

# Extended Multilingual Protest News Detection - Shared Task 1, CASE 2021 and 2022

**Ali Hürriyetoglu**

KNAW Humanities Cluster DHLab  
ali.hurriyetoglu@dh.huc.knaw.nl

**Osman Mutlu**

Koc University  
omutlu@ku.edu.tr

**Fırat Duruşan**

Koc University  
fdurusan@ku.edu.tr

**Onur Uca**

Mersin University  
onuruca@mersin.edu.tr

**Alaeddin Selçuk Gürel**

Huawei  
alaeddinselcukgurel@gmail.com

**Benjamin Radford**

UNC Charlotte  
benjamin.radford@uncc.edu

**Yaoyao Dai**

UNC Charlotte  
yaoyao.dai@uncc.edu

**Hansi Hettiarachchi**

Birmingham City University  
hansi.hettiarachchi@mail.bcu.ac.uk

**Niklas Stoehr**

ETH Zurich  
niklas.stoehr@inf.ethz.ch

**Tadashi Nomoto**

National Institute of Japanese Literature  
nomoto@acm.org

**Milena Slavcheva**

IICT, Bulgarian Academy of Sciences  
milena@lml.bas.bg

**Francielle Vargas**

University of São Paulo  
francielleavargas@usp.br

**Aaqib Javid**

Koc University  
ajavid20@ku.edu.tr

**Fatih Beyhan**

Sabancı University  
fatihbeyhan@sabanciuniv.edu

**Erdem Yörük**

Koc University  
eryoruk@ku.edu.tr

## Abstract

We report results of the CASE 2022 Shared Task 1 on Multilingual Protest Event Detection. This task is a continuation of CASE 2021 that consists of four subtasks that are i) document classification, ii) sentence classification, iii) event sentence coreference identification, and iv) event extraction. The CASE 2022 extension consists of expanding the test data with more data in previously available languages, namely, English, Hindi, Portuguese, and Spanish, and adding new test data in Mandarin, Turkish, and Urdu for Sub-task 1, document classification. The training data from CASE 2021 in English, Portuguese and Spanish were utilized. Therefore, predicting document labels in Hindi, Mandarin, Turkish, and Urdu occurs in a zero-shot setting. The CASE 2022 workshop accepts reports on systems developed for predicting test data of CASE 2021 as well. We observe that the best systems submitted by CASE 2022 participants achieve between 79.71 and 84.06 F1-macro for new languages in a zero-shot setting. The winning approaches are mainly ensembling models and merging data in multiple languages. The best two submissions on CASE 2021 data outperform submissions from last year for Subtask 1 and Subtask 2 in all languages. Only the following scenarios were not outperformed by new submissions on CASE 2021: Subtask 3 Portuguese & Subtask 4 English.

## 1 Introduction

We aim at determining event trigger and its arguments in a text snippet in the scope of an event extraction task. The performance of an automated system depends on the target event type as it may be broad or potentially the event trigger(s) can be ambiguous. The context of the trigger occurrence may need to be handled as well. For instance, the ‘protest’ event type may be synonymous with ‘demonstration’ or not in a specific context. Moreover, the hypothetical cases such as future protest plans may need to be excluded from the results. Finally, the relevance of a protest depends on the actors as in a contentious political event only citizen-led events are in the scope. This challenge is even harder in a cross-lingual and zero-shot setting in case training data are not available in new languages.

We provide a benchmark that consists of four subtasks and multiple languages in the scope of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text at The 2022 Conference on Empirical Methods in Natural Language Processing (CASE @ EMNLP 2022) (Hürriyetoglu et al., 2022).<sup>1</sup> bgenfrhumil: To paraphrase: The work presented

<sup>1</sup><https://emw.ku.edu.tr/case-2022/>, accessed on November 13, 2022.

in this paper is a continuation of the work initiated in CASE 2021 Task 1 (Hürriyetoglu et al., 2021) and consists in adding new documents in already available languages, as well as adding new languages to the evaluation data.

