

Moscow, June 15–18, 2022

RuNNE-2022 Shared Task: Recognizing Nested Named Entities

**Ekaterina Artemova^{1,2}, Maxim Zmeev¹, Natalia Loukachevitch³, Igor Rozhkov³,
Tatiana Batura^{3,4}, Vladimir Ivanov⁵, Elena Tutubalina^{1,6,7}**

¹HSE University ²Huawei Noah's Ark lab,

³Lomonosov Moscow State University, ⁴Novosibirsk State University

⁵Innopolis University, ⁶Kazan Federal University, ⁷Sber AI, Russia

Abstract

The RuNNE Shared Task approaches the problem of nested named entity recognition. The annotation schema is designed in such a way, that an entity may partially overlap or even be nested into another entity. This way, the named entity "The Yermolova Theatre" of type ORGANIZATION houses another entity "Yermolova" of type PERSON. We adopt the Russian NEREL dataset (Loukachevitch et al., 2021) for the RuNNE Shared Task. NEREL comprises news texts written in the Russian language and collected from the Wikinews portal. The annotation schema includes 29 entity types. The nestedness of named entities in NEREL reaches up to six levels. The RuNNE Shared Task explores two setups. (i) In the general setup all entities occur more or less with the same frequency. (ii) In the few-shot setup the majority of entity types occur often in the training set. However, some of the entity types are have lower frequency, being thus challenging to recognize. In the test set the frequency of all entity types is even.

This paper reports on the results of the RuNNE Shared Task. Overall the shared task has received 156 submissions from nine teams. Half of the submissions outperform a straightforward BERT-based baseline in both setups. This paper overviews the shared task setup and discusses the submitted systems, discovering meaning insights for the problem of nested NER. The links to the evaluation platform and the data from the shared task are available in our github repository.

Keywords: nested named entity recognition, few-shot setup, shared task

DOI: 10.28995/2075-7182-2022-21-XX-XX

Соревнование RuNNE-2022: извлечение вложенных именованных сущностей

**Артемova Е. Л.^{1,2}, Змеев М. В.¹, Лукашевич Н. А.³, Рожков И.С.³,
Батура Т. В.^{3,4}, Иванов В. В.⁵, Тутубалина Е. В.^{2,6,7}**

¹Национальный исследовательский университет Высшая школа экономики

²Huawei Noah's Ark lab,

³Московский Государственный Университет им. М.В. Ломоносова

⁴Новосибирский государственный университет

⁵Иннополис, ⁶Казанский федеральный университет, ⁷Sber AI, Россия

Аннотация

Извлечение именованных сущностей – одна из самых востребованных на практике задач извлечения информации – предполагает поиск в тексте упоминаний имен, организаций, топонимов и других сущностей. Соревнование RuNNE посвящено задаче извлечения вложенных именованных сущностей. Разметка данных допускает следующие случаи: внутри одной именованной сущности находится другая именованная сущность. Так, например в сущность класса Organization "Московский драматический театр имени М. Н. Ермоловой" вложена сущность типа Person – "М. Н. Ермоловой". Соревнование проводится на материале корпуса NEREL (Loukachevitch et al., 2021), собранного из новостных текстов WikiNews на русском языке. В корпусе NEREL представлено 29 классов различных сущностей, а глубина вложенности сущностей достигает 6 уровней разметки. В рамках соревнования RuNNE мы предлагаем участникам рассмотреть few shot постановку задачи. Задача предполагает извлечение вложенных именованных сущностей, В обучающем множестве большая часть типов именованных сущностей встречается достаточно часто, а некоторое

количество специально отобранных типов – встречается всего несколько раз. В тестовом множестве все типы сущностей представлены одинаково. В данной статье мы описываем соревнование RuNNE, подводим его итоги и проводим сравнение решений, полученных от участников.

Ключевые слова: извлечение вложенных именованных сущностей, соревнование

1 Introduction

Named entity recognition is one of the most popular tasks in natural language processing. It involves labeling mentions of personal names, organizations, toponyms and other entities in the text. In early works, only “flat” named entities, in which internal named entities within a longer entity are not allowed, were annotated and extracted from texts. However, the assumption that all entities are flat oversimplifies the task and leads to annotation collision. For example, consider such entities as the “Ministry of Education of the Russian Federation” or the “State Duma of the Federal Assembly of the Russian Federation”. In this case, there are two options of flat labeling: a) splitting into minimum spans (“Ministry of Education” and “Russian Federation”), or b) extracting maximal spans. The latter case would miss the mention of the *Russian Federation* as a separate entity. Both labeling approaches lead to losing meaningful information pieces, which, in turn, can be useful for further processing, e.g. relation extraction.

