its associated polarity, failing to extract full sentiment triples.

1 Introduction

The development of specialized language resources and NLP tools to analyze the debate on the environmental crisis and its solutions is still at an early stage. Previous work has mostly taken a narrow view focusing on a single issue, i.e., climate change (Stede and Patz, 2021; Spokoiny et al., 2023; Mullappilly et al., 2023; Ni et al., 2023; Stambach et al., 2023). In this contribution, we take a broader perspective by analyzing Social Media messages in different languages (English, Italian, and Indonesian) covering multiple topics

related to **environmental sustainability** (ES) to better understand the public debate and contribute to the identification of potential areas of interventions (Kirilenko and Stepchenkova, 2014; Veltri and Atanasova, 2017).

To this end, we follow the paradigm of Sentiment Analysis (SA) as a proxy to identify, monitor, and analyze opinions of the public in a more natural setting (Liu, 2015; Zhang et al., 2018). However, we depart from the classical approach and adopt **Structured Sentiment Analysis** (SSA). SSA offers a more fine-grained analysis of the relationship between the holder (of an opinion), the sentiment (expressed by the opinion terms), and the eventual target, thus allowing for the decoding of multiple sentiment triplets in the same message. Our main contributions can be summarized as follows:

- we present ENVIS, a new multilingual annotated dataset for SSA (§3.1 and §3.2) together with a framework for the automatic aggregation of text spans (§3.3);
- we perform a series of experiments comparing four open-source instruction-tuned LLMs to a state-of-the-art to an encoder-based dependency graph parser (Zhai et al., 2023) and show that LLMs are unsuitable for this task (§4);
- we conduct a thorough error analysis showing that LLMs tend to overgenerate SSA tuples while mostly failing to identify sentiment terms and their polarity (§5).

2 Related Works

SA has evolved from assigning global sentiment values to more fine-grained annotations (Wiebe et al., 2005; Pontiki et al., 2014; Peng et al., 2020). It is now common to frame SA tasks as **Aspect-Based SA** (ABSA) (Xu et al., 2020; Barnes et al., 2022). ABSA requires systems to associate the correct *sentiment term* (also called *opinion term*) and

its polarity value to their specific *aspect/target*, usually expressed in the form of attribute/entity (Pontiki et al., 2014, 2015, 2016).

To address the subtask fragmentation of ABSA, **Structured Sentiment Analysis (SSA)** (Deng and Wiebe, 2015; Barnes et al., 2021) proposes a more holistic approach by jointly predicting all elements of a sentiment graph of an *opinion tuple* O . Each O contains four elements: the holder (h), the target (t), the sentiment term expression (e), and the polarity value triggered by the sentiment term (p). Early work has adopted pipeline approaches by extracting each subcomponents of an opinion tuple O and subsequently connecting them. Recent approaches leverage Transformer-based pre-trained language models (PTLMs) to enhance SSA performance, in combination with graph-based methods or multi-task learning (Lin et al., 2022; Chen et al., 2022; Zhai et al., 2023; Zhou et al., 2024). The recasting of the task as a dependency graph parsing problem by Barnes et al. (2021), where the sentiment terms are the roots connecting to their respective holders and targets, has led to the introduction of a new evaluation measure, namely **Sentiment Graph F_1** (SF_1). SF_1 captures the ability of a model to identify the full sentiment graph, rather than its components. In particular, a true positive is an exact graph-level match, calculated by averaging the weighted overlap of predicted and gold spans across all span types of the opinion tuple O .

The availability of datasets for SSA has expanded beyond English and single-text-type sources, such as the MPQA corpus (Deng and Wiebe, 2015) to many other languages and diverse text types, such as reviews covering various topics and social media messages (Agerri et al., 2013; Barnes et al., 2018; Øvrelid et al., 2020; Toprak et al., 2010; Bosco et al., 2023). Most of these datasets are now part of the SemEval 2022 shared task on SSA (Barnes et al., 2022), representing a reference benchmark. We aim at further expanding the variety of SSA datasets by modeling the ES debate in Social Media (Ibrohim et al., 2023). When compared to previous work on this topic, we specifically focus on extracting and evaluating opinions and associated polarity values rather than developing Question-Answering models (Spokoyin et al., 2023; Mullappilly et al., 2023; Ni et al., 2023), detection of claims (Stammach et al., 2023), and political debate framing (Stede and Patz, 2021).

Research on applying LLMs to SSA is limited. Zhou et al. (2024) evaluated ChatGPT-3.5-Turbo

and GPT-4 in zero-shot, 1-shot, and 5-shot settings, finding their performance inferior to an encoder-based fully supervised model. Zhang et al. (2024) conducted an extensive study comparing sequence-to-sequence models (Flan-T5-13B and Flan-UL2) against ChatGPT-3.5-Turbo and the text-davinci-003 model showing that, although not perfect, few-shot in-context learning is somewhat effective for encoder-based LLMs, reaching F1-score between 0.35 and 0.55 according to the dataset. Despite these limitations, LLMs offer potential, particularly in low-resource and multilingual scenarios like ours. We explore their application by analyzing the impact of shot order and quantity in in-context learning, assessing whether strategic prompting can help compensate for limited training data.

3 ENVIS: A Multilingual Corpus for SSA for Environmental Sustainability

ENVIS is the first multilingual SSA corpus on ES from Social Media messages. In this section, we describe how the corpus was collected, the annotation process, and the label aggregation process.

3.1 Data Collection

The starting point for ENVIS is the dataset from Bosco et al. (2023), with only Italian and English messages collected with the Twitter API. The collection is based on 13 keywords for Italian and 120 for English¹ covering 10 ES topics. The Italian subcorpus is composed of 8,756 tweets collected between February, 2nd and March, 4th 2022. The English subcorpus contains more than 490k messages collected between September, 12th and 30th 2022.

A third subcorpus extends ENVIS to Indonesian. Indonesian is the language spoken in one of the fastest growing economies of the Global South² with the 4th largest population in the world,³ where the debate surrounding ES could be meaningfully different when compared to the Global North. The lack of a universally shared definition of sustainable development opens indeed the debate around ES to various interpretations and perceptions between these two macro geo-political areas. English is here considered as a global *lingua franca* not specifically representing a single world area.

The Indonesian data were collected between February, 28th and January, 2nd 2024 with the new

¹The keywords lists are in Appendix A.

²<https://bit.ly/3HRO3rA>

³<https://worldpopulationreview.com/>

version of the Twitter/X API. We have manually translated the English keywords from [Bosco et al. \(2023\)](#) and added 31 keywords related to ES debate in the Indonesian for a total of 159 keywords resulting in 27,279 tweets. Since a preliminary check revealed that many messages were not relevant to ES (e.g. advertising, tweets about cooking and healthy lifestyle), we implemented a multi-stage filtering approach to improve the quality of the data. First, we dropped duplicate tweets due to multiple keywords. Then, we trained a classifier to distinguish whether a tweet is on topic (i.e., related to ES) or not. For this purpose, we have manually annotated 600 tweets,⁴ split them into train (80%) and test (20%) sets and fine-tuned IndoBERT ([Wilie et al., 2020](#)), a monolingual Indonesian PTLM. The classifier returned a macro $F_1 - Score$ of 89.49 at test time - showing that we can quite reliably run it to remove most of the noisy messages. After applying our classifier to the remaining 13,228 collected tweets, we retained 2,300 messages for manual annotation.

3.2 Annotation and Agreement

[Bosco et al. \(2023\)](#) introduce an SSA annotation scheme with four markables: holder (h), target (t), sentiment term (e), and topic (tp). Targets and holders are classified as individuals or organizations, and topics are limited to 10 environmental categories. Sentiment terms are marked only as positive or negative, with neutral sentiments excluded. Relations are annotated as (h, t, e, p, tp) tuples.

In ENVIS, we simplify this scheme by removing the topic annotation and expanding the dataset. For Italian, we introduced a third annotator—a native-speaking Linguistics master’s student—to align with [Bosco et al. \(2023\)](#)’s annotators. This addition helps resolve disagreements and consolidate annotation spans.

For English, the original annotation covered only sentiment expressions in 700 tweets. We expanded this subcorpus to 1,500 messages, including all markables. Using the same participant selection criteria as ([Bosco et al., 2023](#)), we recruited annotators via Prolific.⁵ Three annotators completed the task in two phases: first identifying sentiment terms and their polarity, then marking holders and targets. To ensure quality, we incrementally added

new participants until Fleiss’ κ exceeded 0.4 for each markable ([Landis and Koch, 1977](#)). Holder and target annotations were based on automatically aggregated sentiment spans (§3.3).

For Indonesian, three native speakers were recruited and trained using a translated version of the English guidelines, as Prolific lacked native speakers. They annotated messages following the same two-step strategy as in English, with quality controlled by requiring a Fleiss’ κ score above 0.4.

We computed pairwise Span Cohen’s κ ([Øvrelid et al., 2020](#)) based on proportional span overlap for each annotation layer. Agreement levels are generally consistent across languages despite varying annotator expertise. The holder label shows the lowest agreement (fair to moderate), aligning with [Barnes et al. \(2018\)](#), while other labels achieve moderate to substantial agreement ($\kappa = 0.40$ – 0.67). This reflects the task’s subjectivity and challenge, with most disagreements stemming from span boundary variations rather than differing annotations. Detailed scores are in Appendix B.

Table 1 shows the annotation statistics for the three subcorpora in a disaggregated format. Since SSA annotation is based on spans, disagreements are mostly due to inconsistencies in the length of the spans rather than choices of completely different spans. As the figures show, across all languages, the annotators share the same characteristics in terms of span length and number of spans annotated. As expected, the sentiment expressions have spans (between 2.39 and 4.90 tokens) longer than holders and targets (less than 2 tokens), usually corresponding to proper nouns. Italian, the only pro-drop language, has a remarkable lower number of holders spans being usually encoded in the verb morphology. Italian annotators are those who show the least variation in the number of annotated spans across all markables. English and Indonesian, on the contrary, show greater variation - usually clustered around one annotator. This behavior is clearly due to differences in expertise of the annotators pools.

3.3 Final Label Aggregation

Although disaggregated data are gaining popularity in NLP ([Plank, 2022](#); [Basile et al., 2021](#)), for our annotation strategy and evaluation purposes we need to settle on an aggregated version of the opinion tuples (h, t, e, p) . Considering the task at hand, the aggregation process first focuses on the sentiment term, and then on the target and holder. At the

⁴The annotator is a native speaker and a Master’s student in computer science.