Task 1 consists of the following subtasks that ensure the task is tackled incrementally: i) Document classification, ii) Sentence classification, iii) event sentence coreference identification, and iv) event extraction. The training data consist of documents in English, Portuguese, and Spanish, while the evaluation texts are in English, Hindi, Mandarin, Portuguese, Spanish, Turkish, and Urdu. Subtask 1 ensures documents with relevant senses of event triggers are selected. Next, Subtask 2 focuses on identifying event sentences in a document. Discriminating sentences that are about separate events and grouping them is done in Subtask 3 (Hürriyetoglu et al., 2020, 2022). Finally, the sentences that are about the same events are processed to identify the event trigger and its arguments in Subtask 4. In addition to accomplishing the event extraction task, the subtask division improves significantly the annotation quality, as the annotation team can focus on a specific part of the task and errors in previous levels are corrected during the preparation of the following subtask (Hürriyetoglu et al., 2021). The significance of this specific task division is twofold: i) facilitating the work with a random sample of documents by first identifying relevant documents and sentences before annotating or processing a sample or a complete archive of documents respectively; ii) increasing the generalizability of the automated systems that may be developed using this data (Yörük et al., 2021; Mutlu, 2022).

The current report is about Task 1 in the scope of CASE 2022. Task 2 (Zavarella et al., 2022) and Task 3 (Tan et al., 2022b,a) complement Task 1 by evaluating Task 1 systems on events related to COVID-19 and detecting causality respectively.

The following section, which is Section 2 describes the data we use for the shared task. Next we describe the evaluation setting in Section 3. The results are provided in Section 4. Finally, the Section 5 conclude this report.

## 2 Data

We used the CASE 2021 training data as those for CASE 2022.<sup>2</sup> The CASE 2022 test data are the

<sup>2</sup><https://github.com/emerging-welfare/case-2021-shared-task> for CASE 2021

union of CASE 2021 test data and additional new documents in both available and new languages. The new languages are Mandarin, Turkish, and Urdu.

The new document level data, which are used to extend CASE 2021 data, were randomly sampled from MOT v1.2 (Palen-Michel et al., 2022)<sup>3</sup> and were annotated by co-authors of this report. Documents were annotated by native speakers of the respective language. A single label was attached to each document. The annotation manual followed in the annotation process (Duruşan et al., 2022) was the same as that used in CASE 2021. The total number of CASE 2022 documents with labels is 3,870 for English, 267 for Hindi, 300 for Mandarin, 670 for Portuguese, 399 for Spanish, 300 for Turkish, and 299 for Urdu. Teams that developed systems for Subtasks 2, 3, and 4 evaluated their systems on CASE 2021 test data.

## 3 Evaluation setting

We utilized Codalab for evaluation of Task 1 for CASE 2022.<sup>4</sup> The evaluation for CASE 2021 was performed on an additional scoring page<sup>5</sup> of the original<sup>6</sup> CASE 2021 Codalab page. Moreover, we launched an additional scoring page for CASE 2022 after completion of the official evaluation period.<sup>7</sup> Five submissions per subtask and language pair could be submitted in total for CASE 2022. The additional scoring phase of both CASE 2021 and CASE 2022 allow only one submission per subtask and language combination per day. The test data of CASE 2021 were shared with participants at the same time with the training data. But the CASE 2022 evaluation data were shared around two weeks before the deadline for submission. The same evaluation scores that are F1-macro for Subtasks 1 and 2, CoNLL-2012<sup>8</sup> for Subtask 3, and CoNLL-2000<sup>9</sup> script for Subtask 4 were utilized.

and <https://github.com/emerging-welfare/case-2022-multilingual-event> for CASE 2022.

<sup>3</sup><https://github.com/bltlib/mot>

<sup>4</sup><https://codalab.lisn.upsaclay.fr/competitions/7438>, accessed on November 13, 2022.

<sup>5</sup><https://codalab.lisn.upsaclay.fr/competitions/7126>, accessed on November 13, 2022.

<sup>6</sup><https://competitions.codalab.org/competitions/31247>, which is not accessible due to change of the servers of Codalab.

