

1.1. Автоматическое определение сарказма в текстах на русском языке

А.А.Гурин¹

Т.А.Жуков²

^{1,2}РЭУ им. Г.В. Плеханова. Учебно-научная лаборатория искусственного интеллекта, нейротехнологий и бизнес-аналитики

В данной статье описываются подходы и методы к определению саркастических выражений на русском языке. Было разработано и внедрено решение, позволяющее определять сарказм для микроблоговой платформы Twitter.

Введение

С греческого сарказм (σαρκασμός) переводится как «разрывать плоть». Сарказм является одним из видов сатирического изобличения, язвительной насмешкой, высшей степенью иронии, которая основана не только на усиленном контрасте подразумеваемого и выражаемого, но и на немедленном намеренном обнажении подразумеваемого [Лапина Маталина, Секачев., 2008].. Высказывания и художественные произведения, написанные с сарказмом, утверждают одно, но дают ясно понять, что подразумевают противоположное, например, при помощи издевательской гиперболы или интонации. К сарказму часто прибегали античные ораторы в политической борьбе, чтобы обвинить оппонента.

Сарказм встречается в произведениях, относящихся к эпохе Древней Греции, например, в поэме Гомера «Одиссея» во второй песни: *«Меж тем женихи, изобильный обед учреждая,*

Многими колкими сердце его оскорбляли речами.

Так говорили одни из ругателей дерзко-надменных:

«Нас Телемах поубить не на шутку замыслил; быть может,

Многих он в помощь себе приведет из песчаного Пилоса, многих

Также из Спарты; о том он, мы видим, заботится сильно.

Может случиться и то, что богатую землю Эфиру

Он посетит, чтоб, добывши там яду, смертельного людям,

Здесь отравить им кратеры и разом нас всех уничтожить».

В эпохе Древнего Рима тоже встречается сарказм. Так в словаре латинских выражений Сомова В.П. [Сомов., 1992] встречаются саркастические выражения, например:

An nescis, mi fili, quantilla prudentia mundus regatur? — Сын мой, разве ты не знаешь, как мало надо ума, чтобы управлять миром?

Esse spectaculum dignum, ad quod respiciat intentus operi suo Deus — Вот зрелище, достойное того, чтобы на него оглянулся Бог, созерцая своё творение.

Англо-ирландский писатель-сатирик Джонатан Свифт использовал в своем творчестве приемы иронии и сарказма, так в 1729 году появилась сатира «Скромное предложение». Где автор обращается к правительству и предлагает просто съесть бедных детей Ирландии, чтобы решить проблемы голода и перенаселения. Это саркастическое предложение было адресовано правящему английскому высшему классу того времени. *Один очень образованный американец, с которым я познакомился в Лондоне, уверял меня, что маленький здоровый годовалый младенец, за которым был надлежащий уход, представляет собою в высшей степени восхитительное, питательное и полезное для здоровья кушанье, независимо от того, приготовлено оно в тушёном, жареном, печёном или варёном виде. Я не сомневаюсь, что он также превосходно подойдёт и для фрикасе или рагу.*

Кроме того, сарказм встречается и в классической литературе. Он используется для резкой и жесткой критики лица или объекта действительности. Например: А.С. Грибоедов в комедии «Горе от ума» описывает ситуацию, где отец Софьи Павел Афанасьевич Фамусов, кокетничая со служанкой Лизой, говорит: *«Скромна, а ничего кроме проказ и ветру на уме...»*. Другой пример: Н. В. Гоголь в поэме «Мертвые души» пародийно описывает смерть бедного городского прокурора: *«Послали за доктором, чтобы пустить кровь, но увидели, что прокурор был уже одно бездушное тело. Тогда только с соболезнованием узнали, что у покойника была, точно, душа, хотя он по скромности своей никогда ее не показывал»*. Таким образом, сарказм — это издевка в завуалированной форме, содержащая уничижительную оценку лица, предмета или явления действительности. Сегодня саркастические выражения можно встретить в социальных сетях, микроблогах, сервисах отзывов и многих других платформах.

Распознавание сарказма в текстовых данных является одной из сложных задач обработки естественного языка, в последнее время стала интересной областью исследований из-за ее важности

для интеллектуального анализа эмоциональной окраски текстовых сообщений в социальных сетях [Долбин.,2018]. Данная задача неразрывно связана с проблемой анализа тональности в целом. На сегодняшний день существуют различные методы анализа текста на естественном языке, которые предоставляют возможности по извлечению признаков и установлению семантических связей [Рубцова.,2014],[Gurin.,2020],[Mishra., Dey., Bhattacharyya., 2017]. Однако, автоматическое распознавание тональности сильно отличается от человеческого, в связи с этим задача определения сарказма на естественном языке все еще является актуальной.