⁵Only native English speakers with a 100% work acceptance rate were selected and paid £9 per hour.

Statistic		ENVIS-IT			ENVIS-EN			ENVIS-ID		
		A1	A2	A3	A1	A2	A3	A1	A2	A3
# holder		43	46	62	438	396	269	274	86	174
# target		505	538	547	1318	1331	1400	2848	1609	1198
# negative term		634	491	652	1706	1639	1741	1482	1635	1635
# positive term		517	535	488	956	872	1005	1333	1905	2067
avg. span length (# token)	holder	1.47	1.46	1.39	1.30	1.34	1.38	1.77	1.48	1.45
	target	1.59	1.70	1.31	1.96	1.95	2.07	2.56	2.18	2.17
	neg. term	2.44	2.91	3.70	3.28	4.90	4.29	3.48	2.36	2.35
	pos. term	1.54	2.09	2.84	3.31	4.46	4.01	3.34	2.36	2.21
# tweets no holder		965	957	941	1125	1156	1278	2032	2215	2133
# tweets no target		616	565	539	570	563	481	483	946	1350
# tweets no sentiment term		268	305	137	155	201	130	576	360	424

Table 1: Disaggregated annotation statistics of ENVIS for each language.

Aggregation	ENVIS-IT					ENVIS-EN					ENVIS-ID				
	<i>h</i> (%)	<i>t</i> (%)	<i>e</i> (%)	Avg. (%)	Std.	<i>h</i> (%)	<i>t</i> (%)	<i>e</i> (%)	Avg. (%)	Std.	<i>h</i> (%)	<i>t</i> (%)	<i>e</i> (%)	Avg. (%)	Std.
MVToken	97.65	92.21	75.88	88.58	11.33	85.88	85.36	57.17	76.14	16.43	93.55	73.5	60.21	75.75	16.78
MVSeq	97.55	92.55	79.16	89.75	9.51	85.94	84.83	59.80	76.86	14.78	93.57	72.33	59.48	75.13	17.22
MVSeg	97.50	92.13	71.58	87.07	13.68	85.92	84.69	59.35	76.65	15.00	93.61	72.25	59.48	75.11	17.24
MVOver	96.30	89.36	62.35	82.67	17.94	85.67	81.08	33.23	66.66	29.04	93.35	66.63	48.07	69.35	22.76
MVUnion	89.13	84.41	56.86	76.80	17.43	84.57	77.98	38.96	67.17	24.65	90.63	61.12	46.38	66.04	22.53

Table 2: PCR for holder (*h*), target (*t*), and sentiment term expression (*e*) for each language. The best average for each language in bold.

same time, our aggregation process must be independent of the specific tags, robust to avoid manual intervention, and correct, i.e., each aggregated tag is valid. An intrinsic limitation of automatically aggregating data is that there will be no annotation if there is a complete disagreement. Based on our annotation strategies, this impacts each language differently: messages may lack sentiment terms (all of them), or have a sentiment term and no holder and/or target (English and Indonesian).

Evaluating automatically aggregated spans is particularly challenging. Rodrigues et al. (2014) propose various aggregation methods by comparing them to expert annotations, an approach not feasible in our case. Additionally, aggregated spans may result in non-grammatical phrases - causing further ambiguity for the relation annotations (and their automatic identification). To address this issue and assess the quality of the aggregations, we introduce a new evaluation measure, the **Phrase Completeness Ratio** (PCR). PCR is formulated as the ratio between the number of aggregated spans that correspond to a grammatical phrase and the total number of aggregated spans. A grammatical phrase is any token combination directly connected via a dependency relation to its parent node. The parent token is considered a single token phrase. To illustrate how PCR works, consider the following message and three different annotations of the sentiment term:

- (1) *Internationally uncompetitive energy prices cause industrial production to shift.*
(Anno. 1) uncompetitive
(Anno. 2) uncompetitive energy
(Anno. 3) uncompetitive energy prices

From the dependency parsing,⁶ we obtain 13 grammatical phrases. Focusing only on the first part of the sentence, i.e., until the verb “cause”, the list of grammatical phrases is the following: {internationally uncompetitive, uncompetitive, uncompetitive prices, energy prices, uncompetitive energy prices, prices}. Using an aggregation method based on majority voting at token level, we would obtain the phrase “*uncompetitive energy*” as a candidate sentiment term, with has no matches in the list of grammatical phrases. This will result in a PCR score of 0. On the other hand, if we aggregate by taking the union of all tokens, the resulting phrase would be “*uncompetitive energy prices*”, with a match to our list of grammatical phrases. This will result in a PCR score of 1.0 (one matching phrase divided by one aggregated span). The global score for each aggregation method over each language-specific subcorpus of ENVIS is calculated by computing the average of the PCR score of each message. The pseudo-code for PCR

⁶To obtain the dependency tree, we have used the SpaCy library v3.7.

is in Appendix C.

Overall, five different aggregation methods have been evaluated. The first three (*MVToken*, *MVSeq*, *MVSeg*) are a reimplementation of the baselines in [Rodrigues et al. \(2014\)](#). Note that for the sentiment term, the aggregation also considers the associated polarity value.

MVToken Majority voting at the token level, i.e., the tokens with the most votes (and same polarity for sentiment terms).

MVSeq Majority voting at the sequence level by considering the exact match of a sequence of tokens (and same polarity for sentiment terms).

MVSeg Majority voting over segment level. The aggregation takes place by majority in two steps: first at the segment level, then at the token level like in *MVToken*. Note that with two annotators this measure will produce the same output as *MVSeq*.

MVOver Majority voting by considering the union of overlapping token sequences (including same polarity for the sentiment terms); this method is useful for capturing long phrases.

MVUnion The maximum span sequence of partially overlapping tokens (including the the same polarity for sentiment terms).

Table 2 summarizes the results of each aggregation method per language and per markable. The PCR of sentiment term expressions is generally lower than the holder and target and subject to larger variability across the aggregation methods. This is due to the higher number of tokens involved (and thus more prone to disagreements). *MVOver* and *MVUnion* have the lowest PCR scores: this is expected since they take the longest span, which may overlap over multiple phrases. On the basis of these results, we selected one aggregation method per language - rather than per markable. The selection has been done on the basis of the method returning the highest PCR average across the three markables. This results in *MVSeq* for the aggregation of Italian and English and *MVToken* for Indonesian. Table 3 summarizes the final aggregated annotations statistic of the ENVIS corpus.

As Table 3 shows, Italian is the only language with an almost perfectly balanced distribution between positive and negative sentiment terms. In English, negative sentiment terms predominate, whereas this trend is less pronounced in Indonesian. As for the sentiment term length, Italian has

Corpus Data	ENVIS-IT	ENVIS-EN	ENVIS-ID
# holder	26	333	144
# target	483	2481	3127
# negative term	679	2444	3039
# positive term	563	1180	2783
avg. span length holder	1.06	1.14	1.19
avg. span length target	1.14	1.16	1.17
avg. span length neg. term	2.53	3.14	2.34
avg. span length pos. term	1.61	3.17	2.29
# tweets no holder	980	1251	2204
# tweets no target	656	595	904
# tweets no sentiment term	291	234	479
# tweets - total	1,000	1,500	2,300

Table 3: Data overview of the aggregated ENVIS dataset.

relatively short sentiment spans, while the length for English and Indonesian is comparable. In general, the negative sentiment terms are longer than the positive ones, due to the fact that in many cases the presence of an explicit negation is used to revert the polarity. Holder and target have the same average length across all languages. Italian has 29.10% of messages with no sentiment terms, while this drops to 15.60% for English and 20.83% for Indonesian. On average, 25.61% of the messages across all the languages result with no sentiment terms after the aggregation process because of disagreements, a percentage that reaches 32.05% in English. This corresponds to 8.86% for the holder markables and 20.89% for the targets. Detailed results are reported in Table E in Appendix D.

4 ENVIS Dataset Benchmark

We benchmark ENVIS against four open-source, instruction-tuned LLMs, each with 7B to 8B parameters, to assess how effectively small LLMs can solve SSA tasks related to environmental sustainability through in-context learning. Considering the overall size of ENVIS, 4,800 messages in total, we have reserved 80% of the data for testing and only 20% for training. This setup allows us to explore how well these models, which are typically less resource-intensive than larger models, can handle specialized tasks such as SSA in low-resource settings. To provide a comparative baseline, we also contrast our LLM results with those of the USSA model ([Zhai et al., 2023](#)), a trainable dependency parsing graph initialized with encoder-based representations, which has achieved state-of-the-art performance on the SemEval 2022 SSA shared task. For the evaluation, we have used SF_1 ([Barnes et al., 2022](#)). All experiments have been conducted on one NVIDIA A-100 GPU.

Subcorpus	Model	Zero-shot	1-set (C1)	1-set (C2)	2-sets (C1)	2-sets (C2)	3-sets (C1)	3-sets (C2)
ENVIS-IT	Mistral-7B	1.56	8.11	7.65	9.40	10.31	10.22	11.23
	Llama-3-8B	3.13	8.46	8.86	10.21	9.66	11.77	10.57
	Llamantino-3-8B	3.60	7.61	7.56	7.25	8.10	8.94	8.29
ENVIS-EN	Mistral-7B	0.21	3.76	4.28	3.21	4.19	3.36	4.17
	Llama-3-8B	0.57	4.32	3.18	4.39	3.23	4.58	4.68
ENVIS-ID	Mistral-7B	2.07	5.89	6.29	7.22	7.10	7.02	7.54
	Llama-3-8B	2.64	4.86	5.29	4.59	5.23	5.25	5.35
	SahabatAI-8B	3.24	5.48	5.31	5.48	5.65	6.25	5.64

Table 4: ENVIS LLMs baselines. Best scores per language in bold. All scores correspond to SF_1 .

LLMs Benchmarking We selected these four models in their instruction-tuned versions: Mistral-7B (Jiang et al., 2023), Llama-3-8B (Dubey et al., 2024), Llamantino-3-8B (Polignano et al., 2024) for Italian, and SahabatAI-8B for Indonesian.

Llamantino-3-8B and SahabatAI are language-specific retrained versions of Llama-3-8B. Llamantino-3-8B uses QLORA (Dettrmers et al., 2023) strategy with 4-bits precision and an automatically translated DPO dataset.⁷ SahabatAI is fine-tuned on full parameter setting using Indonesian and two regional languages (Javanese and Sundanese).