<sup>7</sup><https://codalab.lisn.upsaclay.fr/competitions/7768>, accessed on November 13, 2022.

<sup>8</sup><https://github.com/LoicGrobol/scorch>, accessed on November 13, 2022.

<sup>9</sup><https://github.com/sighsmile/conlleval>, accessed on November 13, 2022.

## 4 Results

Eighteen teams were registered for the task and obtained the training and test data for both CASE 2022 and CASE 2021. Ten and seven teams submitted their results for CASE 2021 and CASE 2022 respectively. Seven papers were submitted as system description papers to the CASE 2022 workshop in total. The scores of the submissions are calculated on two different Codalab pages for CASE 2021<sup>10</sup> and CASE 2022<sup>11</sup>. The teams that have participated are ARC-NLP (Sahin et al., 2022), CamPros (Kumari et al., 2022), CEIA-NLP (Fernandes et al., 2022), ClassBases (Wiriyathamabhum, 2022), EventGraph (You et al., 2022), NSUT-NLP (Suri et al., 2022), SPARTA (Müller and Dafnos, 2022). We provide details of the results and submissions of the participating teams for each subtask in the following subsections.<sup>12</sup>

### 4.1 CASE 2022 Subtask 1

The results for CASE 2022 subtask 1 are provided in Table 1. ARC-NLP finetune an ensemble of transformer-based language models and use ensemble learning, varying training data for each target language. They also perform tests with automatic translation of both training and test sets. They achieve 1st place both in Turkish and Mandarin, 2nd place in Portuguese and 3rd to 5th place in other languages. CEIA-NLP finetune XLM-Roberta-base transformers model with all the training data to achieve 1st place in Portuguese, 3rd or 4th places in other languages. ClassBases achieve 1st place in Hindi test data finetuning XLM-Roberta-large model, 5th or 6th places in other languages.

CamPros finetune XLM-Roberta-base model with all training data, and NSUT-NLP finetune mBERT while augmenting the data by translating different languages into each other.

<sup>10</sup><https://codalab.lisn.upsaclay.fr/competitions/7126#results>, accessed on Nov 14, 2022.

<sup>11</sup><https://codalab.lisn.upsaclay.fr/competitions/7438#results>, accessed on Nov 14, 2022.

<sup>12</sup>The results and system descriptions from participants that did not submit a system description paper are provided as well. This shows the capacity of the state-of-the-art systems on our benchmark. These systems are provided with their codalab names that are colabhero, fine\_sunny\_day, gauravsingh, lapardnemihk9989, lizhuoqun2021\_iscas.

### 4.2 CASE 2021 Subtask 1

The extended results for CASE 2021 subtask 1 are provided in Table 2. The boldness indicates CASE 2022 entries. ClassBases finetune XLM-Roberta-large transformers model to perform 1st in Hindi and 2nd in Portuguese test data. They also achieve 5th and 6th places in Spanish and English respectively. Another team that submitted their model to CASE 2021 test data is ARC-NLP, taking 5th, 8th and 9th places in Portuguese, Spanish and English.

### 4.3 Subtask 2

The extended results for CASE 2021 subtask 2 are provided in Table 3. The boldness indicates CASE 2022 entries. ARC-NLP train an ensemble of transformers models using all training data to achieve 4th, 5th and 7th places in Spanish, English and Portuguese respectively. ClassBases finetune mLUKE-base for Portuguese and Spanish placing 5th in both, XLM-Roberta-large for English taking 8th place.<sup>13</sup>

### 4.4 Subtask 3

The extended results for CASE 2021 subtask 3 are provided in Table 4. The boldness indicates CASE 2022 entries. ARC-NLP achieve 1st place in both English and Spanish, 2nd place in Portuguese. They use an ensemble of English transformers models for English, Portuguese and Spanish test data. They train with only English data and translating Portuguese test data into English during model prediction. For Spanish test data, they train with English, translated Portuguese and translated Spanish, and test on translated Spanish data.

