

Social Support Detection and Classification in Social-media Comments

Alan Rodrigo López López^{1,*†}, Christian Carlos González Almonte^{1,*†}, Olga Kolesnikova^{1,*} and Francisco Hiram Calvo Castro^{1,*}

¹Computing Research Center, Av. Juan de Dios Bátiz S/N, Esq. Miguel Othón de Mendizábal, Col. Nueva Industrial Vallejo, Gustavo A. Madero, 07738, Mexico City, Mexico

Abstract

This paper describes the results obtained from our participation in the Social Support Detection shared task at IberLEF 2026, which focuses on detecting social support in social media comments and identifying the target of such support. We addressed the three proposed subtasks: binary support detection, binary target identification, and targeted group multiclass classification. We evaluated several supervised text classification approaches based on traditional machine learning classifiers, sentence-transformer embeddings and pre-trained transformer-based language models. We deployed models such as RoBERTa, BERT, SpanBERTa, and XLM-RoBERTa. Our proposed approaches achieved Macro-F1 scores of 0.7798 for English and 0.8514 for Spanish in Subtask 1, 0.8225 and 0.8123 in Subtask 2, and 0.8096 and 0.8148 in Subtask 3, respectively. The results suggest that contextual language models are effective for binary support-related classification, while sentence embeddings combined with linear classifiers provide a competitive alternative for fine-grained target classification.

Keywords

Social Support Detection, Text Classification, Social Media Comments, Transformer-based Models, Sentence Embeddings, IberLEF 2026

1. Introduction

Moving beyond basic sentiment analysis toward the actionable understanding of prosocial language, the task challenges participants to build models capable of detecting social support in English and Spanish social media comments, as well as classifying the target, distinguishing between individuals and specific communities.

This is particularly challenging due to the unpredictable nature of online content which frequently deviates from standard grammatical structures. Online users heavily rely on slang, disregard grammar and spelling rules, and often employ sarcasm, making it difficult for automated systems to accurately capture the intent behind a message.

To address the noisy nature of user-generated content, the current state of the art relies heavily on Transformer-based pre-trained language models. Models explicitly pre-trained on social media corpora, such as *BERTweet* [1], as well as robust multilingual architectures like *XLM-RoBERTa* [2], have established strong baselines for parsing informal linguistic nuances across different languages.

In this paper, we present our approach to tackling these linguistic complexities and describe the architecture of our submitted model for the SSD-2026 shared task [3].

IberLEF 2026, September 2026, León, Spain

*Corresponding author.

†These authors contributed equally.

✉ alanrodrigolopez@gmail.com (A. R. L. López); chrisalmonte@protonmail.com (C. C. G. Almonte); kolesnikova@cic.ipn.mx (O. Kolesnikova); hcalvo@cic.ipn.mx (F. H. C. Castro)

🌐 <https://github.com/AlanLopez1017> (A. R. L. López); <https://github.com/chrisalmonte> (C. C. G. Almonte)

🆔 0000-0002-1307-1647 (O. Kolesnikova); 0000-0003-2836-2102 (F. H. C. Castro)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

2. Task Description

The objective of this task is to build and evaluate natural language processing systems that detect social support in social media comments and identify the subject of that support [4]. It is motivated by an actionable understanding of prosocial language, allowing for realistic end-to-end evaluation for implementation scenarios.

The shared task is divided into different subtasks, each targeting specific data fields in the provided datasets [5, 6, 7, 8].

2.1. Subtask 1: Binary Support Detection

For both languages, this subtask takes a single social media comment as input. If the comment expresses encouragement, care, admiration, help, or solidarity toward a target, it is classified as “Support”; otherwise, it is labeled as “Not Support”.

2.2. Subtask 2: Binary Target Identification

Its input is a single supportive comment. The objective is to determine whether the comment is supporting specific a person (“Individual”) or a collective entity (“Group”).

2.3. Subtask 3: Targeted Group Multiclass Classification

Given a single supportive comment directed at a specific group, the task is to classify the supported community into one of six categories: Nation, LGBTQ, Black Community, Religion, and Women or Other. As a boundary condition, supportive comments that explicitly promote violence or harm are labeled as “Not Support”.

3. Dataset description

In the Subtask 1 English dataset, the output classes are coded as 0 and 1, for Non-Supportive and Supportive, respectively. The class distribution is shown in Table 1. It can be observed that the dataset is imbalanced, with Non-Supportive comments representing 77.56% of the instances. The total number of instances is 7998.