С точки зрения анализа текста, задача определения сарказма может быть полезна при анализе ответов пользователей для правильного выбора стратегии общения с виртуальным помощником или для автоматического анализа и верной интерпретации комментариев, отзывов, размещенных в социальных медиа-ресурсах [Долбин.,2018].

Существует ряд исследований, посвященных автоматическому распознаванию сарказма в тексте [Bouazizi., Ohtsuk.,2016], [Son., Kumar., Sangwan.,2019], [Долбин.,2018]. Данные комплексные решения можно классифицировать на две группы. Первая группа включает решения, построенные на лингвистической структуре сарказма. Под этим понимается синтаксический разбор предложений по определенным правилам, вследствие чего представляется возможным определить сарказм. Более детально принцип метода изложен в работах [Bharti., Babu., Jena., 2015],[Riloff., Qadir., Surve., 2013].

Вторая группа решений основана на методах машинного обучения. В данную группу входят подходы с использованием машинного обучения с учителем, глубокое обучение на основе сверточных нейросетей, преднатренированные языковые модели, такие как Word2Vec, GloVe, BERT и другие [Devlin., Chang., 2018],[Евдокимова., 2022].

Для оценки качества построенной модели с использованием методов машинного обучения применяются следующие метрики: Accuracy, Precision, Recall и F1.

Метрики рассчитываются по следующим формулам:

$Accuracy = (TP+TN)/(TP+TN+FP+FN)$

$Precision = TP/(TP+FP)$

$Recall = TP/(TP+FN)$

$F1 = 2*((Precision*Recall)/(Precision + Recall))$

Здесь TP – количество документов, по которым классификатор принял правильное решение среди сообщений в которых присутствует сарказм, TN – количество документов, по которым классификатор принял правильное решение среди сообщений, в которых отсутствует сарказм. FP – количество документов, по которым классификатор принял ошибочное решение, классифицировав их, как содержащие сарказм, FN – количество документов, по которым классификатор принял ошибочное решение, классифицировав их, как не содержащие сарказм.

Решения и исследования в области определения сарказма представлены на английском, испанском, итальянском, французском, японском, китайском, корейском и арабском языках [Baroju., 2022],[Edoardo., 2020],[Elsa., Segura-Bedmar., 2021],[Khalid., 2020]. В работах [Bjarke., Mislove., Oegaard., 217],[Buckley., Paltoglou., Cai., 2019] представлены решения для английского языка. Среди них DeepMoji, созданный в Массачусетском технологическом институте [Bjarke., Mislove., Oegaard., 217]. Данное решение было обучено на 1,2 миллиарда сообщений в Twitter, содержащих смайлики, для обучения использовалась модель глубокого обучения, которая учитывает многие нюансы того, как язык используется для выражения эмоций. Например, она хорошо улавливает сарказм и сленг. Предварительно обученная модель DeepMoji доступна для использования и тестирования, а также реализована для фреймворков Keras и pyTorch. Точность определения сарказма по метрике Accuracy составляет 75%. Другой пример: программное обеспечение SentiStrength изначально предназначенное для определения тональности текста и эмоций [Buckley., Paltoglou., Cai., 2019]. На базе данного программного обеспечения существует надстройка для определения сарказма на английском языке. Точность определения сарказма по метрике Accuracy составляет 65% [Eisner., Rocktaschel., Augenstein.,2021]. Существуют работы, в которых используются архитектуры с LSTM (долгосрочная, краткосрочная память) [Arshad., 2022] и преднатренированные векторные представления слов (GloVe, FastText) [Jeff., Socher., 2014],[Arshad., 2022].

В работе [Reynier., Francisco., Delial., 2021] была построена модель определения сарказма для испанского языка, точность определения по метрике F1 составляет 68,3%. В работе [Tommaso., Nicole., Viviana., 2018] представлено решения для итальянского языка, оценка качества работы модели по метрике Precision составляет 51,8%. В работе [Radu., 2021] создано решение для французского языка точность определения сарказма составляет 74,01% по метрике Accuracy. Для японского языка тоже существуют решения. Так в работе [Sayed., Agrawal., Darkunde., 2020] точность определения сарказма составляет 79% по метрике Precision. В исследовании [Zhu., 2015] точность определения сарказма для китайского языка составляет 73,6%, оценка выполнялась по метрике F1. Для корейского языка в исследовании [Raijmakers., 2021] была построена модель, которая обеспечивает точность в 65,4% по метрике F1. В работе [Mohammed., Eltyeb., Elsamani., 2021] было создано решение для арабского языка, точность определения сарказма 87,23% по метрике F1.

Выбор модели и метода реализации зависит от конкретной задачи и набора данных, на котором происходит обучение модели.