In recent years, recognition of nested named entities has received special attention. In case of nested NEs a longer named entity can contain internal named entities. In this research direction, multiple new datasets were published (Plank et al., 2020; Ringland et al., 2019; Ruokolainen et al., 2019). The best performing methods for nested NER utilize specialized neural network architectures (Straková et al., 2019; Yu et al., 2020; Jue et al., 2020). A recent dataset for nested NER in Russian, NEREL, is annotated with 29 entity types (Loukachevitch et al., 2021).

The RuNNE shared task aims to popularize the nested NER task. To this end, we setup two sub-tasks based on the NEREL dataset: (i) the general nested NER evaluation; entities of more or less the same frequency are used (ii) the few-shot nested NER evaluation, where low-frequency entities are used. This paper (i) introduces the NEREL corpus (Section 3), (ii) reports the results of the RuNNE Shared Task (Section 5.1) and (iii) compares submitted systems (Section 5.2).

2 Related NER Datasets and Shared Tasks

Several datasets for named entity recognition in the Russian language are available, e.g. the Gareev’s dataset (Gareev et al., 2013), Persons-1000 and Collection3 (Mozharova and Loukachevitch, 2016; Vlasova et al., 2014), FactRuEval (Starostin et al., 2016), the Russian subset of the BSNLP dataset (Piskorski et al., 2019), and RURED (Gordeev et al., 2020). In particular, FactRuEval and BSNLP datasets were used for corresponding shared tasks. FactRuEval published in 2016 has been the only Russian dataset annotated with nested named entities until recently. For example, a person object in FactRuEval could include name, surname, patronymic, and nickname spans. The dataset is rather small with approximately 250 news texts annotated with 11,700 spans and 7,700 objects. This limits its effectiveness in building large neural network models. Additional descriptions and statistics of Russian NER datasets can be found in (Loukachevitch et al., 2021).

In recent years, the NLP community has held multiple shared tasks on named entity recognition, tackling different aspects of the task:

- **Multilinguality** The series of SlavNER Shared Tasks (Piskorski et al., 2021) focuses on six Slavic languages; the MultiCoNER Shared Task (Malmasi et al., 2022) features 11 world languages. XTREME (Hu et al., 2020) and XGLUE (Liang et al., 2020) data suites are collections of tasks and corresponding datasets for evaluation of zero-shot transfer capabilities of large multilingual models from English to other languages. XGLUE adopts CoNLL-2002 NER (Tjong Kim Sang, 2002) and CoNLL-2003 NER (Tjong Kim Sang and De Meulder, 2003) datasets covering four languages (English, German, Spanish, and Dutch) and four entity types (Person, Location, Organization, and Miscellaneous). XTREME uses Wikiann dataset (Pan et al., 2017), where Person, Location and Organization entities were automatically annotated in Wikipedia pages in 40 languages.
- **Complexity** The MultiCoNER Shared Task (Malmasi et al., 2022) was devoted to extraction of

semantically ambiguous and complex entities, such as movie and book titles in short and low-context settings.

- **Applied domains** The WNUT initiative aims at developing NLP tools for noisy user-generated text. To this end, they run multiple shared tasks on NER and relation extraction from social media (Nguyen et al., 2020; Chen et al., 2020). PharmaCoNER is aimed at clinical and medical NER (Gonzalez-Agirre et al., 2019).
- **Languages, other than English** The Dialogue evaluation campaign has supported two shared tasks for NER in Russian. FactRuEval dataset comprised news texts, Wikipedia pages and posts on social media, annotated according to a two-level schema. The RuReBUS Shared task (Ivanin et al., 2020) approached joint NER and relation extraction in Russian business-related documents.