We have experimented using both the zero-shot and few-shot settings within the in-context learning paradigm. The zero-shot setting helps assess the internal knowledge of the models, while the few-shot experiments allow us to evaluate how sensitive these smaller models are to variations in the number and order of shots (Song et al., 2025). In this way, we aim to understand how the models adapt to different in-context learning configurations and whether optimizing shot order and quantity can enhance their performance on SSA tasks.

For the selection of the examples, we proceed in blocks of four shots. The basic setting, called 1-set, contains four shots as follows: one example where all annotated sentiment terms are positive, one example where all annotated sentiment terms are negative, one example with no sentiment term and one example with mixed sentiment terms. For the other two settings, we have increased the number of shots per label of sentiment terms. This means that for 2-set, we have two instances per sentiment term type, for a total of 8 shots, and for 3-set we have three, for a total of 12 shots. To test the impact of the order of the shots, we compare two contexts. In Context 1, the examples are

given per group of labels related to the sentiment terms: first all positive cases, followed by negatives, neutral, and mixed. In Context 2, the same examples are presented in a randomized order.

As SSA is a complex task, we have carefully designed our prompt that consists of the task description, the instructions to constraint the output format, additional constraint in case where there is no holder, target, and/or sentiment term expression. We have used greedy search as the token generation method and set to 250 the maximum number of generated tokens (max_new_token). Our prompts are in Appendix E.

The LLMs results are reported in Table 4. In no circumstance, the LLMs refused to give an answer. However, in some cases (2.97% on average), LLMs completely failed to follow the instructions and offer explanations for their answers, a behavior already observed in previous work (Han et al., 2023; Varia et al., 2023; Lu et al., 2023). To avoid unnecessary penalization, we performed a post-processing step using regular expressions to extract the answers in the desired format for the evaluation. In general, all LLMs fail to solve the task, especially in zero-shot. In the few-shot settings, we can observe some improvements, however the performances are very low. We observe that increasing the number of shots result in different behaviors across the languages. Italian and Indonesian appear to benefit whereas on English it seems to have a detrimental effect. Concerning the order of presentation of the shots the results indicate a positive trend for Context 2, i.e., random order. With few exceptions for English, Mistral-7B is the best performing model across all settings.

The results for the language-adapted versions, Llamantino-3-8B and SahabatAI-8B, offer interesting insights. In zero-shot, both get higher results than Llama-3-8B, the model from which they are derived. This suggests the benefit of language-

⁷<https://huggingface.co/datasets/mlabonne/orpo-dpo-mix-40k>

specific fine-tuning. However, they behave differently in the few-shot settings. In no case, they give better results than Mistral-7B. Llamantino-3-8B also underperforms compared to Llama-3-8B, while this is not the case for SahabatAI-8B. This can be explained in light of the different strategies used in the re-training process (full parameter setting for SahabatAI-8B vs. QLoRA for Llamantino-3-8B). The decrease in bit precision in the Llamantino-3-8B retraining process may have affected the LLM’s ability to learn the examples given in the prompt.

Comparing LLMs with USSA Zhai et al. (2023) propose a bi-lexical dependency parsing graph method, called Unified table filling scheme for SSA (USSA), by converting bi-lexical dependency parsing graph into a unified 2D table-filling scheme. This addresses issues related to both overlap and discontinuity in span prediction. We ran USSA in a multilingual setting using all tweets on our training split (200 for Italian, 300 for English, and 460 for Indonesian). Zhai et al. (2023) use multilingual BERT (mBERT) (Devlin et al., 2019) as the embedding model backbone. For our experiment, we also used multilingual DeBERTa-v3 (mDeBERTa-v3) (He et al., 2021), which has shown to achieve better results when compared to mBERT. Table 5 shows the comparison of USSA against the Mistral-7B with 12 shots (3-set) randomly presented (Context 2), our best LLMs when considering the average SF_1 across all languages as reference.

Model	ENVIS-IT	ENVIS-EN	ENVIS-ID	Avg.
Mistral-7B (3-set Context 2)	11.23	4.17	7.54	7.65
USSA-mBERT	18.57	2.31	10.50	10.46
USSA-mDeBERTa	14.16	4.41	7.33	8.63

Table 5: SF_1 comparison of our best LLM with USSA. Best ones per language and average in bold.

From Table 5, we can see that USSA-mBERT outperforms USSA-mDeBERTa for Italian and Indonesian but not for English. On average, both USSA-mBERT and USSA-mDeBERTa outperform Mistral-7B even when facing a very low number of training instances, further confirming previous work (Su et al., 2024; Sun et al., 2024; Zhang et al., 2024). Considering the differences in the results across languages, the lowest results are on English, regardless of the model used, while Italian and Indonesian have a more similar behavior. From an analysis of the data, it turns out that English

has the highest number of tuples per instance on average (2.67) compared to Italian (1.26) and Indonesian (2.53). In terms of tuple token variability, English has the highest unique sentiment term expressions token (2,849) than Italian (945) and Indonesian (2,093), adding an additional challenge to for all models. Lastly, for topic variability, English contains the highest number of tweets that discuss multiple ES topics in a single tweet (39.20%) compared to Italian (23.10%) and Indonesian (37.39%). Detailed statistics are in Appendix F.

Our results on the ENVIS are notably lower compared to other SSA benchmarks (Barnes et al., 2022; Zhai et al., 2023; Zhou et al., 2024). Although for the LLMs this indicates inherent problems of the models in performing this task properly, the low results for the USSA model are a consequence of the limited amount of training data. While other SSA benchmarks are trained on datasets with thousands of examples, our reimplementation of the USSA model can only rely on hundreds of messages (the full training set amounts to 960 tweets). Additionally, the number of tuples per instance and also sentiment term expressions token and topic variability affects the USSA model preventing the potential benefit of a multilingual training dataset. A further aspect to take into account is the length of the markable sentiment term spans in ENVIS when compared to other SSA datasets. In ENVIS, the average span length of the sentiment term expression ranges between 1.61 and 3.17, while in most of the current SSA datasets ranges between 1.9 and 2.6. This observation is consistent with trends in other SSA datasets, where longer average spans, such as those in the English MPQA dataset are associated with lower model performance.

5 Error Analysis

We have structured the error analysis along two dimensions: quantitative and qualitative. For the quantitative analysis, we follow the error types categories defined by Oberländer and Klinger (2020). The quantitative error analysis takes into account both the best LLM and USSA model. For the qualitative analysis, we have focused only on the LLM outputs and identified seven categories of errors on the basis of the model behavior.

Quantitative Error Analysis Oberländer and Klinger (2020) identify six error categories considering the spans of the markables. In particular,

they list **false positive** (FP), no span in the ground truth, but the model predicts it, **false negative** (FN), one or more span in the ground truth, but the model predicts nothing, **too early** (TE), the predicted span intersects too early with the ground truth, **too late** (TL), the predicted span intersects too late with the ground truth, **multiple** (M), the ground truth span is predicted as two or more separated spans, or vice versa, and **other** (O), the predicted span is the subset of the ground truth span, or vice versa. Overall, FPs are more prominent for Mistral-7B (33.36% vs. 2.92% for USSA-mBERT). FNs are prominent in USSA-mBERT (48.19%) but also largely present in Mistral-7B (31.56%). This further confirms the ineffectiveness of LLMs for this task where a trained model such as USSA-mBERT can fail to detect many cases but it has a very good Precision. In particular, the high FP error rate for Mistral-7B suggests that the model overgenerating markables and it could be considered as a form of hallucination. Overlapping errors (TE and TL), on the other hand, are very low both models representing less than 1% or errors on average (for each type). On the other hand, M are relatively frequent in both models (on average 32.57% for Mistral-7B and 48.50% for USSA-mBERT) consistent with the FN trend, indicating potential issues with distinguishing between multiple valid predictions. The error distribution is consistent with the general trends for each language. Detailed statics are in Appendix G.

Qualitative Error Types We aggregated the errors in three main categories corresponding to **Missing Information** (MI), corresponding to SSA tuples missing either the polarity value or the sentiment expression, **Incorrect Generation** (IG), corresponding to randomly generated tuples, “hallucinated” sequences (i.e., text not present in the original message), and **Miscellaneous**, including cases such as failure to follow instructions, errors in output format, or failure to do the task. Table 6 summarizes their distributions across all languages.

Error Type	ENVIS-IT	ENVIS-EN	ENVIS-ID	Avg.
Missing Information	78.95	52.00	34.09	55.01
Incorrect Generation	10.53	34.00	45.46	30.00
Miscellaneous	10.53	14.00	20.45	14.99

Table 6: Qualitative errors categories for Mistral-7B 3-sets Context 2. Figures correspond to percentages.

MI errors are the most frequent, mostly corresponding to tuples with missing polarity value (29.11% on average) or sentiment terms (25.91% on aver-

age), particularly in Italian and English. This indicates that the model struggles more with sentiment term components of the SSA tuples, thus affecting the FN rates. The IG errors are relatively fewer than MI (30.00% on average) with the most frequent instance being randomly generated tuples (23.15% on average). Although these cases mostly correspond to text spans in the message, they offer further insights on the overgeneration behavior of the model. In the Miscellaneous category, 11.57% of the overall errors are cases where the LLM generates example-based outputs, possibly due to misinterpretation of the task prompt.

These observations show that the major sources of errors are polarity assignment and identification of sentiment terms, highlighting the need for improved polarity classification or post-processing correction. Tuple generation should be better constrained to prevent long random extractions and hallucinations, while refining prompt formulation could reduce the generation of examples errors.

6 Conclusions and Future Works

We present ENVIS, a new multilingual resource 4.8k tweets for SSA on environmental sustainability. This resource can contribute to an understanding of the on-going debate on environmental sustainability opening up multicultural comparison scenarios. We also introduce a new framework to automatically aggregate span-level annotations that, while preserving the annotators’ perspectives, avoids additional manual intervention, reducing costs while maintaining the annotation quality and the grammaticality of the aggregated spans.

We have benchmarked ENVIS against four instruction-tuned encoder-based LLMs using in-context learning. While we have observed a positive influence of a larger number of shots presented in random order, the overall performance is poor, with LLMs being outperformed by a supervised neural model (USSA) trained with less than 1k instances. The results reinforce the finding from Zhang et al. (2024) by extending these observations to encoder-based models, and question the suitability of in-context learning and LLMs for fine-grained information extraction tasks. Experimenting with fine-tuning seems to provide better results (Dagdelen et al., 2024) but calls for a reflection on the suitability of using this highly energy-consuming technology when better results can be achieved with less complex architectures.