Постановка задачи и исходные данные

Целью данной работы является создание и реализация алгоритма распознавания саркастических выражений на русском языке.

Для достижения цели был выполнен ряд задач, а именно:

- Были собраны наборы данных для обучения моделей для анализа тональности текста и определения сарказма;
- Подготовлены данные для использования в моделях (проведена очистка от специальных символов, приведены все слова в начальную форму);
- Построена модель, определяющая тональность предложений с разбиением на 3 класса (негативный, нейтральный и позитивный);
- Обучена модель определения сарказма.

Так, для построения модели определения сарказма была собрана выборка русскоязычных текстов, которые включают хэш-тэг «сарказм» из социальной сети Twitter. Выборка содержит два набора данных. Первый набор предназначен для тренировки и проверки модели он включает в себя: 6598 сообщений, из которых 5938 предназначены для обучающей выборки и 660 для проверки качества построенной модели. Второй набор предназначался для проверки построенной модели, он не использовался при ее построении и содержит 3044 сообщения. Все сообщения прошли стадию подготовки данных. В них заранее были удалены специальные символы, осуществлен перевод к нижнему регистру всех слов и приведение их в начальную форму с помощью библиотеки `ru morphology2` [Korobov., 2015].

Также, для построения функции анализа тональности был использован общедоступный набор данных `PolSentiLex` [Koltsova., Alexeeva., Kolcov., 2016]. Данный набор содержит 26849 текстов. Из них 10736 относятся к негативному, 13953 к нейтральному и 2160 к позитивному классам. В качестве выделения признаков в тексте использовались языковые модели `Word2Vec`, `Bert` и `ruBert` [Kuratov., Arkhipov., 2019],[Devli., Chang.,2019],[Евдокимова., 2022].

Для дальнейшего выделения саркастически настроенных сообщений, использовался словарь тональности слов [Кулагин., 2022]. Схема работы приложения представлена на рисунке 1.

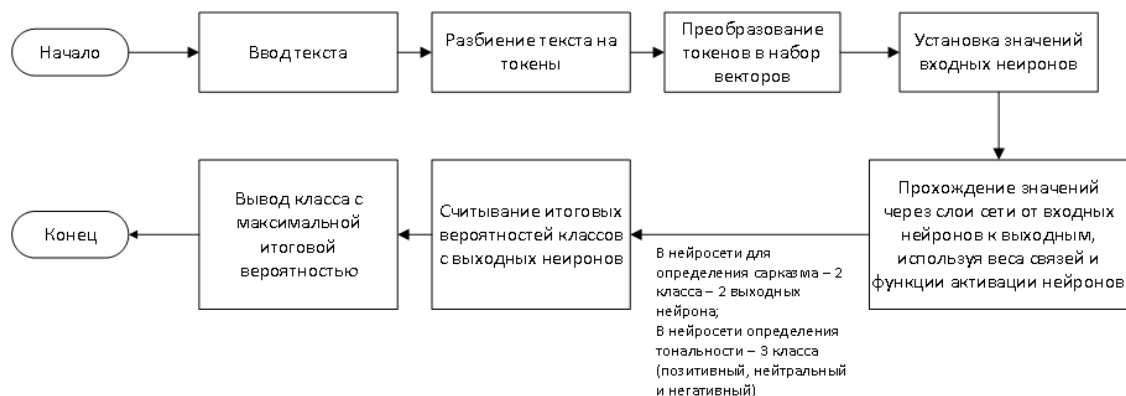


Рис. 1. Схема работы приложения для определения сарказма

Построение решения для определения тональности текста с использованием Word2Vec

В ходе исследования и изучения существующих подходов к распознаванию тональности сообщений была разработана нейросетевая модель, которая является сверточной нейросетью и использует эмбединги `Word2Vec` [Arshad., 2022]. Архитектура данной сверточной нейросети представлена в таблице 1.

Таблица 1. Архитектура сверточной нейросети

Слой	Размерность выхода	Количество параметров	Присоединен к
Входной слой	(N, 2800)	Без представления	Не присоединен
Эмбединг (Представление слов в виде вектора)	(N, 2800, 500)	40353500	Входному слою
Первый слой	(N, 2799, 200)	200200	Эмбедингу
Второй слой	(N, 2798,	300200	Эмбедингу

	200)		
Третий слой	(N, 2797, 200)	400200	Эмбедингу
Четвертый слой	(N, 2796, 200)	500200	Эмбедингу
Пятый слой	(N, 2795, 200)	600200	Эмбедингу
Снижение размерности выхода для первого слоя	(N, 200)	0	Первому слою
Снижение размерности выхода для второго слоя	(N, 200)	0	Второму слою
Снижение размерности выхода для третьего слоя	(N, 200)	0	Третьему слою
Снижение размерности выхода для четвертого слоя	(N, 200)	0	Четвертому слою
Снижение размерности выхода для пятого слоя	(N, 200)	0	Пятому слою
Объединение слоев	(N, 1000)	0	Снижению размерности выхода для слоев:1,2,3,4,5
Метод исключения	(N, 1000)	0	Объединению слоев
Функция активации	(N, 128)	128128	Методу исключения

Всего в модели было использовано 42483015 параметров, из них обучаемых 2129515, необучаемых 40353500.