### 4.5 Subtask 4

The extended results for CASE 2021 subtask 4 are provided in Table 5. The boldness indicates CASE 2022 entries. SPARTA employ two methods. Both of these methods build on pretrained XLM-Roberta-large and use a data augmentation technique (sentence reordering). For English and Portuguese, they gather articles that contain protest events from outside sources and use them for further pretraining. For Spanish, they use an XLM-Roberta-large model that was further pretrained on CoNLL 2002 Spanish data. They take 1st place both in Portuguese and Spanish, 3rd place in English.

<sup>13</sup>CamPros do not describe their model for subtask 2.

| Team                       | English            | Portuguese         | Spanish            | Hindi              | Turkish            | Urdu               | Mandarin           |
|----------------------------|--------------------|--------------------|--------------------|--------------------|--------------------|--------------------|--------------------|
| ARC-NLP                    | 80.74 <sub>4</sub> | 79.85 <sub>2</sub> | 69.44 <sub>5</sub> | 80.08 <sub>4</sub> | 84.06 <sub>1</sub> | 77.99 <sub>3</sub> | 83.39 <sub>1</sub> |
| CEIA-NLP                   | 80.77 <sub>3</sub> | 80.07 <sub>1</sub> | 73.19 <sub>3</sub> | 78.17 <sub>6</sub> | 82.43 <sub>4</sub> | 77.65 <sub>4</sub> | 77.63 <sub>4</sub> |
| CamPros                    | 76.52 <sub>7</sub> | 77.11 <sub>6</sub> | 69.55 <sub>4</sub> | 80.49 <sub>2</sub> | 74.75 <sub>6</sub> | 73.77 <sub>6</sub> | 75.90 <sub>6</sub> |
| ClassBases                 | 78.50 <sub>6</sub> | 77.11 <sub>5</sub> | 69.25 <sub>6</sub> | 80.78 <sub>1</sub> | 78.57 <sub>5</sub> | 75.72 <sub>5</sub> | 77.16 <sub>5</sub> |
| NSUT-NLP                   | 80.62 <sub>5</sub> | 73.02 <sub>7</sub> | 64.45 <sub>7</sub> | 56.71 <sub>7</sub> | 67.02 <sub>7</sub> | 65.55 <sub>7</sub> | 75.45 <sub>7</sub> |
| <i>fine_sunny_day</i>      | 82.22 <sub>2</sub> | 79.05 <sub>4</sub> | 73.84 <sub>2</sub> | 80.11 <sub>3</sub> | 82.91 <sub>2</sub> | 79.71 <sub>1</sub> | 80.99 <sub>3</sub> |
| <i>lizhuoqun2021_iscas</i> | 82.49 <sub>1</sub> | 79.22 <sub>3</sub> | 74.96 <sub>1</sub> | 80.01 <sub>5</sub> | 82.89 <sub>3</sub> | 78.67 <sub>2</sub> | 83.06 <sub>2</sub> |

Table 1: The performance of the submissions in terms of F1-macro and their ranks as a subscript for each language and each team participating in CASE 2022 subtask 1.