3 NEREL Dataset

The NEREL collection is developed for studying three levels of information extraction methods including named entity recognition, relation extraction and entity linking (Loukachevitch et al., 2021). Currently, NEREL is the largest Russian dataset annotated with entities and relations compared to the existing Russian datasets. NEREL comprises 29 named entity types and 49 relation types. At the time of writing, the dataset contains 56K named entities and 39K relations annotated in 900+ person-oriented news articles. NEREL is annotated with relations at three levels: within nested named entities, within and across sentences. Entity linking annotations leverage nested named entities, and each nested named entity can be linked to a separate Wikidata entity.

The NEREL corpus consists mainly of Russian Wikinews articles having size of 1–5 Kb, as such medium-sized texts are more convenient for annotation. The BRAT tool (Stenetorp et al., 2012) was used for annotation. Three levels of annotation – named entities, relations, Wikidata links, – were performed as independent subsequent passes. NEREL dataset is publicly available. It is subdivided into train, dev, and test parts, which were used in the RuNNE evaluation. The training part was specially reduced for few-shot NER evaluation. Evaluation was carried out on the NEREL’s dev and test sets. Table 1 contains examples of nested named entities in the NEREL dataset. Entity type statistics of the dataset used for the RuNNE competition are shown in Table 2. During the evaluation phase the NEREL github repository was temporarily closed.

4 RuNNE Task

The RuNNE competition offered two tasks: general nested named entity recognition and few-shot setting. In the training set, most named entity types occur quite often, and a certain number of specially selected types occur only a few times. In the test set, all entity types are represented equally.

As a quality metric in the RuNNE competition, macro averaging of the F1-measure is used in two versions: by classes of known entities (general formulation of the problem of extracting nested named entities) and by classes of new named entities (few-shot formulation). Three entity types annotated in the NEREL dataset were selected for the few-shot setting: DISEASE, WORK_OF_ART, PENALTY.

As a baseline we employed a RuBERT model (Kuratov and Arkhipov, 2019) with a fully connected output layer, predicting each of the classes in the data, from where only the flat most internal entities were taken. For named entity encoding the IOBES (Inside-Out-Begin-End-Single) scheme was used.

5 Evaluation

5.1 Participants’ submissions

We have received 156 submissions from nine teams. Table 3 presents with the final scores of the submitted systems. Four out of nine systems outperformed the baseline. Below, we give an overview of these approaches.

Team **Ksmith (Pullenti)** achieved the best results by using a rule-based approach. Customized rules were written for every entity type (Kozerenko et al., 2018). This team made the highest number of entries to the leaderboard.

Entity	Example	Annotation
General entities		
PROFESSION	<i>governor of California</i>	[governor of [California] _{STATE_OR_PROVINCE}] _{PROFESSION}
	<i>head of Gazprom</i>	[head of [Gazprom] _{ORGANIZATION}] _{PROFESSION}
ORGANIZATION	<i>physics department of Lomonosov Moscow State University</i>	[physics department of [[Lomonosov] _{PERSON} [Moscow] _{CITY} State University] _{ORG}] _{ORG}
	<i>Russian Government</i>	[[Russian] _{COUNTRY} government] _{ORG}
NATIONALITY	<i>citizen of Russia</i>	[citizen of [Russia] _{COUNTRY}] _{NATIONALITY}
	<i>Russians</i>	[Russians] _{NATIONALITY}
	<i>Russian writer</i>	[Russian] _{NATIONALITY} [writer] _{PROFESSION}
LAW	<i>Yarovaya law</i>	[[Yarovaya] _{PERSON} law] _{LAW}
	<i>article 84 of the Constitution of Kyrgyzstan</i>	[article [84] _{ORDINAL} of the [[Constitution] _{LAW} COUNTRY] _{LAW}] _{LAW}
CRIME	<i>complicity in murder of VGTRK journalists</i>	[[complicity in murder] _{CRIME} of [[VGTRK] _{ORG} [journalists] _{PROFESSION}] _{CRIME}
	<i>armed attack on passers-by attempt to oust Erdogan</i>	[[armed attack] _{CRIME} on passers-by] _{CRIME} [attempt to oust [Erdogan] _{PERSON}] _{CRIME}
PRODUCT	<i>Boeing-737 MAX</i>	[[Boeing] _{ORG} -[737] _{NUMBER} MAX] _{PRODUCT}
	<i>Apple Watch</i>	[Apple] _{ORG} Watch] _{PRODUCT}
AWARD	<i>Merit for the Fatherland Order</i>	[Merit for the Fatherland Order] _{AWARD}
	<i>gold of the Olympic Games</i>	[gold of the [Olympic Games] _{EVENT}] _{AWARD}
	<i>champion of the Olympic Games</i>	[champion of the [Olympic Games] _{EVENT}] _{AWARD}
	<i>Miss Russia-2017</i>	[Miss [Russia] _{COUNTRY} -[2017] _{DATE}] _{AWARD}
	<i>gold medal</i>	[gold medal] _{AWARD}
EVENT	<i>40th Moscow International Film Festival</i>	[[40th] _{ORDINAL} [[Moscow] _{CITY} International Film Festival] _{EVENT}] _{EVENT}
	<i>UEFA champions league</i>	[UEFA] _{ORG} champions league] _{EVENT}
	<i>Sochi-2014</i>	[[Sochi] _{CITY} -[2014] _{DATE}] _{EVENT}
Few-shot entities		
PENALTY	<i>\$129 million fine imprisonment for 3 years</i>	[[[\$129 million] _{MONEY} [fine] _{PENALTY}] _{PENALTY} [[imprisonment] _{PENALTY} for [3 years] _{DATE}] _{PENALTY}
DISEASE	<i>died from Covid thyroid cancer</i>	[[died] _{EVENT} from [Covid] _{DISEASE}] _{EVENT} [thyroid [cancer] _{DISEASE}] _{DISEASE}
WORK_OF_ART	<i>the host of TV show “Vzglyad”</i>	[the host of TV show [“Vzglyad”] _{WORK_OF_ART}] _{PROFESSION}