Данная модель использует 500-мерный вектор параметров слова (embedding) для снижения размерности параметров сети [Kuzmenko, Kutuzov., 2016]. Параметры обучения модели: Обучение происходило на 49 941 примерах, валидация осуществлялась на 5 550 примерах, использовалось 3 эпохи обучения, среднее время эпохи обучения на CPU составляет 2,1 часа, общее время обучения на CPU 6,3 часа.

Проведено обучение и тестирование моделей, основанных на модели НС с различными гипер-параметрами для сетей со сверточными слоями, представлено в таблице 2. В результате тестирования наилучший результат на тестовой выборке показала модель под номером 6.

Табл.2. Результаты обучения и тестирования Word2Vec моделей

Параметр\Модель	1	2	3	4	5	6	7
Количество фильтров	5	10	5	5	5	5	5
Размерность фильтра	200	200	200	200	100	400	200
Метод исключения	0,1; 0,2	0,1; 0,2	0,3	0,5	0,1;0,2	0,1;0,2	0,1; 0,2
Дополнительные признаки							Эмодзи
Точность на обучающем наборе (в %)	75,38	75,39	80,84	75,06	78,73	78,44	77,15
Точность на валидационном наборе (в %)	73,37	73,32	72,40	72,72	72,90	73,96	73,69
Точность на тестовом наборе (в %)	74,13	73,19	73,64	73,47	73,45	74,23	74,11
Время обучения (GPU)	0:14:28	0:31:13	0:20:06	0:21:00	0:08:20	0:26:15	0:13:20
Количество эпох	3	4	6	6	4	4	4
Размерность пакета	32	32	32	32	32	32	32

При оценке точности использовалась мера Ассигасу, которая вычисляется по формуле: Точность = P/N . Здесь P – количество документов, по которым классификатор принял правильное решение, а N – размер выборки.

Также была обучена модель с эмодзи эмбедингом, с итоговой точностью на тестовой выборке в 74,11%, не повысив результат по сравнению с базовым показателем в 74,13% во второй модели. В ней использован 300-мерный эмбединг эмодзи из emoji2vec [Eisner., Rocktaschel., 2021].

До построения вектора признаков текста, собранный набор данных, проходит предобработку, которая включает в себя несколько этапов, а именно:

- Удаление специальных символов (точек, запятых, тире и др);

- Преобразование всех слов и букв в словах к одному регистру, как правило, к нижнему;
- Удаление стоп-слов, которые не несут в себе никакой информации о содержании;
- Процесс лемматизации всех слов. То есть преобразование всех слов к начальной форме.

Технология Word2Vec работает с большим текстовым корпусом и по определенным правилам присваивает каждому слову уникальный набор чисел, так называемый семантический вектор. Библиотека NLTK [Bird., Klein., Loper., 2022] предоставляет инструменты для обработки естественного языка. В их числе есть метод, который позволяет исключать некоторые слова, которые не имеют эмоционального окраса. Смысл этой операции состоит в том, чтобы понизить количество признаков, извлеченных из текста, тем самым сделать модель более точной. Таким образом, был выполнен эксперимент для 6 модели с использованием стоп-слов и без их использования. Результаты данного эксперимента приведены в таблице 3.

Табл.3. Сравнение процесса обучения моделей с фильтрацией стоп-слов

Модель	6 модель без фильтрации стоп-слов		6 модель с фильтрацией стоп-слов	
Выборка	обучающая	валидационная	обучающая	валидационная
1 эпоха	66,20%	71,26%	64,84%	69,55%
2 эпоха	72,45%	72,92%	71,37%	72,77%
3 эпоха	75,47%	73,17%	74,33%	71,48%
4 эпоха	78,44%	73,96%	78,19%	72,43%
Максимум по эпохам		73,96% (4 эпоха)		72,77% (2 эпоха)

При тестировании полученных моделей были получены следующие результаты: при исключении стоп-слов из списка NLTK точность распознавания на тестовой выборке снижается до 72,10% (на 2,1%) по сравнению с моделью без исключений этих стоп-слов (74,24% точности) и на валидационной выборке до 72,77% (на 1,19%) против 73,96%.