| Team                       | English             | Hindi               | Portuguese          | Spanish             |
|----------------------------|---------------------|---------------------|---------------------|---------------------|
| ALEM                       | 80.82 <sub>10</sub> | N/A                 | 72.98 <sub>11</sub> | 46.47 <sub>13</sub> |
| AMU-EuraNova               | 53.46 <sub>15</sub> | 29.66 <sub>11</sub> | 46.47 <sub>14</sub> | 46.47 <sub>13</sub> |
| DAAI                       | 84.55 <sub>3</sub>  | 77.07 <sub>6</sub>  | 82.43 <sub>4</sub>  | 69.31 <sub>10</sub> |
| DaDeFrTi                   | 80.69 <sub>11</sub> | 78.77 <sub>3</sub>  | 77.22 <sub>10</sub> | 73.01 <sub>7</sub>  |
| FKIE_itf_2021              | 73.90 <sub>13</sub> | 54.24 <sub>10</sub> | 62.39 <sub>12</sub> | 68.20 <sub>11</sub> |
| HSAIR                      | 77.58 <sub>12</sub> | 59.55 <sub>9</sub>  | 81.21 <sub>7</sub>  | 69.84 <sub>9</sub>  |
| IBM MNLP IE                | 83.93 <sub>4</sub>  | 78.53 <sub>5</sub>  | 84.00 <sub>3</sub>  | 77.27 <sub>3</sub>  |
| SU-NLP                     | 81.75 <sub>8</sub>  | N/A                 | N/A                 | N/A                 |
| NoConflict                 | 51.94 <sub>16</sub> | N/A                 | N/A                 | N/A                 |
| jitin                      | 67.39 <sub>14</sub> | 70.49 <sub>8</sub>  | 52.23 <sub>13</sub> | 62.05 <sub>12</sub> |
| <b>ARC-NLP</b>             | 81.35 <sub>9</sub>  | N/A                 | 81.73 <sub>5</sub>  | 72.42 <sub>8</sub>  |
| <b>ClassBases</b>          | 82.30 <sub>6</sub>  | 80.78 <sub>1</sub>  | 85.39 <sub>2</sub>  | 73.48 <sub>5</sub>  |
| <b>colabhero</b>           | 82.34 <sub>5</sub>  | 74.21 <sub>7</sub>  | 81.73 <sub>5</sub>  | 73.27 <sub>6</sub>  |
| <b>fine_sunny_day</b>      | 85.00 <sub>2</sub>  | N/A                 | 80.74 <sub>8</sub>  | 82.45 <sub>1</sub>  |
| <b>gauravsingh</b>         | 82.28 <sub>7</sub>  | 78.60 <sub>4</sub>  | 79.41 <sub>9</sub>  | 73.86 <sub>4</sub>  |
| <b>lizhuoqun2021_iscas</b> | 85.12 <sub>1</sub>  | 80.01 <sub>2</sub>  | 85.87 <sub>1</sub>  | 81.19 <sub>2</sub>  |

Table 2: The performance of the submissions in terms of F1-macro and their ranks as a subscript for each language and each team participating in CASE 2021 subtask 1. Bold teams indicate CASE 2022 entries.

| Team                       | English             | Portuguese          | Spanish             |
|----------------------------|---------------------|---------------------|---------------------|
| ALEM                       | 79.67 <sub>9</sub>  | 42.79 <sub>15</sub> | 45.30 <sub>15</sub> |
| AMU-EuraNova               | 75.64 <sub>14</sub> | 81.61 <sub>11</sub> | 76.39 <sub>11</sub> |
| DaDeFrTi                   | 79.28 <sub>10</sub> | 86.62 <sub>6</sub>  | 85.17 <sub>6</sub>  |
| FKIE_itf_2021              | 64.96 <sub>16</sub> | 75.81 <sub>13</sub> | 70.49 <sub>14</sub> |
| HSAIR                      | 78.50 <sub>11</sub> | 85.06 <sub>8</sub>  | 83.25 <sub>8</sub>  |
| IBM MNLP IE                | 84.56 <sub>4</sub>  | 88.47 <sub>3</sub>  | 88.61 <sub>2</sub>  |
| IIIT                       | 82.91 <sub>7</sub>  | 79.51 <sub>12</sub> | 75.78 <sub>12</sub> |
| SU-NLP                     | 83.05 <sub>6</sub>  | N/A                 | N/A                 |
| NoConflict                 | 85.32 <sub>3</sub>  | 87.00 <sub>4</sub>  | 79.97 <sub>10</sub> |
| jiawei1998                 | 76.14 <sub>13</sub> | 84.67 <sub>9</sub>  | 83.05 <sub>9</sub>  |
| jitin                      | 66.96 <sub>15</sub> | 69.02 <sub>14</sub> | 72.94 <sub>13</sub> |
| <b>ARC-NLP</b>             | 83.77 <sub>5</sub>  | 86.53 <sub>7</sub>  | 87.20 <sub>4</sub>  |
| <b>CamPros</b>             | 77.94 <sub>12</sub> | 81.63 <sub>10</sub> | 83.69 <sub>7</sub>  |
| <b>ClassBases</b>          | 81.12 <sub>8</sub>  | 86.83 <sub>5</sub>  | 87.10 <sub>5</sub>  |
| <b>fine_sunny_day</b>      | 85.75 <sub>2</sub>  | 89.67 <sub>1</sub>  | 88.78 <sub>1</sub>  |
| <b>lizhuoqun2021_iscas</b> | 85.93 <sub>1</sub>  | 88.86 <sub>2</sub>  | 88.61 <sub>2</sub>  |