Table 1: Examples of annotating nested named entities of different types

Team **Fulstock (MSU RCC)** applied the Machine reading comprehension model (MRC) (Li et al., 2020). The MRC model treats NER as a question-answering task. Entity types are translated into Russian. Next, their definitions are gathered from dictionaries in Russian and used as questions. The MRC model comprises three binary classifiers over the output of the last hidden layer from RuBERT. The first classifier determines the starting position of a named entity. The second classifier decides about the end position of a named entity (perhaps another) of the same class. The third classifiers classifies, whether the start-end pairs are a single entity. These classifiers are trained for each of the entity type.

Team **Abrosimov_Kirill (Saldon)** used the span-based Sodner model (Li et al., 2021) that can recognize both overlapped and discontinuous entities jointly. The model is based on the graph convolutional network architecture. It includes two steps. First, entity fragments are recognized by traversing over all possible text spans. Second, relation classification is implemented to judge to detect overlapping or succession. The team used a pre-trained RuBERT model as contextualized encoder and the Natasha syntax parser.¹ Additionally, selected sentences for few-shot entity types were manually annotated.

¹<https://github.com/natasha/slovnet#syntax>

Entity	train	dev	test	Entity	train	dev	test
PROFESSION	4593	860	854	IDEOLOGY	300	36	43
PERSON	4518	947	961	LOCATION	272	64	62
ORGANIZATION	4059	616	675	PRODUCT	238	30	53
EVENT	2879	707	690	CRIME	181	64	35
COUNTRY	2521	355	456	MONEY	171	29	43
DATE	2276	527	523	TIME	154	29	47
CITY	1102	208	239	DISTRICT	98	18	25
NUMBER	1026	186	230	RELIGION	94	9	24
ORDINAL	565	102	107	PERCENT	82	9	7
AGE	554	137	138	LANGUAGE	43	7	8
NATIONALITY	394	59	66	DISEASE	32	117	57
LAW	392	77	62	PENALTY	32	57	18
FACILITY	371	84	63	WORK_OF_ART	30	104	93
STATE_OR_PROVINCE	343	99	112	FAMILY	17	7	14
AWARD	328	43	121	Total	27665	5587	5826

Table 2: NEREL dataset statistics for the RuNNE competition.