Limitations

The dataset used in this study was collected during 2022 and 2023. Since online discourse and attitudes can greatly vary over time, the findings drawn from this dataset may not reflect the previous or future landscape and online behavior toward environmental sustainability.

The dataset focuses specifically on three languages, limiting its generalizability to other languages and cultures. The sentiment about the environment present in Italian, English, and Indonesian Social Media users may not align with those found in different linguistic and cultural contexts.

The paper reports on the use of a range of models for Sentiment Analysis experiments. The performance and results obtained may be influenced by the specific characteristics of these models and their training data. Other models or approaches might yield different results, and the generalizability of the results to other models or architectures should be further investigated.

The limitations or biases arising from the dataset creation process, including data collection and annotation, should be considered in terms of the specific involvement of the annotators and the potential power dynamics that may have influenced the creation of the dataset.

Ethical reflections

The study presented in the paper can raise ethical considerations that should be carefully taken into account when collecting, analyzing, and disseminating the data and results.

This study on the creation and use of a dataset as a benchmark aims to analyze the application of Sentiment Analysis to the ongoing debate on environmental sustainability. In collecting and annotating the dataset, there is a risk of reinforcing or perpetuating existing biases about the issues raised in the collected data. The potential impact of the research on marginalized communities and the broader social implications related to the different perceptions of the observed phenomena should be carefully considered. We did our best to address this aspect by considering data and annotators from the Global North and South.

It is important to consider the possible misuse or unintended consequences of NLP tools. Care should be taken to avoid using systems that unintentionally and disproportionately target particular perspectives or promote misinformation on

environmental issues. We can address this aspect by considering annotations even in disaggregated form, but a thorough analysis of the ethical implications of the tools developed should be conducted. Our work highlights the need to consider and incorporate the subjectivity of annotators in NLP applications and encourages thinking about the different perspectives encoded in annotated datasets to minimize the amplification of biases.

In building the proposed resource, we have taken measures to protect annotators' privacy, and our data processing protocols are designed to protect personal information (e.g., anonymizing users' mentions).

As for the annotation process, we have endeavored to pay annotators fairly, as reported in the paper.

To ensure responsible and ethical use, we intend to implement mechanisms to track the use of the dataset. By recording who accesses and uses the dataset, we aim to promote a better understanding of its impact, encourage collaboration, and potentially address concerns that may arise from its use. The dataset will be made available for research purposes only. To maintain transparency and accountability, we will distribute the dataset under the Creative Commons Attribution 4.0 International (CC BY 4.0) license.

References

- Rodrigo Agerri, Montse Cuadros, Sean Gaines, and German Rigau. 2013. Opener: Open polarity enhanced named entity recognition. *Procesamiento del Lenguaje Natural*, (51):215–218.
- Jeremy Barnes, Toni Badia, and Patrik Lambert. 2018. [MultiBooked: A corpus of Basque and Catalan hotel reviews annotated for aspect-level sentiment classification](#). In *Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)*, Miyazaki, Japan. European Language Resources Association (ELRA).
- Jeremy Barnes, Robin Kurtz, Stephan Oepen, Lilja Øvrelid, and Erik Velldal. 2021. [Structured sentiment analysis as dependency graph parsing](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 3387–3402, Online. Association for Computational Linguistics.
- Jeremy Barnes, Laura Ana Maria Oberländer, Enrica Troiano, Andrey Kutuzov, Jan Buchmann, Rodrigo Agerri, Lilja Øvrelid, and Erik Velldal. 2022. SemEval-2022 task 10: Structured sentiment analysis.

788	In <i>Proceedings of the 16th International Workshop on Semantic Evaluation (SemEval-2022)</i> , Seattle. Association for Computational Linguistics.	845
789		846
790		847
791	Valerio Basile, Michael Fell, Tommaso Fornaciari, Dirk Hovy, Silviu Paun, Barbara Plank, Massimo Poesio, Alexandra Uma, et al. 2021. We need to consider disagreement in evaluation. In <i>Proceedings of the 1st workshop on benchmarking: past, present and future</i> , pages 15–21. Association for Computational Linguistics.	848
792		849
793		850
794		851
795		852
796		853
797		854
798	Cristina Bosco, Muhammad Okky Ibrohim, Valerio Basile, and Indra Budi. 2023. How green is sentiment analysis? environmental topics in corpora at the university of turin. In <i>Proceeding of CLiC-it 2023</i> , Venice, Italy.	855
799		856
800		857
801		858
802		859
803	Cong Chen, Jiansong Chen, Cao Liu, Fan Yang, Guanglu Wan, and Jinxiong Xia. 2022. MT-speech at SemEval-2022 task 10: Incorporating data augmentation and auxiliary task with cross-lingual pretrained language model for structured sentiment analysis . In <i>Proceedings of the 16th International Workshop on Semantic Evaluation (SemEval-2022)</i> , pages 1329–1335, Seattle, United States. Association for Computational Linguistics.	860
804		861
805		862
806		863
807		864
808		865
809		866
810		867
811		868
812	John Dagdelen, Alexander Dunn, Sanghoon Lee, Nicholas Walker, Andrew S Rosen, Gerbrand Ceder, Kristin A Persson, and Anubhav Jain. 2024. Structured information extraction from scientific text with large language models. <i>Nature Communications</i> , 15(1):1418.	869
813		870
814		871
815		872
816		873
817		874
818	Lingjia Deng and Janyce Wiebe. 2015. MPQA 3.0: An entity/event-level sentiment corpus . In <i>Proceedings of the 2015 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 1323–1328, Denver, Colorado. Association for Computational Linguistics.	875
819		876
820		877
821		878
822		879
823		880
824		881
825	Tim Dettmers, Artidoro Pagnoni, Ari Holtzman, and Luke Zettlemoyer. 2023. Qlora: Efficient finetuning of quantized llms. <i>arXiv preprint arXiv:2305.14314</i> .	882
826		883
827		884
828	Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: Pre-training of deep bidirectional transformers for language understanding . In <i>Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)</i> , pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.	885
829		886
830		887
831		888
832		889
833		890
834		891
835		892
836		893
837	Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Al-lonsius, Daniel Song, Danielle Pintz, Danny Livshits, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Graeme Nail, Gregoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Ishan Misra, Ivan Evtimov, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Mahadeokar, Jeet Shah, Jelmer van der Linde, Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden Alwala, Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer, Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Lauren Rantala-Yeary, Laurens van der Maaten, Lawrence Chen, Liang Tan, Liz Jenkins, Louis Martin, Lovish Madaan, Lubo Malo, Lukas Blecher, Lukas Landzaat, Luke de Oliveira, Madeline Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar Paluri, Marcin Kardas, Mathew Oldham, Mathieu Rita, Maya Pavlova, Melanie Kambadur, Mike Lewis, Min Si, Mitesh Kumar Singh, Mona Hassan, Naman Goyal, Narjes Torabi, Nikolay Bashlykov, Nikolay Bogoychev, Niladri Chatterji, Olivier Duchenne, Onur Celebi, Patrick Alrassy, Pengchuan Zhang, Pengwei Li, Petar Vasic, Peter Weng, Prajjwal Bhargava, Pratik Dubal, Praveen Krishnan, Punit Singh Koura, Puxin Xu, Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj Ganapathy, Ramon Calderer, Ricardo Silveira Cabral, Robert Stojnic, Roberta Raileanu, Rohit Girdhar, Rohit Patel, Roman Sauvestre, Ronnie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sahana Chennabasappa, Sanjay Singh, Sean Bell, Seohyun Sonia Kim, Sergey Edunov, Shao-liang Nie, Sharan Narang, Sharath Raparthy, Sheng Shen, Shengye Wan, Shruti Bhosale, Shun Zhang, Simon Vandenhende, Soumya Batra, Spencer Whitman, Sten Sootla, Stephane Collet, Suchin Gururangan, Sydney Borodinsky, Tamar Herman, Tara Fowler, Tarek Sheasha, Thomas Georgiou, Thomas Scialom, Tobias Speckbacher, Todor Mihaylov, Tong Xiao, Ujjwal Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh Ramanathan, Viktor Kerkez, Vincent Gouget, Virginie Do, Vish Vogeti, Vladan Petrovic, Weiwei Chu, Wenhan Xiong, Wenxin Fu, Whitney Meers, Xavier Martinet, Xiaodong Wang, Xiaoqing Ellen Tan, Xinfeng Xie, Xuchao Jia, Xuwei Wang, Yaelle Goldschlag, Yashesh Gaur, Yasmine Babaei, Yi Wen, Yiwen Song, Yuchen Zhang, Yue Li, Yuning Mao, Zacharie Delpierre Coudert, Zheng Yan, Zhengxing Chen, Zoe Papakipos, Aaditya Singh, Aaron Grattafiori, Abha Jain, Adam Kelsey, Adam Shajnfeld, Adithya Gangidi, Adolfo Victoria, Ahuva	900
838		901
839		902
840		903
841		904
842		905
843		906
844		907
		908