Table 3: The performance of the submissions in terms of F1-macro and their ranks as a subscript for each language and each team participating in subtask 2. Bold teams indicate CASE 2022 entries.

ARC-NLP finetune an ensemble of transformers models for each language. They use all training data for Portuguese and Spanish, and only English for English test data. They achieve 2nd place in all languages. EventGraph aim to solve event ex-

| Team                   | English            | Portuguese         | Spanish            |
|------------------------|--------------------|--------------------|--------------------|
| DAAI                   | 80.40 <sub>4</sub> | 90.23 <sub>6</sub> | 81.83 <sub>6</sub> |
| FKIE_itf_2021          | 77.05 <sub>7</sub> | 91.33 <sub>4</sub> | 82.52 <sub>4</sub> |
| Handshakes AI Research | 79.01 <sub>5</sub> | 90.61 <sub>5</sub> | 81.95 <sub>5</sub> |
| IBM MNLP IE            | 84.44 <sub>2</sub> | 92.84 <sub>3</sub> | 84.23 <sub>2</sub> |
| NUS-IDS                | 81.20 <sub>3</sub> | 93.03 <sub>1</sub> | 83.15 <sub>3</sub> |
| SU-NLP                 | 78.67 <sub>6</sub> | N/A                | N/A                |
| <b>ARC-NLP</b>         | 85.11 <sub>1</sub> | 93.00 <sub>2</sub> | 85.25 <sub>1</sub> |

Table 4: The performance of the submissions in terms of CoNLL-2012 average score Pradhan et al. (2014) and their ranks as a subscript for each language and each team participating in subtask 3. Bold teams indicate CASE 2022 entries.

traction as semantic graph parsing. They use a graph encoding method where the labels for triggers and arguments are represented as node labels, also splitting multiple triggers. They use the pre-trained XLM-Roberta-large as their encoder. They achieve 4th place both in English and Portuguese, 5th place in Spanish. ClassBases take 9th place in all languages finetuning XLM-Roberta-base transformers model.

| Team                    | Scores             |                    |                    |
|-------------------------|--------------------|--------------------|--------------------|
|                         | English            | Portuguese         | Spanish            |
| AMU-EuraNova            | 69.96 <sub>7</sub> | 61.87 <sub>8</sub> | 56.64 <sub>8</sub> |
| Handshakes AI Research  | 73.53 <sub>5</sub> | 68.15 <sub>6</sub> | 62.21 <sub>6</sub> |
| IBM MNLP IE             | 78.11 <sub>1</sub> | 73.24 <sub>3</sub> | 66.20 <sub>3</sub> |
| SU-NLP                  | 2.58 <sub>10</sub> | N/A                | N/A                |
| jitin                   | 66.43 <sub>8</sub> | 64.19 <sub>7</sub> | 58.35 <sub>7</sub> |
| <b>ARC-NLP</b>          | 77.83 <sub>2</sub> | 73.84 <sub>2</sub> | 67.99 <sub>2</sub> |
| <b>ClassBases</b>       | 46.88 <sub>9</sub> | 12.52 <sub>9</sub> | 37.09 <sub>9</sub> |
| <b>EventGraph</b>       | 74.76 <sub>4</sub> | 71.72 <sub>4</sub> | 64.48 <sub>5</sub> |
| <b>SPARTA</b>           | 76.60 <sub>3</sub> | 74.56 <sub>1</sub> | 69.86 <sub>1</sub> |
| <b>lapardnemihk9989</b> | 72.18 <sub>6</sub> | 70.98 <sub>5</sub> | 64.83 <sub>4</sub> |

Table 5: The performance of the submissions in terms of F1 score based on CoNLL-2003 (Tjong Kim Sang and De Meulder, 2003) and their ranks as a subscript for each language and each team participating in subtask 4. Bold teams indicate CASE 2022 entries.