User	Team	# of runs	F1 _{full_set}	F1 _{few-shot}	System Summary
<i>Baseline</i>		-	0.674	0.447	RuBERT
<i>Participating Teams</i>					
ksmith	Pullenti	44	0.811	0.710	Rule-based
abrosimov_kirill	Saldon	20	<u>0.741</u>	<u>0.644</u>	Sodner model, labelling
fulstock	MSU-RCC	7	<u>0.749</u>	<u>0.604</u>	MRC model
svetlan		23	0.607	<u>0.572</u>	n/a
LIORI		6	0.653	0.433	n/a
botbot		8	0.460	0.414	n/a
bond005	SibNN	20	<u>0.743</u>	0.404	Siamese network, Viterbi alg.
Stud2022		24	0.477	0.395	n/a
mojesty		2	0.619	0.172	n/a

Table 3: Macro-averaged F-scores on the official test sets. The best results are in bold. Results above the baseline are underlined.

Team **Bond005 (SibNN)** fine-tuned a pre-trained BERT-based transformer two times. First, the contextualized encoder is fine-tuned as a Siamese neural network using the supervised contrastive learning loss. This step brings together embeddings of named entities of the same class and moves apart named entities of different classes. The second step fine-tunes the contextualized encoder for sequence labeling using a task-specific loss function. Finally, the team applied the Viterbi algorithm to smooth prediction probabilities. The team defined the transition probabilities manually.

5.2 Results

The results (Table 3) indicate that:

1. the conventional rule-based approach is capable of outperforming recent neural models. However, the rule-based approach requires careful configuration, leading to more attempts to submit to the leaderboard;

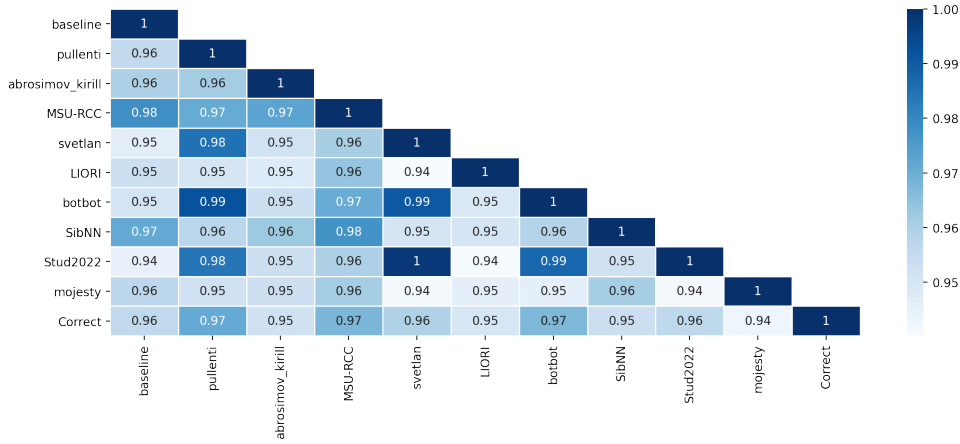


Figure 1: Agreement (Cohen’s κ) among participating systems in annotating entity types. The “Correct” row corresponds to the test dataset annotations.

- the MRC approach benefits from additional textual descriptions and outperforms other learnable approaches;
- the few-shot setup benefits from additional manual labelling and hand-crafted rules.

5.3 Detailed comparison of systems

Entity Type	baseln.	pullenti	abros.	MSU.	svetlan	LIORI	botbot	SibNN	Stud2022	mojesty
ORGANIZATION	23	21	28	21	20	18	17	29	14	24
COUNTRY	8	36	19	12	27	21	15	12	19	34
WORK_OF_ART	22	3	20	16	1	16	2	24	1	23
STATE_OR_PROV.	11	6	21	9	5	18	4	18	4	29
EVENT	10	9	17	8	20	7	5	13	15	8
LOCATION	9	8	12	8	6	18	8	5	6	10
FACILITY	6	9	11	11	5	7	2	10	3	10
NATIONALITY	8	2	6	7	8	10	2	9	5	7
PRODUCT	9	1	5	6	-	5	1	17	-	7
CITY	6	5	7	5	6	9	3	5	3	11
PERSON	12	2	11	4	5	7	1	6	2	8
FAMILY	9	-	10	6	2	7	1	8	1	7
DISTRICT	7	2	4	3	1	10	1	6	1	5
AWARD	8	1	9	3	2	3	1	4	1	5
DISEASE	2	3	1	3	5	3	2	6	4	7
PROFESSION	5	1	5	5	1	3	-	5	1	6
Total span matches	4401	4732	4948	4853	3511	4283	2363	4618	2361	4595

Table 4: Number of test instances with incorrectly recognized entity type. Top-16 entities cover approximately 85% of all mismatches. The last row shows the number of correctly recognized spans for each system. Symbol ‘-’ means that all system’s predictions were correct. All calculations were performed for the *full_set* test set.