Table 1

Class distribution for Subtask 1 in the English dataset

Class	Label	Instances	Percentage
Non-Supportive	0	6203	77.56%
Supportive	1	1795	22.44%

The class distribution for Subtask 1 in Spanish is similar to that of its English counterpart, where the predominant class is Non-Supportive, accounting for 77.31% of the data. The dataset comprises a total of 2600 instances. The distribution is shown in Table 2.

Table 2

Class distribution for Subtask 1 in the Spanish dataset.

Class	Label	Instances	Percentage
Non-Supportive	0	2010	77.31%
Supportive	1	590	22.69%

In the Subtask 2 English dataset, the output classes correspond to whether the support is directed towards an individual or a group. Filtering strictly for instances with valid individual or group labels. We can see that the dataset is heavily imbalanced, with the Group class representing 80.78% of the instances. The total number of instances for this subtask is 1795.

Table 3

Class distribution for Subtask 2 in the English dataset.

Class	Label	Instances	Percentage
Group	Group	1450	80.78%
Individual	Individual	345	19.22%

As for the Spanish dataset (train-spanish.csv), the class distribution is shown in Table 4. We can see that the dataset is heavily imbalanced, with the Group class representing 78.35% of the instances. The total number of instances for this subtask is 619.

Table 4

Class distribution for Subtask 2 in the Spanish dataset.

Class	Label	Instances	Percentage
Group	Group	485	78.35%
Individual	Individual	134	21.65%

The class distribution for Subtask 3 in English is shown in Table 5. The majority class is Nation, with 786 instances, followed by Other (416 instances) and LGBTQ (123 instances). Conversely, the Religion class contains only 15 instances. The total number of instances is 1450.

Table 5

Class distribution for Subtask 3 in the English dataset.

Class	Instances	Percentage
Nation	786	54.21%
Other	416	28.69%
LGBTQ	123	8.48%
Black Community	91	6.28%
Women	19	1.31%
Religion	15	1.03%

As for the Spanish dataset, the class distribution is shown in Table 6. We can see that the dataset is imbalanced, with the LGBTQ class being the predominant group, representing 46.19% of the instances. The total number of instances for this subtask is 485.

Table 6

Class distribution for Subtask 3 in the Spanish dataset.

Class	Instances	Percentage
LGBTQ	224	46.19%
Other	82	16.91%
Religion	60	12.37%
Women	48	9.90%
Nation	47	9.69%
Black Community	24	4.95%

4. Proposed Methodology

4.1. Text Preprocessing

A minimal preprocessing strategy was applied to preserve the informal linguistic cues commonly found in social media texts. Furthermore, the provided data for the English dataset was already lemmatized, allowing for less intervention during cleaning. Preprocessing included conversion to lowercase and normalization of whitespace.

4.2. Feature Extraction

For Subtask 1 and 2 in English, we used the pre-trained roberta-base model as the feature extraction backbone. The input was tokenized using the corresponding RoBERTa tokenizer.

For Subtask 1, each comment was truncated to a maximum length of 128 tokens, while no limit was applied for Subtask 2. Dynamic padding was applied during batching using a data collator, allowing each batch to be padded according to the longest sequence in that batch.

For Subtask 1 in Spanish, the backbone was the pre-trained dccuchile/bert-base-spanish-wwm-cased model. As in English, each comment was truncated to a maximum length of 128 tokens and dynamic padding was also applied.

For Subtask 2 in Spanish, the backbone was the pre-trained skimai/spanberta-base-cased model, which was trained over the “OSCAR” corpus [9], using its corresponding tokenizer.

For Subtask 3 in English, each input comment was represented using sentence-level embeddings generated with all-MiniLM-L6-v2. The embeddings were computed with a batch size of 32. These embeddings were used as a fixed feature extractor.

For Subtask 3 in Spanish, the pre-trained XLM-RoBERTa model was used as a backbone feature extractor. Its corresponding tokenizer was used. No limits were imposed on the number of tokens for each comment.

4.3. Classification Models

For Subtask 1 the classification models were based on FacebookAI/roberta-base for English and dccuchile/bert-base-spanish-wwm-cased for Spanish, both adjusted for binary classification, where classification heads with two output neurons were added on top of each pre-trained encoder to predict the two labels. The tag 1 was considered the positive class, corresponding to the category *Supportive*, for both languages.