## 5 Conclusion

The CASE 2022 extension consists of expanding the test data with more data in previously available languages, namely, English, Hindi, Portuguese, and Spanish, and adding new test data in Mandarin, Turkish, and Urdu for Sub-task 1, document classification. The training data from CASE 2021 in English, Portuguese and Spanish were utilized. Therefore, predicting document labels in Hindi, Mandarin, Turkish, and Urdu occurs in a zero-shot setting.

The CASE 2022 workshop accepts reports on systems developed for predicting test data of CASE 2021 as well. We observe that the best systems submitted by CASE 2022 participants achieve between 79.71 and 84.06 F1-macro for new languages in a zero-shot setting. The winning approaches are mainly ensembling models and merging data in multiple languages. The best two submissions on CASE 2021 data outperform submissions from last year for Subtask 1 and Subtask 2 in all languages. Only the following scenarios were not outperformed by new submissions on CASE 2021: Subtask 3 Portuguese & Subtask 4 English.

We aim at increasing number of languages and subtasks such as event coreference resolution (Hürriyetoğlu et al., 2022) and event type classification (Hürriyetoğlu et al., 2021) in the scope of following edition of this shared task.

## References

- Fırat Duruşan, Ali Hürriyetoğlu, Erdem Yörük, Osman Mutlu, Çağrı Yoltar, Burak Gürel, and Alvaro Comin. 2022. [Global contentious politics database \(glocon\) annotation manuals](#).
- Diogo Fernandes, Adalberto Junior, Gabriel da Mata Marques, Anderson da Silva Soares, and Arlindo Rodrigues Galvao Filho. 2022. CEIA-NLP at CASE 2022 Task 1: Protest News Detection for Portuguese. In *Proceedings of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2022)*, online. Association for Computational Linguistics (ACL).
- Ali Hürriyetoğlu, Osman Mutlu, Erdem Yörük, Farhana Ferdousi Liza, Ritesh Kumar, and Shyam Ratan. 2021. [Multilingual protest news detection - shared task 1, CASE 2021](#). In *Proceedings of the 4th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2021)*, pages 79–91, Online. Association for Computational Linguistics.
- Ali Hürriyetoğlu, Hristo Tanev, Vanni Zavarella, Reyhan Yeniterzi, Osman Mutlu, and Erdem Yörük. 2022. Challenges and applications of automated extraction of socio-political events from text (case 2022): Workshop and shared task report. In *Proceedings of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2022)*, online. Association for Computational Linguistics (ACL).
- Ali Hürriyetoğlu, Osman Mutlu, Fatih Beyhan, Fırat Duruşan, Ali Safaya, Reyhan Yeniterzi, and Erdem Yörük. 2022. [Event coreference resolution for contentious politics events](#).
- Ali Hürriyetoğlu, Erdem Yörük, Osman Mutlu, Fırat Duruşan, Çağrı Yoltar, Deniz Yüret, and Burak Gürel. 2021. [Cross-Context News Corpus for Protest Event-Related Knowledge Base Construction](#). *Data Intelligence*, 3(2):308–335.
- Ali Hürriyetoğlu, Vanni Zavarella, Hristo Tanev, Erdem Yörük, Ali Safaya, and Osman Mutlu. 2020. [Automated extraction of socio-political events from news \(AESPEN\): Workshop and shared task report](#). In *Proceedings of the Workshop on Automated Extraction of Socio-political Events from News 2020*, pages 1–6, Marseille, France. European Language Resources Association (ELRA).
- Neha Kumari, Mrinal Anand, Tushar Mohan, and Ponnurangam Kumaraguru. 2022. CamPros at CASE 2022 Task 1: Transformer-based Multilingual Protest News Detection. In *Proceedings of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2022)*, online. Association for Computational Linguistics (ACL).
- Osman Mutlu. 2022. [Utilizing coarse-grained data in low-data settings for event extraction](#).
- Arthur Müller and Andreas Dafnos. 2022. SPARTA at CASE 2021 Task 1: Evaluating Different Techniques to Improve Event Extraction. In *Proceedings of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2022)*, online. Association for Computational Linguistics (ACL).
- Chester Palen-Michel, June Kim, and Constantine Lignos. 2022. [Multilingual open text release 1: Public domain news in 44 languages](#). In *Proceedings of the Thirteenth Language Resources and Evaluation Conference*, pages 2080–2089, Marseille, France. European Language Resources Association.
- Sameer Pradhan, Xiaoqiang Luo, Marta Recasens, Eduard Hovy, Vincent Ng, and Michael Strube. 2014. [Scoring coreference partitions of predicted mentions: A reference implementation](#). In *Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 30–35, Baltimore, Maryland. Association for Computational Linguistics.