We have compared solutions of participating systems in two aspects. First, we compare errors in entity type annotations given a span of a named entity was correctly detected. Second, we calculate pairwise agreements of systems’ answers, again only for the annotations that have exact match of spans. Both evaluations are done on the test set. In Table 4 we provide top entities for which systems give wrong entity type prediction. We found that around a half of errors come from six entity types: ORGANIZATION, COUNTRY, WORK_OF_ART, STATE_OR_PROVINCE, EVENT, and LOCATION. The most common mismatches are the following (correct entity is on the left of the ‘ $\rightarrow$ ’ sign):

- ORGANIZATION → FACILITY / PRODUCT / CITY,
- COUNTRY → NATIONALITY / ORGANIZATION,
- WORK_OF_ART → ORGANIZATION / LAW / EVENT,
- STATE_OR_PROVINCE → CITY / COUNTRY / LOCATION,
- EVENT → ORGANIZATION / CRIME / DISEASE,
- LOCATION → STATE_OR_PROVINCE / CITY/ COUNTRY.

To assess pairwise agreements between systems, we build lists of spans that have exactly matching boundaries, but may have different type annotations. Then we calculate Cohen’s κ for each pair of systems (Fig. 1). Depending on a pair of systems, the number of matching spans varies between 1,294 and 5,826. Due to this difference, values in the Figure 1 slightly deviate from the ranking of the systems. Three teams, showing mediocre performance (**svetlan**, **botbot** and **Stud2022**) have highest agreement with each other. Also these three teams have the highest agreement with the **pullenti** team. This indicates that these systems learn named entity patterns, which can be described by rules, too. Other systems may gain more generalization abilities and thus are less correlated with the rule-based system. Finally, the main source of errors for all approaches may be attributed to ambiguous or complex cases.

6 Conclusion

In this paper we described the RuNNE Shared Task, aimed at extracting nested named entities from texts in Russian. We used NEREL (Loukachevitch et al., 2021) as a source dataset. RuNNE comprises two sub-tasks: (i) the general nested NER evaluation and (ii) the few-shot nested NER evaluation. Comparing submissions of the RuNNE participants, we found that the classic rule-based approach is capable to outperform recent neural models. However, it requires a special framework for writing rules and iterative improvement of rules, leading to more attempts in leaderboard submits. In both subtasks, the best model among machine learning approaches applied without additional manual data labeling was Machine Reading Comprehension model (Li et al., 2020).

Acknowledgments

The project is supported by the Russian Science Foundation, grant # 20-11-20166. The experiments were partially carried out on computational resources of HPC facilities at HSE University (Kostenetskiy et al., 2021) and the shared research facilities of HPC computing resources at Lomonosov Moscow State University. Ekaterina Artemova was supported by the framework of the HSE University Basic Research Program.

References

- Chacha Chen, Chieh-Yang Huang, Yaqi Hou, Yang Shi, Enyan Dai, and Jiaqi Wang. 2020. TEST_POSITIVE at W-NUT 2020 shared task-3: Cross-task modeling. // *Proceedings of the Sixth Workshop on Noisy User-generated Text (W-NUT 2020)*, P 499–504, Online, November. Association for Computational Linguistics.
- Rinat Gareev, Maksim Tkachenko, Valery Solovyev, Andrey Simanovsky, and Vladimir Ivanov. 2013. Introducing Baselines for Russian Named Entity Recognition. // *International Conference on Intelligent Text Processing and Computational Linguistics*, P 329–342.
- Aitor Gonzalez-Agirre, Montserrat Marimon, Ander Intxaurre, Obdulia Rabal, Marta Villegas, and Martin Krallinger. 2019. PharmaCoNER: Pharmacological substances, compounds and proteins named entity recognition track. // *Proceedings of The 5th Workshop on BioNLP Open Shared Tasks*, P 1–10, Hong Kong, China, November. Association for Computational Linguistics.
- Denis Gordeev, Adis Davletov, Alexey Rey, Galiya Akzhigitova, and Georgiy Geymbukh. 2020. Relation extraction dataset for the Russian language. // *Computational Linguistics and Intellectual Technologies*.
- Junjie Hu, Sebastian Ruder, Aditya Siddhant, Graham Neubig, Orhan Firat, and Melvin Johnson. 2020. XTREME: A massively multilingual multi-task benchmark for evaluating cross-lingual generalisation. // *International Conference on Machine Learning*, P 4411–4421.