Similarly, for Subtask 2, FacebookAI/roberta-base and skimai/spanberta-base-cased were used as foundational models for English and Spanish, respectively. A classification head with two output neurons was added on top of the pre-trained models. The tag 1 was considered the positive class, corresponding to the category *Individual*, for both languages.

The same strategy was used for Subtask 3 in Spanish, where a classification head with six output neurons was added on top of FacebookAI/xlm-roberta-base.

For Subtask 3 in English, a multiclass Linear Support Vector Classifier was trained using the sentence embeddings as input.

4.4. Experimental Settings

For Subtask 1 in both English and Spanish the model was evaluated using five-fold stratified cross-validation. For each fold, the model was tuned for three epochs with learning rate of 2×10^{-5} , a batch size of 8 for both training and evaluation, and a weight decay of 0.01. The maximum sequence length was 128. The best model checkpoint was selected according to the validation F1-score.

For Subtask 2 in both English and Spanish, the model was tuned for 15 epochs with a learning rate of 2×10^{-5} , a batch size of 16 for both training and evaluation, and a weight decay of 0. The best model checkpoint was selected according to the validation average loss.

The Subtask 3 English classifier was evaluated using 5-fold stratified cross-validation. An embedding batch size of 32 was used. It was implemented using a regularization parameter of $C = 0.1$, squared_hinge loss.

For Subtask 3 in Spanish the model was tuned for 15 epochs with a learning rate of 2×10^{-5} , a batch size of 16 for both training and evaluation, and a weight decay of 0. The best model checkpoint was selected according to the validation average loss.

5. Results

5.1. Subtask 1: Binary Support Detection

5.1.1. English dataset

Subtask 1 for the English dataset was evaluated using six different classifiers. Among the evaluated models, the best performing model was RoBERTa-base, achieving a Macro F1 score of 0.8148 ± 0.0108 . The second-best result was obtained by tuned Word TF-IDF + Logistic Regression model, with a Macro F1 of 0.7971 ± 0.0079 . The remaining classifiers included SVM, and SGD. The corresponding results for all models are shown in Table 7.

Table 7

Cross-validation results for Subtask 1 on the English dataset.

Model	Macro F1	Std.
RoBERTa-base	0.8148	0.0108
Word TF-IDF + Logistic Regression tuned	0.7971	0.0079
Word TF-IDF + Logistic Regression	0.7852	0.0072
Word TF-IDF + LinearSVC	0.7753	0.009
Char TF-IDF + SGD	0.7725	0.0069
Word TF-IDF + MultinomialNB	0.7008	0.0094

Since RoBERTa-base was the best-performing model, the specific classification results are presented in detail. The model across the five cross-validation folds correctly classified 5652 out of 6203 Non-Supportive comments and 1303 out of 1795 Supportive comments. 492 Supportive instances were misclassified as Non-Supportive, indicating that the detection of supportive comments remains more challenging.

The detailed cross-validation results obtained for the RoBERTa model are provided in the Table 8. The evaluation metric employed is the macro-averaged F1-score, which achieved a value of 0.8148 ± 0.0108 . In addition to this, the model obtained a Weighted F1 of 0.8703 ± 0.0074 .

Table 8

Cross-validation results for the RoBERTa model on Subtask 1 for English dataset

Metric	Result
Accuracy	0.8696 ± 0.0074
Precision (Supportive)	0.7029 ± 0.0169
Recall (Supportive)	0.7259 ± 0.0216
F1 (Supportive)	0.7141 ± 0.0168
Macro F1	0.8148 ± 0.0108
Weighted F1	0.8703 ± 0.0074

Finally, on the official Codabench test set, the submitted RoBERTa model achieved a **Macro F1 Score** of:

$$\text{Macro F1} = 0.7798$$

5.1.2. Spanish dataset

Subtask 1 for the Spanish dataset was evaluated using eight different classifiers. The best-performing model was BETO, achieving a Macro F1 of 0.8351 ± 0.0132 . It was followed by XLM-RoBERTa (XLM-R) and tuned TF-IDF + LinearSVC model, which obtained Macro F1 score of 0.8201 ± 0.0101 and 0.8125 ± 0.0065 , respectively. The corresponding results for all models are shown in Table 9.

Table 9

Cross-validation results for Subtask 1 on the Spanish dataset.