909	Goldstand, Ajay Menon, Ajay Sharma, Alex Boesenberg, Alex Vaughan, Alexei Baeviski, Allie Feinstein, Amanda Kallet, Amit Sangani, Anam Yunus, Andrei Lupu, Andres Alvarado, Andrew Caples, Andrew Gu, Andrew Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchandani, Annie Franco, Aparajita Saraf, Arkabandhu Chowdhury, Ashley Gabriel, Ashwin Bharambe, Assaf Eisenman, Azadeh Yazdan, Beau James, Ben Maurer, Benjamin Leonhardi, Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi Paranjape, Bing Liu, Bo Wu, Boyu Ni, Braden Hancock, Bram Wasti, Brandon Spence, Brani Stojkovic, Brian Gamido, Britt Montalvo, Carl Parker, Carly Burton, Catalina Mejia, Changan Wang, Changkyu Kim, Chao Zhou, Chester Hu, Ching-Hsiang Chu, Chris Cai, Chris Tindal, Christoph Feichtenhofer, Damon Civin, Dana Beaty, Daniel Kreymer, Daniel Li, Danny Wyatt, David Adkins, David Xu, Davide Testuggine, Delia David, Devi Parikh, Diana Liskovich, Didem Foss, Dingkan Wang, Duc Le, Dustin Holland, Edward Dowling, Eissa Jamil, Elaine Montgomery, Eleonora Presani, Emily Hahn, Emily Wood, Erik Brinkman, Esteban Arcaute, Evan Dunbar, Evan Smothers, Fei Sun, Felix Kreuk, Feng Tian, Firat Ozgenel, Francesco Caggioni, Francisco Guzmán, Frank Kanayet, Frank Seide, Gabriela Medina Florez, Gabriella Schwarz, Gada Badeer, Georgia Sweet, Gil Halpern, Govind Thattai, Grant Herman, Grigory Sizov, Guangyi, Zhang, Guna Lakshminarayanan, Hamid Shojanazeri, Han Zou, Hannah Wang, Hanwen Zha, Haroun Habeeb, Harrison Rudolph, Helen Suk, Henry Aspegren, Hunter Goldman, Ibrahim Damla, Igor Molybog, Igor Tufanov, Irina-Elena Veliche, Itai Gat, Jake Weissman, James Geboski, James Kohli, Japhet Asher, Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jennifer Chan, Jenny Zhen, Jeremy Reizenstein, Jeremy Teboul, Jessica Zhong, Jian Jin, Jingyi Yang, Joe Cummings, Jon Carvill, Jon Shepard, Jonathan McPhie, Jonathan Torres, Josh Ginsburg, Junjie Wang, Kai Wu, Kam Hou U, Karan Saxena, Karthik Prasad, Kartikay Khandelwal, Katayoun Zand, Kathy Matosich, Kaushik Veeraraghavan, Kelly Michelena, Keqian Li, Kun Huang, Kunal Chawla, Kushal Lakhotia, Kyle Huang, Lailin Chen, Lakshya Garg, Lavender A, Leandro Silva, Lee Bell, Lei Zhang, Liangpeng Guo, Licheng Yu, Liron Moshkovich, Luca Wehrstedt, Madian Khabza, Manav Avalani, Manish Bhatt, Maria Tsim-poukelli, Martynas Mankus, Matan Hasson, Matthew Lennie, Matthias Reso, Maxim Groshev, Maxim Naumov, Maya Lathi, Meghan Keneally, Michael L. Seltzer, Michal Valko, Michelle Restrepo, Mihir Patel, Mik Vyatskov, Mikayel Samvelyan, Mike Clark, Mike Macey, Mike Wang, Miquel Jubert Hermoso, Mo Metanat, Mohammad Rastegari, Munish Bansal, Nandhini Santhanam, Natascha Parks, Natasha White, Navyata Bawa, Nayan Singhal, Nick Egebo, Nicolas Usunier, Nikolay Pavlovich Laptev, Ning Dong, Ning Zhang, Norman Cheng, Oleg Chernoguz, Olivia Hart, Omkar Salpekar, Ozlem Kalinli, Parkin Kent, Parth Parekh, Paul Saab, Pavan Balaji, Pedro Rittner, Philip Bontrager, Pierre Roux, Piotr Dollar, Polina Zvyagina, Prashant Ratan-chandani, Pritish Yuvraj, Qian Liang, Rachad Alao,	
	Rachel Rodriguez, Rafi Ayub, Raghotham Murthy, Raghu Nayani, Rahul Mitra, Raymond Li, Rebekkah Hogan, Robin Battey, Rocky Wang, Rohan Maheswari, Russ Howes, Ruty Rinott, Sai Jayesh Bondu, Samyak Datta, Sara Chugh, Sara Hunt, Sargun Dhillon, Sasha Sidorov, Satadru Pan, Saurabh Verma, Seiji Yamamoto, Sharadh Ramaswamy, Shaun Lindsay, Shaun Lindsay, Sheng Feng, Shenghao Lin, Shengxin Cindy Zha, Shiva Shankar, Shuqiang Zhang, Shuqiang Zhang, Sinong Wang, Sneha Agarwal, Soji Sajuyigbe, Soumith Chintala, Stephanie Max, Stephen Chen, Steve Kehoe, Steve Satterfield, Sudarshan Govindaprasad, Sumit Gupta, Sungmin Cho, Sunny Virk, Suraj Subramanian, Sy Choudhury, Sydney Goldman, Tal Remez, Tamar Glaser, Tamara Best, Thilo Kohler, Thomas Robinson, Tianhe Li, Tianjun Zhang, Tim Matthews, Timothy Chou, Tzook Shaked, Varun Vontimitta, Victoria Ajayi, Victoria Montanez, Vijai Mohan, Vinay Satish Kumar, Vishal Mangla, Vitor Albiero, Vlad Ionescu, Vlad Poenaru, Vlad Tiberiu Mihailescu, Vladimir Ivanov, Wei Li, Wenchen Wang, Wenwen Jiang, Wes Bouaziz, Will Constable, Xiaocheng Tang, Xiaofang Wang, Xiaojian Wu, Xiaolan Wang, Xide Xia, Xilun Wu, Xinbo Gao, Yanjun Chen, Ye Hu, Ye Jia, Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi, Youngjin Nam, Yu, Wang, Yuchen Hao, Yundi Qian, Yuzi He, Zach Rait, Zachary DeVito, Zef Rosnbrick, Zhaoduo Wen, Zhenyu Yang, and Zhiwei Zhao. 2024. The llama 3 herd of models .	973 974 975 976 977 978 979 980 981 982 983 984 985 986 987 988 989 990 991 992 993 994 995 996 997 998 999 1000 1001 1002
	Ridong Han, Tao Peng, Chaohao Yang, Benyou Wang, Lu Liu, and Xiang Wan. 2023. Is information extraction solved by chatgpt? an analysis of performance, evaluation criteria, robustness and errors .	1003 1004 1005 1006
	Pengcheng He, Jianfeng Gao, and Weizhu Chen. 2021. Debertav3: Improving deberta using electra-style pre-training with gradient-disentangled embedding sharing .	1007 1008 1009 1010
	Muhammad Okky Ibrohim, Cristina Bosco, and Valerio Basile. 2023. Sentiment analysis for the natural environment: A systematic review . <i>ACM Comput. Surv.</i> , 56(4).	1011 1012 1013 1014
	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 .	1015 1016 1017 1018 1019 1020 1021
	Andrei P Kirilenko and Svetlana O Stepchenkova. 2014. Public microblogging on climate change: One year of twitter worldwide. <i>Global environmental change</i> , 26:171–182.	1022 1023 1024 1025
	J. Richard Landis and Gary G. Koch. 1977. The measurement of observer agreement for categorical data . <i>Biometrics</i> , 33(1):159–174.	1026 1027 1028
	Yangkun Lin, Chen Liang, Jing Xu, Chong Yang, and Yongliang Wang. 2022. ZHIXIAOBAO at	1029 1030

1031	SemEval-2022 task 10: Approaching structured sentiment with graph parsing. In <i>Proceedings of the 16th International Workshop on Semantic Evaluation (SemEval-2022)</i> , pages 1343–1348, Seattle, United States. Association for Computational Linguistics.	1087
1032		1088
1033		1089
1034		1090
1035		1091
1036	Bing Liu. 2015. <i>Sentiment analysis: Mining opinions, sentiments, and emotions</i> . Cambridge University Press.	1092
1037		1093
1038		1094
1039	Sheng Lu, Irina Bigoulaeva, Rachneet Sachdeva, Harish Tayyar Madabushi, and Iryna Gurevych. 2023. Are emergent abilities in large language models just in-context learning? <i>arXiv preprint arXiv:2309.01809</i> .	1095
1040		1096
1041		1097
1042		1098
1043		
1044	Sahal Mullappilly, Abdelrahman Shaker, Omkar Thawakar, Hisham Cholakkal, Rao Anwer, Salman Khan, and Fahad Khan. 2023. Arabic miniclimategpt: A climate change and sustainability tailored arabic llm. In <i>Findings of the Association for Computational Linguistics: EMNLP 2023</i> , pages 14126–14136.	1099
1045		1100
1046		1101
1047		1102
1048		1103
1049		1104
1050		1105
1051	Jingwei Ni, Julia Bingler, Chiara Colesanti-Senni, Mathias Kraus, Glen Gostlow, Tobias Schimanski, Dominik Stambach, Saeid Ashraf Vaghefi, Qian Wang, Nicolas Webersinke, Tobias Wekhof, Tingyu Yu, and Markus Leippold. 2023. <i>CHATREPORT: Democratizing sustainability disclosure analysis through LLM-based tools</i> . In <i>Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing: System Demonstrations</i> , pages 21–51, Singapore. Association for Computational Linguistics.	1106
1052		1107
1053		1108
1054		1109
1055		1110
1056		1111
1057		1112
1058		
1059		1113
1060		1114
1061		1115
1062	Laura Ana Maria Oberländer and Roman Klinger. 2020. <i>Token sequence labeling vs. clause classification for English emotion stimulus detection</i> . In <i>Proceedings of the Ninth Joint Conference on Lexical and Computational Semantics</i> , pages 58–70, Barcelona, Spain (Online). Association for Computational Linguistics.	1116
1063		1117
1064		1118
1065		1119
1066		1120
1067		1121
1068	Lilja Øvrelid, Petter Mæhlum, Jeremy Barnes, and Erik Velldal. 2020. <i>A fine-grained sentiment dataset for Norwegian</i> . In <i>Proceedings of the Twelfth Language Resources and Evaluation Conference</i> , pages 5025–5033, Marseille, France. European Language Resources Association.	1122
1069		1123
1070		1124
1071		1125
1072		
1073		1126
1074	Haiyun Peng, Lu Xu, Lidong Bing, Fei Huang, Wei Lu, and Luo Si. 2020. <i>Knowing what, how and why: A near complete solution for aspect-based sentiment analysis</i> . <i>Proceedings of the AAAI Conference on Artificial Intelligence</i> , 34(05):8600–8607.	1127
1075		1128
1076		1129
1077		1130
1078		1131
1079	Barbara Plank. 2022. The "problem" of human label variation: On ground truth in data, modeling and evaluation. In <i>Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing</i> , pages 10671–10682.	1132
1080		1133
1081		1134
1082		1135
1083		1136
1084	Marco Polignano, Pierpaolo Basile, and Giovanni Semeraro. 2024. <i>Advanced natural-based interaction for the italian language: Llamantino-3-anita</i> .	1137
1085		1138
1086		1139
		1140
		1141
		1142
		1143
	Maria Pontiki, Dimitris Galanis, Haris Papageorgiou, Ion Androutsopoulos, Suresh Manandhar, Mohammad AL-Smadi, Mahmoud Al-Ayyoub, Yanyan Zhao, Bing Qin, Orphée De Clercq, Véronique Hoste, Marianna Apidianaki, Xavier Tannier, Natalia Loukachevitch, Evgeniy Kotelnikov, Nuria Bel, Salud María Jiménez-Zafra, and Gülşen Eryiğit. 2016. <i>SemEval-2016 task 5: Aspect based sentiment analysis</i> . In <i>Proceedings of the 10th International Workshop on Semantic Evaluation (SemEval-2016)</i> , pages 19–30, San Diego, California. Association for Computational Linguistics.	
	Maria Pontiki, Dimitris Galanis, Haris Papageorgiou, Suresh Manandhar, and Ion Androutsopoulos. 2015. <i>SemEval-2015 task 12: Aspect based sentiment analysis</i> . In <i>Proceedings of the 9th International Workshop on Semantic Evaluation (SemEval 2015)</i> , pages 486–495, Denver, Colorado. Association for Computational Linguistics.	
	Maria Pontiki, Dimitris Galanis, John Pavlopoulos, Haris Papageorgiou, Ion Androutsopoulos, and Suresh Manandhar. 2014. <i>SemEval-2014 task 4: Aspect based sentiment analysis</i> . In <i>Proceedings of the 8th International Workshop on Semantic Evaluation (SemEval 2014)</i> , pages 27–35, Dublin, Ireland. Association for Computational Linguistics.	
	Filipe Rodrigues, Francisco Pereira, and Bernardete Ribeiro. 2014. <i>Sequence labeling with multiple annotators</i> . <i>Machine Learning</i> , 95(2):165–181.	
	Mingyang Song, Mao Zheng, and Xuan Luo. 2025. <i>Can many-shot in-context learning help LLMs as evaluators? a preliminary empirical study</i> . In <i>Proceedings of the 31st International Conference on Computational Linguistics</i> , pages 8232–8241, Abu Dhabi, UAE. Association for Computational Linguistics.	
	Daniel Spokoyny, Tanmay Laud, Tom Corringham, and Taylor Berg-Kirkpatrick. 2023. <i>Towards answering climate questionnaires from unstructured climate reports</i> .	
	Dominik Stambach, Nicolas Webersinke, Julia Bingler, Mathias Kraus, and Markus Leippold. 2023. <i>Environmental claim detection</i> . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)</i> , pages 1051–1066, Toronto, Canada. Association for Computational Linguistics.	
	Manfred Stede and Ronny Patz. 2021. <i>The climate change debate and natural language processing</i> . In <i>Proceedings of the 1st Workshop on NLP for Positive Impact</i> , pages 8–18, Online. Association for Computational Linguistics.	
	Guixin Su, Mingmin Wu, Zhongqiang Huang, Yongcheng Zhang, Tongguan Wang, Yuxue Hu, and Ying Sha. 2024. <i>Refine, align, and aggregate: Multi-view linguistic features enhancement for aspect sentiment triplet extraction</i> . In <i>Findings of the Association for Computational Linguistics: ACL 2024</i> ,	