- Vitaly Ivanin, Ekaterina Artemova, Tatiana Batura, Vladimir Ivanov, Veronika Sarkisyan, Elena Tutubalina, and Ivab Smurov. 2020. Rurebus-2020 shared task: Russian relation extraction for business. // *Computational Linguistics and Intellectual Technologies*, P 416–431.
- Wang Jue, Lidan Shou, Ke Chen, and Gang Chen. 2020. Pyramid: A Layered Model for Nested Named Entity Recognition. // *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, P 5918–5928.
- Pavle Kostenetskiy, Roman Chulkevich, and Viacheslav Kozyrev. 2021. HPC resources of the Higher School of Economics. // *Journal of Physics: Conference Series*, volume 1740, P 012050. IOP Publishing.
- Elena Kozerenko, Konstantin Kuznetsov, and Dmitrii Romanov. 2018. Semantic Processing of Unstructured Textual Data Based on the Linguistic Processor PullEnti. *Informatika i Ee Primeneniya [Informatics and its Applications]*, 12(3):91–98.
- Yuri Kuratov and Mikhail Arkhipov. 2019. Adaptation of deep bidirectional multilingual transformers for Russian language. // *Computational Linguistics and Intellectual Technologies*, P 333–339.
- Xiaoya Li, Jingrong Feng, Yuxian Meng, Qinghong Han, Fei Wu, and Jiwei Li. 2020. A Unified MRC Framework for Named Entity Recognition. // *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, P 5849–5859.
- Fei Li, ZhiChao Lin, Meishan Zhang, and Donghong Ji. 2021. A Span-Based Model for Joint Overlapped and Discontinuous Named Entity Recognition. // *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)*, P 4814–4828.
- Yaobo Liang, Nan Duan, Yeyun Gong, Ning Wu, Fenfei Guo, Weizhen Qi, Ming Gong, Linjun Shou, Daxin Jiang, Guihong Cao, Xiaodong Fan, Ruofei Zhang, Rahul Agrawal, Edward Cui, Sining Wei, Taroon Bharti, Ying Qiao, Jiun-Hung Chen, Winnie Wu, Shuguang Liu, Fan Yang, Daniel Campos, Rangan Majumder, and Ming Zhou. 2020. XGLUE: A new benchmark dataset for cross-lingual pre-training, understanding and generation. // *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, P 6008–6018, Online, November. Association for Computational Linguistics.
- Natalia Loukachevitch, Ekaterina Artemova, Tatiana Batura, Pavel Braslavski, Ilia Denisov, Vladimir Ivanov, Suresh Manandhar, Alexander Pugachev, and Elena Tutubalina. 2021. NEREL: A Russian dataset with nested named entities, relations and events. // *Proceedings of the International Conference on Recent Advances in Natural Language Processing (RANLP 2021)*, P 876–885, Held Online, September. INCOMA Ltd.
- Shervin Malmasi, Anjie Fang, Besnik Fetahu, Sudipta Kar, and Oleg Rokhlenko. 2022. SemEval-2022 Task 11: Multilingual Complex Named Entity Recognition (MultiCoNER). // *Proceedings of the 16th International Workshop on Semantic Evaluation (SemEval-2022)*. Association for Computational Linguistics.
- Valerie Mozharova and Natalia Loukachevitch. 2016. Two-stage Approach in Russian Named Entity Recognition. // *International FRUCT Conference on Intelligence, Social Media and Web (ISMW FRUCT)*, P 1–6.
- Dat Quoc Nguyen, Thanh Vu, Afshin Rahimi, Mai Hoang Dao, Linh The Nguyen, and Long Doan. 2020. WNUT-2020 task 2: Identification of informative COVID-19 English tweets. // *Proceedings of the Sixth Workshop on Noisy User-generated Text (W-NUT 2020)*, P 314–318, Online, November. Association for Computational Linguistics.
- Xiaoman Pan, Boliang Zhang, Jonathan May, Joel Nothman, Kevin Knight, and Heng Ji. 2017. Cross-lingual name tagging and linking for 282 languages. // *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, P 1946–1958, Vancouver, Canada, July. Association for Computational Linguistics.
- Jakub Piskorski, Laska Laskova, Michał Marcińczuk, Lidia Pivovarov, Pavel Přibán, Josef Steinberger, and Roman Yangarber. 2019. The Second Cross-lingual Challenge on Recognition, Normalization, Classification, and Linking of Named Entities across Slavic Languages. // *Proceedings of the 7th Workshop on Balto-Slavic Natural Language Processing*, P 63–74.
- Jakub Piskorski, Bogdan Babych, Zara Kancheva, Olga Kanishcheva, Maria Lebedeva, Michał Marcińczuk, Preslav Nakov, Petya Osenova, Lidia Pivovarov, Senja Pollak, Pavel Přibán, Ivaylo Radev, Marko Robnik-Sikonja, Vasyl Starko, Josef Steinberger, and Roman Yangarber. 2021. Slav-NER: the 3rd cross-lingual challenge on recognition, normalization, classification, and linking of named entities across Slavic languages. // *Proceedings of the 8th Workshop on Balto-Slavic Natural Language Processing*, P 122–133, Kiyv, Ukraine, April. Association for Computational Linguistics.