Model	Macro F1	Std.
BETO	0.8351	0.0132
XLM-R	0.8201	0.0101
Char TF-IDF + LinearSVC tuned	0.8125	0.0065
Char TF-IDF + LinearSVC	0.7986	0.0164
Word TF-IDF + Logistic Regression	0.7946	0.0190
Word TF-IDF + LinearSVC	0.7862	0.0164
Char TF-IDF + SGD	0.7823	0.0151
Word TF-IDF + ComplementNB	0.7474	0.0257

Since BETO was the best-performing model, we present its results obtained. The model correctly classified 1874 out of 2010 Non-Supportive comments and 430 out of 590 Supportive. As in English dataset, Supportive class remains more difficult to identify.

The detailed cross-validation results obtained for the BETO model are provided in the Table 10. The model achieved an Accuracy of 0.8862 ± 0.0071 , a Macro F1 of 0.8351 ± 0.0132 , and a Weighted of 0.8852 ± 0.0082 .

Table 10

Cross-validation results for the BETO model on Subtask 1 for the Spanish dataset.

Metric	Result
Accuracy	0.8862 ± 0.0071
Precision (Supportive)	0.7601 ± 0.0149
Recall (Supportive)	0.7288 ± 0.0444
F1 (Supportive)	0.7434 ± 0.0226
Macro F1	0.8351 ± 0.0132
Weighted F1	0.8852 ± 0.0082

The result obtained on the Codabench platform for the test dataset was:

$$\text{Macro F1} = 0.8514.$$

5.2. Subtask 2: Binary Target Identification

Unlike the previous subtasks, cross-validation was not utilized for this configuration. Instead, we applied a single stratified split exclusively to the supportive comments. The train dataset was partitioned into 75% and 25% for validation.

5.2.1. English

The English model achieved its lowest validation loss of 0.2773 at epoch 3. The validation results are described in Table 11.

Table 11

Validation results for the RoBERTa model on Task 2

Metric	Result
Accuracy	0.8832
Precision (Individual)	0.70
Recall (Individual)	0.71
F1 (Individual)	0.70
Macro F1	0.81
Weighted F1	0.88

The result obtained on the Codabench platform for the test dataset was:

$$\text{Macro F1} = 0.8225.$$

5.2.2. Spanish

Experiments were made using both SpanBERTa and XLM-RoBERTa. Table 12 shows a comparison of their validation results.

Table 12

Model comparison for Subtask 2 in Spanish

Model	Macro F1
SpanBERTa	0.79
XLM-RoBERTa	0.73

The SpanBERTa model reached its optimal validation average loss of 0.3225 at epoch 4, while XLM-RoBERTa reached it at epoch 7 with a value of 0.4049. The validation results for the best model (SpanBERTa) are described in Table 13.

Table 13

Validation results for the SpanBERTa model on Task 2.

Metric	Result
Accuracy	0.8512
Precision (Individual)	0.64
Recall (Individual)	0.69
F1 (Individual)	0.67
Macro F1	0.79
Weighted F1	0.85

The result obtained on the Codabench platform for the test dataset was:

$$\text{Macro F1} = 0.8123.$$

5.3. Subtask 3: Targeted Group Multiclass Classification

5.3.1. English

Subtask 3 for the English dataset was evaluated using nine configurations. The best validation result was achieved by the tuned all-MiniLM-L6-v2 + LinearSVC, with a Macro F1 of 0.7648 ± 0.0498 . This result was very close to the tuned Word TF-IDF with LinearSVC, which obtained a Macro F1 of 0.7639 ± 0.0686 . The corresponding results for all models are shown in Table 14.

Table 14

Cross-validation results for Subtask 3 on the English dataset.

Model	Macro F1	Std.
all-MiniLM-L6-v2 + LinearSVC tuned	0.7648	0.0498
Word TF-IDF + LinearSVC tuned	0.7639	0.0686
all-MiniLM-L6-v2 + LinearSVC	0.7605	0.0549
all-MiniLM-L6-v2 + Logistic Regression	0.7516	0.0525
Word TF-IDF + Logistic Regression	0.7369	0.0595
Word TF-IDF + LinearSVC	0.7345	0.0691
Char TF-IDF + LinearSVC	0.7243	0.0545
Char TF-IDF + SGD	0.6990	0.0414
Word TF-IDF + ComplementNB	0.5892	0.0495

The embedding-based LinearSVC was the best-performing model. The model achieved very strong performance for the *Black Community* and *LGBTQ* class, correctly classifying 91 out of 91 and 122 out of 123 instances, respectively. Most errors occurred in the *Other*, *Religion* and *Women* classes, which is consistent with the lower number of training examples for those categories.