Построение решения для определения тональности текста с использованием языковой модели BERT

BERT — это языковая модель, основанная на архитектуре «трансформер», предназначенная для предобучения языковых представлений с целью их последующего применения в широком спектре задач обработки естественного языка [Devlin., Chang. 2019]. По аналогии с Word2Vec используется для выделения признаков. Отличия BERT от Word2Vec заключаются в следующем:

- Модели Word2Vec генерируют вложения, которые не зависят от контекста: то есть — для каждого слова существует только одно векторное (числовое) представление. Разные значения слова (если они есть) объединяются в один вектор. BERT генерирует эмбединг, который позволяет использовать более одного векторного представления для одного и того же слова в зависимости от контекста, в котором это слово используется;
- Вложения Word2Vec не учитывают позицию слова. Модель BERT принимает в качестве входных данных позицию каждого слова в предложении перед вычислением его встраивания;
- Модель Word2Vec изучает вложения на уровне «слова». То есть, модель Word2Vec была обучена на корпусе из 1 миллиона уникальных слов, тогда модель сгенерирует 1 миллион вложений слов — по одному вектору для каждого слова в словаре. Однако такие представления не могут генерировать векторы для слов, встречающихся за пределами словарного пространства. Так, Word2Vec не поддерживает слова вне словаря, что является одним из его основных недостатков. BERT, с другой стороны, изучает репрезентации в «подслово». Подслова рассматриваются как среднее между вложениями на уровне символов и вложениями на уровне слов. Таким образом, модель BERT может иметь словарный запас около 50 тысяч слов, несмотря на то, что она обучена на корпусе из 1 миллиона уникальных слов [Zhu., 2015]. Этот вид моделирования стал очень популярным, потому что модель может генерировать векторное представление для любого слова и не ограничивается словарным пространством.

Таким образом, было проведено обучение нейросетей, основанных на Ru-BERT эмбединге с использованием фреймворка deeppavlov [Blog DeepPavlov., 2021] и BERT [Devlin., Chang. 2019]. Результаты тестирования для различных эмбедингов BERT: многоязычная — 75% точности, русская (deeppavlov), дополненная токенами эмодзи — 78,8% точности. Значения гиперпараметров для наилучшей модели, результаты представлены в таблицах 4 и 5.

Табл.4. Результаты обучения и тестирования deeppavlov RuBERT моделей

Параметр\Модель	1	2	3	4
-----------------	---	---	---	---

Максимальная длина сообщения в токенах	64	200	512	64
Вероятность сохранения веса связей между слоями	0,5	0,5	0,5	0,2
Скорость изменения весов при обучении	1,0E-05	1,0E-05	1,0E-05	1,0E-05
Максимальное количество итераций обучения без уменьшения	5	5	5	5
Точность на обучающем наборе (в %)	84,260	85,020	84,670	83,570
Точность на валидационном наборе (в %)	78,610	79,440	78,280	78,460
Точность на тестовом наборе (в %)	78,820	79,284	78,820	78,310
Время обучения (GPU)	0:18:36	1:14:44	1:21:01	0:19:40
Количество эпох обучения	1	1	1	1
Размерность пакета	64	32	8	64

Табл.5. Результаты обучения и тестирования deerpravlov RuBERT моделей (продолжение)

Параметр\Модель	5	6	7	8
Максимальная длина сообщения в токенах	64	64	64	64
Вероятность сохранения веса связей между слоями	0,5	0,5	0,5	0,5
Скорость изменения весов при обучении	5,0E-06	2,0E-05	1,0E-05	1,0E-05
Максимальное количество итераций обучения без уменьшения	5	5	2	10
Точность на обучающем наборе (в %)	78,830	78,430	78,900	78,360
Точность на валидационном наборе (в %)	85,660	85,500	84,430	84,100
Точность на тестовом наборе (в %)	78,790	78,710	79,000	78,600
Время обучения (GPU)	1:45:59	0:17:31	0:46:10	0:17:42
Количество эпох обучения	2	1	1	1
Размерность пакета	64	64	64	64

Для обучения использовался токенизатор «bert preprocessor» со словарями токенов, включенных в эмбединги. По итогам тестирования была определена наилучшая модель под номером 2 с точностью на тестовой выборке 79,284%. Матрица ошибок для каждого класса представлена в таблице 6.

Табл.6. Матрица ошибок для определения тональности текста

		Актуальные результаты		
		Негативный класс	Позитивный класс	Нейтральный класс
Прогнозируемые результаты	Негативный класс	1813	132	658
	Позитивный класс	150	3590	1089

Нейтральный
класс

731

1072

9263

Модель учета сарказма

В процессе разработки, была составлена обучающая и тестовая выборки из собранного набора данных. Обучающая выборка строится по следующему принципу: если с помощью словаря окраски слов [Кулагин., 2022] текст оценен положительно по 3 классам (негативных, позитивных, нейтральных), и если в тексте положительных слов больше отрицательных, но при этом текст размечен как негативный, то в нем присутствует сарказм. При подсчете слов был использован словарь тональностей русского языка [Кулагин., 2022]. Блок-схема алгоритма определения сарказма в выборке для обучения модели представлена на рисунке 2.