1144	pages 3212–3228, Bangkok, Thailand. Association	Wenxuan Zhang, Yue Deng, Bing Liu, Sinno Pan, and	1199
1145	for Computational Linguistics.	Lidong Bing. 2024. Sentiment analysis in the era	1200
		of large language models: A reality check . In <i>Find-</i>	1201
1146	Qiao Sun, Liujia Yang, Minghao Ma, Nanyang Ye, and	<i>ings of the Association for Computational Linguis-</i>	1202
1147	Qinying Gu. 2024. MiniConGTS: A near ultimate	<i>tics: NAACL 2024</i> , pages 3881–3906, Mexico City,	1203
1148	minimalist contrastive grid tagging scheme for as-	Mexico. Association for Computational Linguistics.	1204
1149	pect sentiment triplet extraction . In <i>Proceedings of</i>		
1150	<i>the 2024 Conference on Empirical Methods in Natu-</i>	Chengjie Zhou, Bobo Li, Hao Fei, Fei Li, Chong Teng,	1205
1151	<i>ral Language Processing</i> , pages 2817–2834, Miami,	and Donghong Ji. 2024. Revisiting structured sen-	1206
1152	Florida, USA. Association for Computational Lin-	timent analysis as latent dependency graph parsing .	1207
1153	guistics.	In <i>Proceedings of the 62nd Annual Meeting of the</i>	1208
		<i>Association for Computational Linguistics (Volume 1:</i>	1209
1154	Cigdem Toprak, Niklas Jakob, and Iryna Gurevych.	<i>Long Papers)</i> , pages 10178–10191, Bangkok, Thai-	1210
1155	2010. Darmstadt service review corpus .	land. Association for Computational Linguistics.	1211
1156	Siddharth Varia, Shuai Wang, Kishaloy Halder, Robert		
1157	Vacareanu, Miguel Ballesteros, Yassine Benajiba,		
1158	Neha Anna John, Rishita Anubhai, Smaranda Mure-		
1159	san, and Dan Roth. 2023. Instruction tuning for few-		
1160	shot aspect-based sentiment analysis. In <i>Proceedings</i>		
1161	<i>of the 13th Workshop on Computational Approaches</i>		
1162	<i>to Subjectivity, Sentiment and Social Media Analysis</i> .		
1163	Association for Computational Linguistics.		
1164	Giuseppe A Veltri and Dimitrinka Atanasova. 2017. Cli-		
1165	mate change on twitter: Content, media ecology and		
1166	information sharing behaviour. <i>Public understanding</i>		
1167	<i>of science</i> , 26(6):721–737.		
1168	Janyce Wiebe, Theresa Wilson, and Claire Cardie. 2005.		
1169	Annotating expressions of opinions and emotions		
1170	in language. <i>Language resources and evaluation</i> ,		
1171	39:165–210.		
1172	Bryan Wilie, Karissa Vincentio, Genta Indra Winata,		
1173	Samuel Cahyawijaya, Xiaohong Li, Zhi Yuan Lim,		
1174	Sidik Soleman, Rahmad Mahendra, Pascale Fung,		
1175	Syafri Bahar, and Ayu Purwarianti. 2020. IndoNLU:		
1176	Benchmark and resources for evaluating Indonesian		
1177	natural language understanding . In <i>Proceedings of</i>		
1178	<i>the 1st Conference of the Asia-Pacific Chapter of the</i>		
1179	<i>Association for Computational Linguistics and the</i>		
1180	<i>10th International Joint Conference on Natural Lan-</i>		
1181	<i>guage Processing</i> , pages 843–857, Suzhou, China.		
1182	Association for Computational Linguistics.		
1183	Lu Xu, Hao Li, Wei Lu, and Lidong Bing. 2020.		
1184	Position-aware tagging for aspect sentiment triplet		
1185	extraction . In <i>Proceedings of the 2020 Conference on</i>		
1186	<i>Empirical Methods in Natural Language Processing</i>		
1187	<i>(EMNLP)</i> , pages 2339–2349, Online. Association for		
1188	Computational Linguistics.		
1189	Zepeng Zhai, Hao Chen, Ruifan Li, and Xiaojie Wang.		
1190	2023. USSA: A unified table filling scheme for struc-		
1191	tured sentiment analysis . In <i>Proceedings of the 61st</i>		
1192	<i>Annual Meeting of the Association for Computational</i>		
1193	<i>Linguistics (Volume 1: Long Papers)</i> , pages 14340–		
1194	14353, Toronto, Canada. Association for Computa-		
1195	tional Linguistics.		
1196	Lei Zhang, Shuai Wang, and Bing Liu. 2018. Deep		
1197	learning for sentiment analysis: A survey . <i>WIREs</i>		
1198	<i>Data Mining and Knowledge Discovery</i> , 8(4):e1253.		

A **Keywords Used to Collect the Dataset**

<i>transizione energetica</i> (energy turnaround)	agenda 2030	<i>crisi climatica</i> (climate crisis)	<i>combustibili fossili</i> (fossil fuel)	<i>deforestazione</i> (deforestation)	greenwashing
<i>riscaldamento globale</i> (global warming)	<i>impatto ambientale</i> (environmental impact)	climate change	green deal	<i>sviluppo sostenibile</i> (sustainability)	COP26
<i>energie rinnovabili</i> (renewable energy)					

Table A: Keywords used by [Bosco et al. \(2023\)](#) to collect Italian Twitter. Some English keywords were directly used to scrap the data.

carbon dioxide	carbon footprint	carbon leakage	carbon taxation	CH4	CO2
decarbonization	GHG	green house	methane	carbon credit	carbon price
act on climate	climate effect	global warming	alternative energy	clean energy	energy future
energy generation	energy production	energy transition	energy saving	fossil fuel	algal energy
green energy	power plant	nuclear matters	nuclear power	renewable energy	solar panel
sustainable energy	wind energy	wind farm	wind power	wind turbine	air quality
environmental conflict	deforestation	environmentalist	environment footprint	environment friendly	environment protection
environment regulation	environment saving	natural environment	world environment day	livable places	abandoned area
abandoned land	blighted area	blighted land	brownfield	contaminated land	empty land
greyfield	polluted land	undeveloped land	unsustainable land	unused land	urban vacancy
urban vacant lots	vacant area	vacant land	vacant parcel	urban park	urban planning
water crisis	water scarcity	water issue	water quality	alternative meat	food contamination
food poisoning	food quality	food safety	gluten	GMO food	GMO fruit
man mad meat	organic agriculture	organic farming	organic food	beyond burger	beyond meat
plant based	vegan	plant meat	green consumerism	green governance	off shore oil production
off shore platform	oil and gas decommissioning	green hotel	green park	green tourism	green area
green spaces	genetically modified organism	GMO	net zero	oil spill	pollution
sustainable agriculture	SDGs	sustainability	sustainable development goals	sustainable energy consumption	sustainable food consumption
sustainable hotel	sustainable tourism	sustainable transport	urban mobility	urban system	sanitation waste
sewage waste	waste collection	waste crisis	waste issue	waste management	reduce reuse recycle

Table B: Keywords used by [Bosco et al. \(2023\)](#) to collect English Twitter.