The detailed cross-validation results obtained for the selected multiclass classifier are shown in the Table 15. The model achieved an Accuracy of 0.8690 ± 0.0189 , a Macro F1 of 0.7648 ± 0.0498 and a Weighted F1 of 0.8674 ± 0.0206 . We can see that the difference between Macro and Weighted F1 reflects the class imbalance in the dataset, where *Nation* and *Other* have a stronger influence.

Table 15

Cross-validation results for the selected all-MiniLM-L6-v2 + LinearSVC model on Subtask 3 for the English dataset.

Metric	Result
Accuracy	0.8690 ± 0.0189
Macro Precision	0.7651 ± 0.0375
Macro Recall	0.7980 ± 0.0622
Macro F1	0.7648 ± 0.0498
Weighted Precision	0.8735 ± 0.0187
Weighted Recall	0.8690 ± 0.0189
Weighted F1	0.8674 ± 0.0206
Micro F1	0.8690 ± 0.0189

The result obtained on the Codabench platform for the test dataset was:

$$\text{Macro F1} = 0.8096.$$

5.3.2. Spanish

Using the same training strategy as in Subtask 2, the dataset was partitioned into 75% and 25% for validation. Experiments were made using both SpanBERTa and XLM-RoBERTa. Table 16 shows a comparison of their validation results.

Table 16

Model comparison for Subtask 3 in Spanish

Model	Macro F1
SpanBERTa	0.79
XLM-RoBERTa	0.95

For SpanBERTa, the lowest validation loss was achieved at epoch 4, with a value of 0.3203, while XLM-RoBERTa achieved it at epoch 10, with a value of 0.3009. Table 17 shows the validation results for the best model (XLM-RoBERTa).

Table 17

Validation results for the XLM-RoBERTa model on Task 3.

Metric	Result
Accuracy	0.95
F1 (Nation)	0.93
F1 (Other)	0.90
F1 (LGBTQ)	0.97
F1 (Black Community)	1.00
F1 (Religion)	0.95
F1 (Women)	0.93
Macro F1	0.95
Weighted F1	0.95

The result obtained on the Codabench platform for the test dataset was:

$$\text{Macro F1} = 0.8148.$$

6. Discussion

The experimental results show that Transformer-based models (RoBERTa and BETO) provided the strongest performance across most subtasks, while an embedding-based linear classifier achieved the best result for Subtask 3 on the English dataset. This suggests that contextual representations are useful for detecting supportive language, which often depends not only on specific lexical cues but also on the semantic relation between the comment and its target.

Regarding the subtasks in Spanish, BETO and SpanBERTa achieved a better performance than XLM-RoBERTa for Subtask 1 and 2, indicating that a language-specific model trained on Spanish data can be advantageous for text in Spanish for their respective tasks. However, for Subtask 3 in Spanish, XLM-RoBERTa outperformed SpanBERTa. This suggests that the cross-lingual representations acquired during multilingual pre-training are beneficial for distinguishing the nature of the collective subjects.

For Subtask 1 both models performed better on the Non-Supportive class that may be expected given the class imbalance. This point is relevant because the main objective of the task is to detect supportive content. Therefore, it reveals that identifying supportive comments is more challenging than non-supportive ones.

A comparison between validation and official test performance also suggests a moderate generalization gap. In English Subtask 1, RoBERTa achieved a cross-validation Macro F1 of 0.8148 ± 0.0108 , while the official Codabench test was 0.7798. This difference may be due to distributional differences between the training and hidden test sets.

For Subtask 2, adopting the best-performing checkpoints led to official challenge Macro F1-scores of 0.8123 and 0.8225 for Spanish and English, respectively. Despite these strong overall results, both models exhibited lower predictive power on the minority class (*Individual*) compared to the majority class (*Group*) due to baseline dataset imbalances. This suggests that while pre-trained language models adapt quickly to the binary target classification task, class-weighting or specialized sampling strategies could further stabilize minority class representation.

The results for Subtask 3 indicate that classes with a larger number of instances performed well. Conversely, classes characterized by data scarcity posed significant challenges for the classifier in establishing stable decision boundaries.