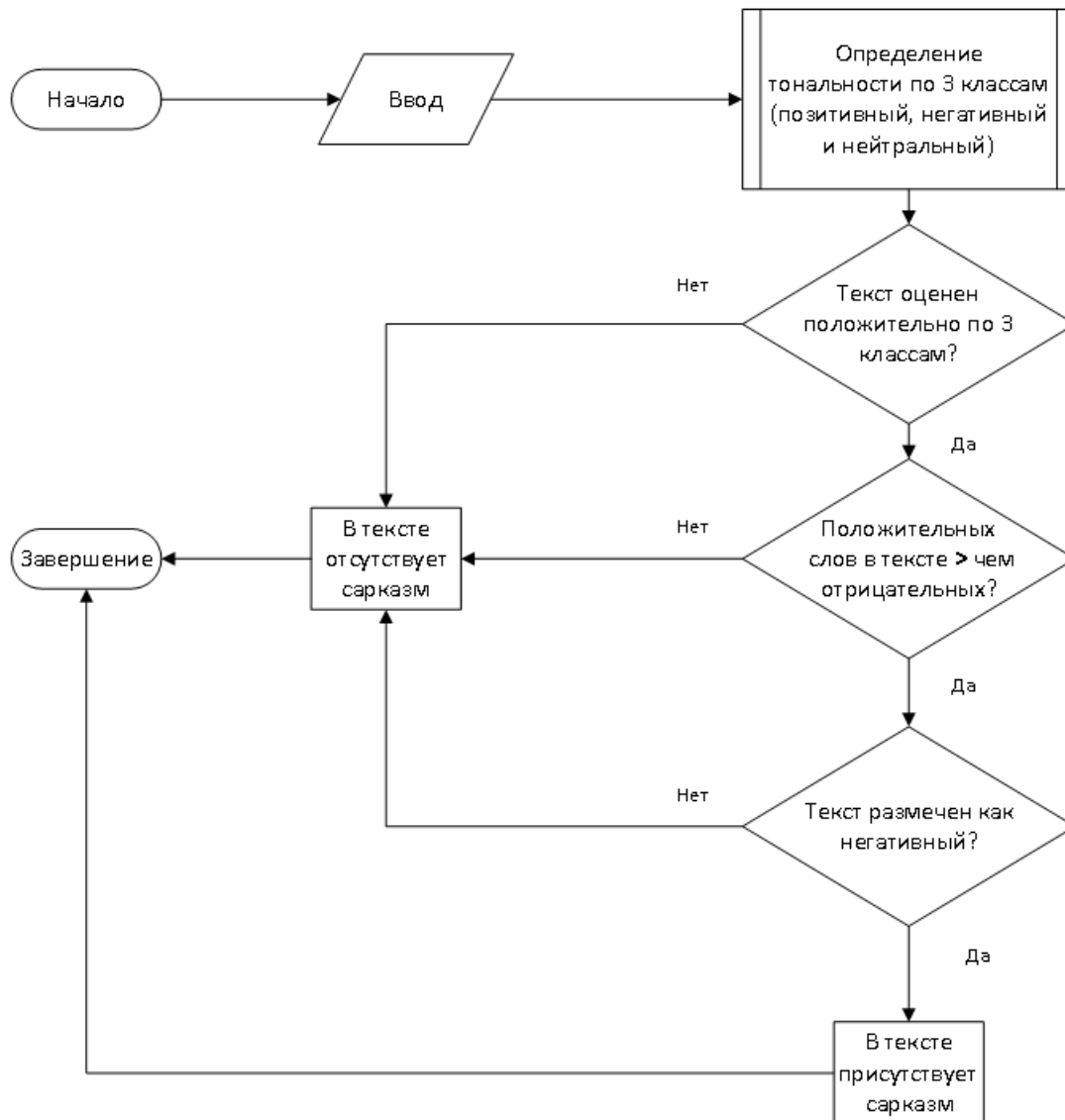


Рис. 2. Блок-схема алгоритма определения сарказма для выборки

Модель была обучена со следующими гиперпараметрами: размерность пакета = 16, максимальная длина сообщения в токенах = 128, вероятность сохранения веса связей между слоями = 0.5, скорость изменения весов при обучении = $1E-05$, максимальное количество итераций обучения без уменьшения = 5. По результатам тестирования итоговая точность модели составила: 78,57%.