<i>lingkungan alam</i> (natural environment)	<i>hari lingkungan hidup sedunia</i> (world environment day)	<i>jejak lingkungan</i> (environmental footprint)	<i>ramah lingkungan</i> (environment friendly)	<i>perlindungan lingkungan</i> (environment protection)	<i>peraturan lingkungan</i> (environment regulation)
<i>regulasi lingkungan</i> (environment regulation)	<i>penghematan lingkungan</i> (environmental saving)	<i>penyelamatan lingkungan</i> (environmental saving)	<i>pecinta lingkungan</i> (environmentalist)	<i>penggundulan hutan</i> (deforestation)	<i>konflik lingkungan</i> (environmental conflict)
<i>kualitas udara</i> (air quality)	<i>masalah air</i> (water issue)	<i>kualitas air</i> (water quality)	<i>krisis air</i> (water crisis)	<i>kelangkaan air</i> (water scarcity)	<i>perencanaan kota</i> (urban planning)
<i>konstruksi perkotaan</i> (urban construction)	<i>tanah kosong</i> (vacant land)	<i>daerah kosong</i> (vacant area)	<i>bidang kosong</i> (vacant parcel)	<i>kekosongan perkotaan</i> (urban vacancy)	<i>lahan kosong perkotaan</i> (urban vacant lots)
<i>tanah rusak</i> (blighted land)	<i>daerah rusak</i> (blighted area)	<i>tanah terlantar</i> (abandoned area)	<i>lahan bekas industri</i> (brownfield)	<i>lahan industri</i> (greyfield)	<i>tanah tercemar</i> (polluted land)
<i>pencemaran tanah</i> (contaminated land)	<i>tanah terkontaminasi</i> (contaminated land)	<i>tanah tidak terpakai</i> (unused land)	<i>tanah belum dikembangkan</i> (undeveloped land)	<i>lahan kosong</i> (empty land)	<i>lahan tidak berkelanjutan</i> (unsustainable land)
<i>tempat layak huni</i> (livable place)	<i>taman kota</i> (urban park)	<i>taman hijau</i> (green park)	<i>lahan hijau</i> (green area)	<i>ruang hijau</i> (green space)	<i>wisata hijau</i> (green tourism)
<i>wisata ramah lingkungan</i> (green tourism)	<i>hotel hijau</i> (green hotel)	<i>hotel ramah lingkungan</i> (green hotel)	<i>konsumerisme ramah lingkungan</i> (green consumerism)	<i>pemerintahan ramah lingkungan</i> (green governance)	<i>anjuan lepas pantai</i> (off shore platform)
<i>anjuan minyak lepas pantai</i> (off shore oil platform)	<i>anjuan minyak dan gas</i> (oil and gas platform)	<i>produksi minyak lepas pantai</i> (off shore oil production)	<i>keberlanjutan</i> (sustainability)	<i>tujuan pembangunan berkelanjutan</i> (sustainable development goals)	SDGs
<i>sistem perkotaan</i> (urban system)	<i>mobilitas perkotaan</i> (urban mobility)	<i>transportasi berkelanjutan</i> (sustainable transport)	<i>pariwisata berkelanjutan</i> (sustainable tourism)	<i>perhotelan berkelanjutan</i> (sustainable hotel)	<i>pertanian berkelanjutan</i> (sustainable agriculture)
<i>konsumsi pangan berkelanjutan</i> (sustainable food consumption)	<i>konsumsi energi berkelanjutan</i> (sustainable energy consumption)	<i>kualitas makanan</i> (food quality)	<i>keamanan pangan</i> (food safety)	<i>pencemaran makanan</i> (food contamination)	<i>kontaminasi makanan</i> (food contamination)
<i>keracunan makanan</i> (food poisoning)	<i>makanan organik</i> (organic food)	<i>pertanian organik</i> (organic agriculture)	<i>perkebunan organik</i> (organic farming)	<i>bebas gula</i> (gluten free)	<i>daging alternatif</i> (alternative meat)
<i>daging buatan manusia</i> (man-made meat)	<i>daging dari tumbuhan</i> (plant meat)	<i>berbasis tanaman</i> (plant-based)	<i>daging nabati</i> (beyond meat)	<i>burger nabati</i> (beyond burger)	vegan
<i>veganisme</i> (veganism)	<i>makanan GMO</i> (GMO food)	<i>buah GMO</i> (GMO fruit)	<i>produk rekayasa genetika</i> (genetically modified organism)	GMO	<i>perubahan iklim</i> (climate change)
<i>aksi iklim</i> (act on climate)	<i>darurat iklim</i> (climate emergency)	<i>krisis iklim</i> (climate crisis)	<i>pemanasan global</i> (global warming)	<i>pengaruh iklim</i> (climate effect)	<i>jejak karbon</i> (carbon footprint)
<i>kebocoran karbon</i> (carbon leakage)	<i>dekarbonisasi</i> (decarbonisation)	<i>karbon dioksida</i> (carbon dioxide)	CO2	GHG	CH4
<i>metana</i> (methane)	<i>rumah kaca</i> (green house)	<i>pajak karbon</i> (carbon tax)	<i>perpajakan karbon</i> (carbon taxation)	<i>kredit karbon</i> (carbon credit)	<i>harga karbon</i> (carbon price)
<i>produksi energi</i> (energy production)	<i>transisi energi</i> (energy transation)	<i>masa depan energi</i> (energy future)	<i>pembangkit listrik</i> (energy generation)	<i>energi alternatif</i> (alternative energy)	<i>energi bersih</i> (clean energy)
<i>bahan bakar fosil</i> (fossil fuel)	<i>industri perminyakan</i> (oil industry)	<i>industri batu bara</i> (coal industry)	<i>pembangkit listrik tenaga batu bara</i> (coal plant)	<i>PLTU batu bara</i> (coal plant)	<i>pembangkit listrik tenaga gas</i> (gas plant)
<i>PLTG</i> (gas plant)	<i>gas alam</i> (natural gas)	<i>energi angin</i> (wind energy)	<i>tenaga angin</i> (wind power)	<i>ladang angin</i> (wind farm)	<i>turbin angin</i> (wind turbine)
<i>energi nuklir</i> (nuclear energy)	<i>tenaga nuklir</i> (nuclear power)	<i>permasalahan nuklir</i> (nuclear matters)	<i>energi terbarukan</i> (renewable energy)	<i>aksi energi terbarukan</i> (renewable energy act)	<i>panel surya</i> (solar panel)
<i>energi surya</i> (solar energy)	<i>tenaga surya</i> (solar power)	<i>kebijakan feed-in tariff</i> (feed-in tariff)	<i>kebijakan feed-in remuneration</i> (feed-in remuneration)	<i>energi panas bumi</i> (geothermal energy)	<i>energi termal</i> (thermal energy)
<i>bahan bakar nabati</i> (biofuel)	<i>energi hijau</i> (green energy)	<i>energi ramah lingkungan</i> (green energy)	<i>pembangkit listrik</i> (power plant)	<i>energi alga</i> (alga energy)	<i>energi berkelanjutan</i> (sustainable energy)
<i>penghematan energi</i> (energi saving)	<i>persoalan sampah</i> (waste issue)	<i>permasalahan sampah</i> (waste issue)	<i>krisis limbah</i> (waste crisis)	<i>cangkir menstruasi</i> (menstrual cup)	<i>limbah plastik</i> (plastic waste)
<i>polusi plastik</i> (plastic pollution)	<i>sampah makanan</i> (food waste)	<i>air limbah</i> (sewage waste)	<i>limbah sanitasi</i> (sanitation waste)	<i>pengumpulan sampah</i> (waste collection)	<i>mengurangi, menggunakan kembali, daur ulang</i> (reduce, reuse, recycle)
<i>manajemen limbah</i> (waste management)	<i>manajemen sampah</i> (trash management)	<i>larangan plastik</i> (plastic ban)	<i>larangan polietilena</i> (polythene ban)	<i>polusi</i> (pollution)	<i>nol emisi karbon</i> (net zero)
<i>tumpahan minyak</i> (oil spill)	<i>polusi udara</i> (air pollution)	<i>emisi</i> (emission)			

Table C: Keywords used to collect Indonesian Twitter by translating and expanding English keywords used by Bosco et al. (2023)

Statistic	ENVIS-IT			ENVIS-EN			ENVIS-EN		
	A1 vs. A2	A1 vs. A3	A2 vs. A3	A1 vs. A2	A1 vs. A3	A2 vs. A3	A1 vs. A2	A1 vs. A3	A2 vs. A3
holder	0.28	0.30	0.35	0.57	0.50	0.47	0.26	0.30	0.46
target	0.58	0.56	0.58	0.62	0.56	0.56	0.51	0.40	0.50
negative term	0.46	0.46	0.43	0.47	0.44	0.45	0.64	0.52	0.67
positive term	0.45	0.41	0.40	0.45	0.40	0.40	0.57	0.44	0.64

Table D: Pairwise Span Cohen’s κ details.

Algorithm 1 PCR

```

1:  $pcr\_list = []$ 
2: for  $doc$  in  $aggregated\_dataset$  do
3:    $total\_span =$  Count number of aggregated span in  $doc$ 
4:   if  $total\_span$  is 0 then
5:     if No agreement in document level then
6:        $pcr\_doc = 1$ 
7:     else
8:        $pcr\_doc = 0$ 
9:     end if
10:  else
11:     $correct\_span = 0$ 
12:     $generated\_phrase =$  Generate all phrases from  $doc$  based on dependency tree.
13:    for each aggregated  $span$  in  $doc$  do
14:      if  $span$  in  $generated\_phrase$  then
15:         $correct\_span++ = 1$ 
16:      end if
17:    end for
18:     $pcr\_doc = correct\_span/total\_span$ 
19:  end if
20:  Append  $pcr\_doc$  to  $pcr\_list$ 
21: end for
22: return  $mean(pcr\_list)$ 

```

Statistic	Disagreement (%)			
	ENVIS-IT	ENVIS-EN	ENVIS-ID	Avg.
holder	7.04	12.87	6.67	8.86
target	22.56	23.19	16.92	20.89
sentiment term	25.77	32.05	19.00	25.61

Table E: Percentage of tweets without holder, target, and sentiment term expression caused by disagreement.

Compito: Determinare tutte le espressioni di sentimento e la loro polarità del sentimento nell'input fornito. Per ogni espressione di sentimento estratta, se presente, determinare anche il detentore e il destinatario dell'espressione di sentimento. Si prega di non spiegare la risposta.