7. Conclusions and Future Work

Overall, the results suggest that contextual transformer models are effective for binary support detection. For multiclass classification, fixed sentence embeddings combined with a linear classifier offer a strong alternative. Nevertheless, class imbalance remains a critical factor to consider.

Although both language-specific and multilingual models exhibit robust classification capabilities, determining the most advantageous approach for specific downstream tasks remains an open area for evaluation.

Future work will explore data augmentation strategies for minority classes, incorporate class balancing techniques, investigate more robust classifiers and evaluate different fine-tuning strategies, which may provide an increase in overall performance.

Acknowledgments

The authors thank the organizers of the SSD-2026 shared task at IberLEF 2026 for providing the dataset and the evaluation framework. This work was supported in part by grants 20260496 and 20260569 from the Secretary of Research and Postgraduate Studies (SIP) of the Instituto Politécnico Nacional, Mexico.

Declaration on Generative AI

During the preparation of this work, the authors used GPT-5.5 for LaTeX formatting and manuscript organization. After using this tool, the authors reviewed and edited the content as needed, and take full responsibility for the publication's content.

References

- [1] D. Q. Nguyen, T. Vu, A. T. Nguyen, BERTweet: A pre-trained language model for English Tweets, in: Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations, 2020, pp. 9–14.
- [2] A. Conneau, K. Khandelwal, N. Goyal, V. Chaudhary, G. Wenzek, F. Guzmán, E. Grave, M. Ott, L. Zettlemoyer, V. Stoyanov, Unsupervised cross-lingual representation learning at scale, in: Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics, 2020, pp. 8440–8451.
- [3] M. Shahiki Tash, L. Ramos, A. Zahra, A. Y. Barrera Animas, F. Balouchzahi, O. Kolesnikova, G. Sidorov, A. Gelbukh, R. Monroy, F. Vargas, Overview of SSD-2026 at IberLEF 2026: Social Support Detection, in: Proceedings of the Iberian Languages Evaluation Forum (IberLEF 2026), co-located with the 42nd Conference of the Spanish Society for Natural Language Processing (SEPLN 2026), CEUR-WS.org, 2026.
- [4] A. Bonet-Jover, J. A. González-Barba, L. Chiruzzo, Overview of IberLEF 2026: Natural Language Processing Challenges for Spanish and other Iberian Languages, in: Proceedings of the Iberian Languages Evaluation Forum (IberLEF 2026), co-located with the 42nd Conference of the Spanish Society for Natural Language Processing (SEPLN 2026), CEUR-WS.org, 2026.
- [5] Z. Ahani, M. S. Tash, F. Balouchzahi, L. Ramos, G. Sidorov, A. Gelbukh, Social support detection from social media texts, PLOS ONE (2026). URL: <https://doi.org/10.1371/journal.pone.0337476>. doi:10.1371/journal.pone.0337476.
- [6] O. Kolesnikova, M. S. Tash, Z. Ahani, A. Agrawal, R. Monroy, G. Sidorov, Advanced machine learning techniques for social support detection on social media, Heliyon 11 (2025). URL: <https://doi.org/10.1016/j.heliyon.2025.e43437>. doi:10.1016/j.heliyon.2025.e43437.
- [7] M. S. Tash, L. Ramos, Z. Ahani, R. Monroy, H. Calvo, G. Sidorov, Online social support detection in Spanish social media texts, arXiv preprint arXiv:2502.09640 (2025). URL: <https://arxiv.org>.
- [8] M. Shahiki-Tash, Z. Ahani, L. Ramos, et al., Linguistic analysis in different types of support using LIWC for Spanish and English online communities, Research Square (2024). URL: <https://doi.org/10.21203/rs.3.rs-9557287/v1>. doi:10.21203/rs.3.rs-9557287/v1, preprint.
- [9] P. J. Ortiz Suarez, B. Sagot, L. Romary, Asynchronous pipelines for processing huge corpora on medium to low resource infrastructures, in: P. Bański, A. Barbaresi, H. Biber, E. Breiteneder, S. Clematide, M. Kupietz, H. Lungen, C. Iliadi (Eds.), Proceedings of the Workshop on Challenges in the Management of Large Corpora (CMLC-7) 2019. Cardiff, 22nd July 2019, 2019, pp. 9 – 16. doi:10.14618/ids-pub-9021.