- Umitcan Sahin, Oguzhan Ozcelik, Izzet Emre Kucukkaya, and Cagri Toraman. 2022. ARC-NLP at CASE 2022 Task 1: Ensemble Learning for Multilingual Protest Event Detection. In *Proceedings of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2022)*, online. Association for Computational Linguistics (ACL).
- Manan Suri, Krish Chopra, and Adwita Arora. 2022. NSUT-NLP at CASE 2022 Task 1: Multilingual Protest Event Detection using Transformer-based Models. In *Proceedings of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2022)*, online. Association for Computational Linguistics (ACL).
- Fiona Anting Tan, Ali Hürriyetoğlu, Tommaso Caselli, Nelleke Oostdijk, Hansi Hettiarachchi, Tadashi Nomoto, Onur Uca, and Farhana Ferdousi Liza. 2022a. Event causality identification with causal news corpus - shared task 3, CASE 2022. In *Proceedings of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2022)*, Online. Association for Computational Linguistics.
- Fiona Anting Tan, Ali Hürriyetoğlu, Tommaso Caselli, Nelleke Oostdijk, Tadashi Nomoto, Hansi Hettiarachchi, Iqra Ameer, Onur Uca, Farhana Ferdousi Liza, and Tiancheng Hu. 2022b. [The causal news corpus: Annotating causal relations in event sentences from news](#). In *Proceedings of the Language Resources and Evaluation Conference*, pages 2298–2310, Marseille, France. European Language Resources Association.
- Erik F. Tjong Kim Sang and Fien De Meulder. 2003. [Introduction to the conll-2003 shared task: Language-independent named entity recognition](#). In *Proceedings of the Seventh Conference on Natural Language Learning at HLT-NAACL 2003 - Volume 4, CONLL '03*, page 142–147, USA. Association for Computational Linguistics.
- Peratham Wiriathamabhum. 2022. ClassBases at the CASE-2022 Multilingual Protest Event Detection Task: Multilingual Protest News Detection and Automatically Replicating Manually Created Event Datasets. In *Proceedings of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2022)*, online. Association for Computational Linguistics (ACL).
- Huiling You, David Samuel, Samia Touileb, and Lilja Øvrelid. 2022. EventGraph at CASE 2021 Task 1: A General Graph-based Approach to Protest Event Extraction. In *Proceedings of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2022)*, online. Association for Computational Linguistics (ACL).
- Erdem Yörük, Ali Hürriyetoğlu, Fırat Duruşan, and Çağrı Yoltar. 2021. [Random sampling in corpus design: Cross-context generalizability in automated multicountry protest event collection](#). *American Behavioral Scientist*, 0(0):00027642211021630.
- Vanni Zavarella, Hristo Tanev, Ali Hürriyetoğlu, Peratham Wiriathamabhum, and Bertrand De Longueville. 2022. Tracking COVID-19 protest events in the United States. Shared Task 2: Event Database Replication, CASE 2022. In *Proceedings of the 5th Workshop on Challenges and Applications of Automated Extraction of Socio-political Events from Text (CASE 2022)*, online. Association for Computational Linguistics (ACL).