Istruzioni:

- La polarità del sentimento può essere solo "negativa" o "positiva".
- Il formato di output deve essere ["titolare", "destinazione", "frase del sentimento", "polarità del sentimento"].
- Se non è presente alcun titolare e/o target da una particolare espressione di sentimento estratta, compilare con una stringa vuota "".
- Se non è presente alcuna frase esplicativa, allora rispondi solo [].
- Non fornire alcuna spiegazione della tua risposta.

Ingresso:

- Testo: "{text}"

Risposta:

Figure A: Zero-shot prompt for ENVIS-IT.

Compito: Determinare tutte le espressioni di sentimento e la loro polarità del sentimento nell'input fornito. Per ogni espressione di sentimento estratta, se presente, determinare anche il detentore e il destinatario dell'espressione di sentimento. Vi verranno forniti anche alcuni esempi. Si prega di non spiegare la risposta.

Istruzioni:

- La polarità del sentimento può essere solo "negativa" o "positiva".
- Il formato di output deve essere ["titolare", "destinazione", "frase del sentimento", "polarità del sentimento"].
- Se non è presente alcun titolare e/o target da una particolare espressione di sentimento estratta, compilare con una stringa vuota "".
- Se non è presente alcuna frase esplicativa, allora rispondi solo [].
- Gli esempi possono aiutarti a determinare la risposta.
- Non fornire alcuna spiegazione della tua risposta.

Esempi:

1. "{text_example_1}". Risposta: "{answer_example_1}"
- ...
- n. "{text_example_n}". Risposta: "{answer_example_n}"

Il tuo turno

Ingresso:

- Testo: "{text}"

Risposta:

Figure B: Few-shot prompt for ENVIS-IT.

Task: Determine all sentiment phrases and their sentiment polarity in the given input. For each extracted sentiment phrase, if any, determine also the holder and target of the sentiment phrase. Please do not explain your answer.

Instructions:

- The sentiment polarity can only be "negative" or "positive".
- The output format should be ["holder", "target", "sentiment phrase", "sentiment polarity"].
- If there is no holder and/or target from a particular extracted sentiment phrase, please fill with an empty string "".
- If there is no sentiment phrase at all, then only answer [].
- Do not give any explanation of your answer.

Input:

- Text: "{text}"

Answer:

Figure C: Zero-shot prompt for ENVIS-EN.

Task: Determine all sentiment phrases and their sentiment polarity in the given input. For each extracted sentiment phrase, if any, determine also the holder and target of the sentiment phrase. You are also provided with some examples. Please do not explain your answer.

Instructions:

- The sentiment polarity can only be "negative" or "positive".
- The output format should be ["holder", "target", "sentiment phrase", "sentiment polarity"].
- If there is no holder and/or target from a particular extracted sentiment phrase, please fill with an empty string "".
- If there is no sentiment phrase at all, then only answer [].
- The examples may assist you in determining the answer.
- Do not give any explanation of your answer.

Examples:

1. "{text_example_1}". Answer: "{answer_example_1}"
- ...
- n. "{text_example_n}". Answer: "{answer_example_n}"

Your Turn

Input:

- Text: "{text}"

Answer:

Figure D: Few-shot prompt for ENVIS-EN.

Tugas: Tentukan semua frasa sentimen dan polaritas sentimennya pada input berikut. Untuk setiap frasa sentimen yang diekstrak, jika ada, tentukan juga pemegang dan target dari frasa sentiment. Mohon untuk tidak menjelaskan jawaban Anda.

Instruksi:

- Polaritas sentimen hanya dapat berupa "negatif" atau "positif".
- Format output harus berupa ["pemegang", "target", "frasa sentimen", "polaritas sentimen"].
- Jika tidak ada pemegang dan/atau target dari suatu frasa sentimen yang dieksteaksi, mohon isi dengan string kosong "".
- Jika sama sekali tidak ada frasa sentimen, maka hanya jawab [].
- Jangan berikan penjelasan apapun terkait jawaban Anda.

Input:

- Teks: "{text}"

Jawaban:

Figure E: Zero-shot prompt for ENVIS-ID.

Tugas: Tentukan semua frasa sentimen dan polaritas sentimennya pada input berikut. Untuk setiap frasa sentimen yang diekstrak, jika ada, tentukan juga pemegang dan target dari frasa sentiment. Anda juga diberikan beberapa contoh. Mohon untuk tidak menjelaskan jawaban Anda.

Instruksi:

- Polaritas sentimen hanya dapat berupa "negatif" atau "positif".
- Format output harus berupa ["pemegang", "target", "frasa sentimen", "polaritas sentimen"].
- Jika tidak ada pemegang dan/atau target dari suatu frasa sentimen yang dieksteaksi, mohon isi dengan string kosong "".
- Jika sama sekali tidak ada frasa sentimen, maka hanya jawab [].
- Contoh yang diberikan mungkin dapat membantu Anda dalam menentukan jawaban.
- Jangan berikan penjelasan apapun terkait jawaban Anda.

Contoh:

1. "{text_example_1}". Jawaban: "{answer_example_1}"
- ...
- n. "{text_example_n}". Jawaban: "{answer_example_n}"

Input:

- Teks: "{text}"

Jawaban:

Figure F: Few-shot prompt for ENVIS-ID.

F Token and Topic Variability Details

1221

Statistic	# of Unique Token		
	ENVIS-IT	ENVIS-EN	ENVIS-ID
holder	42	113	87
target	229	1245	1404
sentiment term	945	2849	2093

Table F: Number of unique tokens for holder, target, and sentiment term expression.

Statistic	ENVIS-IT	ENVIS-EN	ENVIS-ID
Environment	2.80	18.53	37.83
Green	100.00	15.40	7.35
Sustainability	0.80	12.53	4.30
Food	0.30	15.93	14.09
Organism	0.10	10.67	4.74
Climate Change	6.80	16.53	5.74
Carbon	2.90	15.13	16.43
Energy	7.70	17.60	40.09
Waste	1.20	13.87	21.04
Pollution	4.30	15.33	1.65
tweets discussing multi-topics	23.10	39.20	37.39

Table G: Topic discussed in ENVIS, classified based on keyword-matching (Appendix A). All scores are in percent correspond to total tweets for each language. ENVIS-IT has 100% in topic “Green” as [Bosco et al. \(2023\)](#) focus on tweets that discussing “green” keyword.

G Quantitative Error Types Details

1222

For the group of error types, we follow the categorization defined by [Oberländer and Klinger \(2020\)](#), which is also followed by [Barnes et al. \(2022\)](#) in their SSA error analysis. The statistic details for each error type in each language for best LLMs and USSA can be seen in Table H.

1223

1224

1225

Model	Error Type	ENVIS-IT			ENVIS-EN			ENVIS-ID			Avg. (%)
		<i>h</i> (%)	<i>t</i> (%)	<i>e</i> (%)	<i>h</i> (%)	<i>t</i> (%)	<i>e</i> (%)	<i>h</i> (%)	<i>t</i> (%)	<i>e</i> (%)	
Mistral-7B (3-sets Context 2)	False Positive	36.84	25.1	33.65	58.1	11.76	15.31	77.48	13.14	28.88	33.36
	False Negative	31.58	36.53	31.4	20.73	43.9	34.65	11.07	42.38	31.81	31.56
	Too Early	0.00	0.20	1.21	0.29	0.10	1.10	0.19	1.11	1.54	0.64
	Too Late	0.00	0.41	0.69	0.15	0.10	0.82	0.19	0.54	3.05	0.66
	Multiple	31.58	37.76	32.96	20.73	44.1	37.48	11.07	42.81	34.67	32.57
	Other	0.00	0.00	0.09	0.00	0.05	10.64	0.00	0.04	0.04	1.21
USSA-mBERT	False Positive	0.00	0.00	0.00	1.56	6.45	1.99	0.00	11.89	4.43	2.92
	False Negative	50.00	50.00	50.00	49.09	46.53	48.48	50.00	42.78	46.87	48.19
	Too Early	0.00	0.00	0.00	0.26	0.20	0.06	0.00	1.08	0.99	0.29
	Too Late	0.00	0.00	0.00	0.00	0.10	0.22	0.00	0.20	0.18	0.08
	Multiple	50.00	50.00	50.00	49.09	46.73	49.21	50.00	43.97	47.51	48.50
	Other	0.00	0.00	0.00	0.00	0.00	0.03	0.00	0.08	0.02	0.01

Table H: Quantitative error types details for best LLM and best USSA. Most ones (in average) per error type per model are in bold.

H Qualitative Error Types Examples

(2) Example of Missing Information:

- Input:
"india reduced emission intensity of its gdp by 24 per cent in 11 yrs through 2016"
- Ground Truth:
[["", "india", "reduced emission", "positive"]]
- Generated Output:
[["", "India", "reduced emission intensity of its gdp by 24 per cent in 11 yrs through 2016"]]

(3) Example of Incorrect Generation:

- Input:
"they want you to spill the oil that taylor extracted from the ground all over the ground bc taylor wants to cause an environmental disaster along with the help of her fans"
- Ground Truth:
[["", "taylor", "wants to cause an environmental disaster", "negative"]]
- Generated Output:
[["they", "want", "bc", "they", "want", "to cause", "an", "environmental", "disaster", "along", "with", "the", "help", "of", "her", "fans", "she", "wants"]]

(4) Example of Miscellaneous:

- Input:
"animal agriculture as we know it will die out once plant based and cultured meat alternatives find broad acceptance the big boys cargill adm tyson are already getting in position because current practices are so inefficient and harmful to the environment"
- Ground Truth:
[["we", "", "die out", "negative"], ["", "current practices", "inefficient and harmful", "negative"]]
- Generated Output:
- "[["animal agriculture", "", "", "will die out"], ["animal agriculture", "", "", "once plant based"], ["animal agriculture", "", "", "once cultured meat alternatives"], ["the big boys", "", "", "are already getting in position"], ["cargill", "", "", "are already getting in position"], ["adm", "", "", "are already getting in position"], ["", "", "because current practices", "", ""], ["", "", "are so inefficient", "", ""], ["", "", "are harmful to the environment", "", ""]]"
- sentiment polarity: positive.
explanation:
the sentiment phrases are "animal agriculture will die out", "once plant based", "once cultured meat alternatives", and "are already getting in position". the sentiment polarity for all these phrases is positive. the holders for these phrases are "animal agriculture", "the big boys", and empty string for the last two phrases. the targets for these phrases are empty strings for the first three phrases and "current practices" and "the environment" for the last two phrases.