После чего выборки для модели обнаружения сарказма в сообщениях были переформированы с учетом обоих вариантов сарказма, как отрицательного, выраженного положительными словами, так и положительного, выраженного отрицательными. После чего модель сарказма была переобучена на этих новых выборках. Итоговая точность распознавания: обучающая выборка 95,49%, валидационная выборка 90,75%, тестовая выборка 69,19%. Матрица ошибок представлена в таблице 7.

Табл.7. Матрица ошибок саркастических выражений

		Актуальные результаты	
		Присутствует сарказм	Отсутствует сарказм
Прогнозируемые результаты	Присутствует сарказм	1048	464
	Отсутствует сарказм	474	1058

По итогам выполненных работ были сделаны следующие выводы:

Для обучения и оценки моделей распознавания сарказма необходима размеченная вручную, отдельная выборка в которой имеется большинство комбинаций слов и других символов, которые можно интерпретировать как саркастические. Без данной выборки невозможно достоверно оценить любой из предлагаемых методов распознавания сарказма.

Также проблема распознавания сарказма очень часто зависит от контекста, в котором передается информация и картины мира говорящего, что делает невозможным правильное распознавание без использования большого количества дополнительной информации, которая не содержится в обучающей выборке.

Исключение или удаление стоп-слов, а также добавление пространства со смайликами негативно сказывается на точности определения модели распознавания тональности текста.

Для дальнейшего тестирования модели было решено собрать другой список саркастических выражений, которые встречаются в различных произведениях, фильмах, книгах и протестировать функцию определения сарказма.

Построенная функция определения сарказма доступна по адресу: <http://passare.ru/sarcasm.php>

Набор данных для тестирования доступен по ссылке: <https://disk.yandex.ru/d/B825xhjg9HC97w>

Точность на собранном наборе данных составляет 68 %. Всего 423 примера, верно определено 288 примеров, ошибочно 135. В таблице 8 представлена численная оценка качества построенной модели.

Табл.8. Оценка качества модели.

Метрика	Точность в %
Accuracy	69,18
Precision	69,31
Recall	68,85
F1	69,08

Правильно были определены следующие примеры (10 штук из выборки):

*У вас ничего не получается с первого раза. Парашютный спорт как раз для вас.
Продолжайте закатывать глаза. Может там вы найдете свой мозг.
Бог юморист. Не верите? Посмотрите в зеркало.
Не каждая серая масса обладает мозгом. Правда?
Я не съел ваш кусок торта. Я просто спасал вашу фигуру.
Сходят с ума только те, у кого он есть.
Может быть, вам дать еще ключ от квартиры, где деньги лежат?
Я хлопал не потому, что мне понравилось выступление, а по той причине, что оно закончилось.
Несмотря на выражения моего лица вы продолжаете говорить?
К счастью для вас зеркала не могут смеяться.*

Ошибочно были определены следующие примеры (10 штук из выборки):

*Продолжайте говорить. Я всегда зеваю, когда мне интересно.
Ваше мнение очень интересно. Давайте поговорим об этом в следующем месяце.
Безусловно я работаю меньше вас. Просто я делаю все правильно с первого раза.
Вам следует записывать все интересные мысли и идеи. На одном стикере для заметок вам
хватило бы места.*

*Ты мне так же нравишься, как понедельники.
Я терпелив, ведь здесь слишком много свидетелей.
Я понимаю, что твои друзья круче, ведь они невидимы.
Ваша пылкая речь не делает вас правым.
Когда мои глаза закрыты, ты лучше всего выглядишь.
Если вы уйдете, то я не буду скучать.*

Заключение

В процессе выполнения работы были достигнуты следующие результаты:

- Был разработан алгоритм определения сарказма в русскоязычных текстах;
- Построена модель определения тональности сообщений на русском языке по трем классам: позитивному, негативному и нейтральному;
- Реализован алгоритм на языке Python 3.7 для определения тональности текстов и сарказма;
- Разработано несколько подходов на языке Python 3.7 определения тональности сообщений;
- На основе сверточных нейронных сетей с использованием Word2Vec эмбедингов с итоговой точностью 74,23%;
- На основе фреймворка deepraiion и эмбединга RuBERT с итоговой точностью 79,44%;
- Разработана модель определения сарказма в сообщениях из Twitter с точностью 69,19%.