- Barbara Plank, Kristian Nørgaard Jensen, and Rob van der Goot. 2020. DAN+: Danish Nested Named Entities and Lexical Normalization. // *Proceedings of the 28th International Conference on Computational Linguistics*, P 6649–6662.
- Nicky Ringland, Xiang Dai, Ben Hachey, Sarvnaz Karimi, Cecile Paris, and James R Curran. 2019. NNE: A Dataset for Nested Named Entity Recognition in English Newswire. // *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, P 5176–5181.
- Teemu Ruokolainen, Pekka Kauppinen, Miikka Silfverberg, and Krister Lindén. 2019. A Finnish News Corpus for Named Entity Recognition. *Language Resources and Evaluation*, P 1–26.
- Anatoly Starostin, Victor Bocharov, Svetlana Alexeeva, Anastasiya Bodrova, Alexander Chuchunkov, Stanislav Dzhumaev, Irina Efimenko, Dmitry Granovsky, Viktor Khoroshevsky, Irina Krylova, et al. 2016. FactRuEval 2016: evaluation of named entity recognition and fact extraction systems for Russian. *Computational Linguistics and Intellectual Technologies: Proceedings of the International Conference “Dialog”*, P 702–720.
- Pontus Stenetorp, Sampo Pyysalo, Goran Topić, Tomoko Ohta, Sophia Ananiadou, and Jun’ichi Tsujii. 2012. BRAT: a web-based tool for nlp-assisted text annotation. // *Proceedings of the Demonstrations at the 13th Conference of the European Chapter of the Association for Computational Linguistics*, P 102–107.
- Jana Straková, Milan Straka, and Jan Hajic. 2019. Neural Architectures for Nested NER through Linearization. // *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, P 5326–5331.
- Erik F. Tjong Kim Sang and Fien De Meulder. 2003. "introduction to the CoNLL-2003 shared task: Language-independent named entity recognition". // *Proceedings of the Seventh Conference on Natural Language Learning at HLT-NAACL 2003*, P 142–147.
- Erik F. Tjong Kim Sang. 2002. "introduction to the CoNLL-2002 shared task: Language-independent named entity recognition". // *COLING-02: The 6th Conference on Natural Language Learning 2002 (CoNLL-2002)*.
- Natalia Vlasova, Elena Suleymanova, and Igor Trofimov. 2014. Report on Russian corpus for personal name retrieval. // *Proceedings of TEL’2014 Conference on Computational and Cognitive Linguistics*, P 36–40.
- Juntao Yu, Bernd Bohnet, and Massimo Poesio. 2020. Named Entity Recognition as Dependency Parsing. // *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, P 6470–6476.