Однако, все еще остается ряд вопросов, предложений и улучшений к уже существующим решениям. Так, например, существуют решения для английского языка, которые определяют, стоит ли понимать буквально смысл постов в Twitter, Instagram и прочих социальных медиа. Для этого решение анализирует не только текст, но и изображение, прикрепленное к публикации [Emspak., 2022]. Другой пример, основывается на анализе лексических индикаторов, языковых маркерах и информации о контексте, всех предыдущих твитов и действия пользователя в соцсети. Ввиду того, что изучение только текста недостаточно. Немаловажный контекст обеспечивают изображения. Это особенно актуально для таких соцсетей, как Twitter, Instagram и Tumblr, в которых изображения изначально несут более важную смысловую нагрузку, чем текст. Для обнаружения сарказма в Twitter, Instagram и Tumblr также предложены две различные вычислительные структуры, которые объединяют анализ текстовой и визуальной информации. Первый подход основан на методе опорных векторов, как и в большинстве подобных исследований. Этот метод был дополнен для работы не только с текстовой, но и с визуальной информацией. Второй подход основан на глубинном обучении нейросети на базе изображений ImageNet. Таким образом, лучшие результаты (80-88% распознавания) позволяет получить сочетание двух методов [Emspak., 2022].

Литература

1. Долбин, А. В. Распознавание сарказма в задаче определения тональности текста на естественном языке. Молодой ученый. 2018. № 49 (235). — С. 13-17.
2. Евдокимова И.С., Евдокимов Д.А. Модель идентификации саркастических предложений в естественно-языковом тексте. Интернаука: электрон. научн. журн. 2022. № 22 (245).
3. Кулагин Д. И. Открытый тональный словарь русского языка КартаСловСент URL: https://github.com/dkulagin/kartaslov/tree/master/dataset/emo_dict (Дата обращения: 10.01.2023).
4. Лапина, И. К. Большой энциклопедический словарь. / И. К. Лапина, Е. Н. Маталина, Р. Г. Секачев. — М.: АСТ, Астрель, 2008. — 553 с.
5. Рубцова Ю.В. Методы автоматического извлечения терминов в динамически обновляемых коллекциях для построения словаря эмоциональной лексики на основе микроблоговой платформы Twitter. Доклады Томского государственного университета систем управления и радиоэлектроники. 2014, № 3 (33). — С.140-144.
6. Сомов В. П. По-латыни между прочим. Словарь латинских выражений. — М.: Гитис, 1992. — 230 с.
7. Arshad S., «Sentiment Analysis / Text Classification Using CNN URL: <https://towardsdatascience.com/cnn-sentiment-analysis-1d16b7c5a0e7>. (Дата обращения: 20.01.2023).
8. Baroiu. A. Automatic Sarcasm Detection. Systematic Literature Review. Information 2022, vol.13, p. 399.

9. Bharti S. K., Babu K. S., Jena S. K. Parsing-based Sarcasm Sentiment Recognition in Twitter Data. Proceedings of the 2015 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining., ACM, 2015, p. 1373-1380.
10. Bird S., Klein E., Loper E. Natural Language Processing with Python — O'Reilly Media Inc, 2009. — ISBN 0-596-51649-5
11. Bjarke F, Mislove A., Søgaard A., Using millions of emoji occurrences to learn any-domain representations for detecting sentiment, emotion, and sarcasm. Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing. 2017, p. 1615–1625.
12. Bouazizi M., Ohtsuki T. O. A pattern-based approach for sarcasm detection on Twitter. IEEE Access. 2016, vol. 4, p. 5477-5488.
13. Buckley M, Paltoglou K, Cai G, Kappas, A. Sentiment strength detection in short informal text. Journal of the American Society for Information Science and Technology, issue 61 (12), 2019, p. 2544–2558.
14. DeepPavlov: an open source conversational AI framework. (б.д.) [Электронный ресурс]. – URL: <http://deeppavlov.ai/> (Дата обращения: 20.01.2023).
15. Edoardo S. An Exploration of Sarcasm Detection Using Deep Learning. Corso di laurea magistrale in Ingegneria Informatica (Computer Engineering), 2020, p.68.
16. Eisner B, Rocktäschel T., Augenstein I., Bošnjak M., and Riedel S., emoji2vec: Learning Emoji Representations from their Description. URL: <https://deepemoji.mit.edu/> (Дата обращения: 12.01.2023).
17. Elsa S., Segura-Bedmar I. Sarcasm Detection with BERT. Procesamiento del Lenguaje Natural, Revista nº 67, septiembre de 2021, p. 13–25.
18. Emspak J. Computers Can Sense Sarcasm? Yeah, Right– URL:<https://www.scientificamerican.-com/article/computers-can-sense-sarcasm-yeah-right/> (Дата обращения: 15.01.2023).
19. Gurin A.A. Methods for Automatic Sentiment Detection. Proceedings of the 10th International Scientific and Practical Conference named after A. I. Kitov "Information Technologies and Mathematical Methods in Economics and Management (IT&MM2020)". Moscow, Russia, October 15-16, 2020. CEUR Workshop Proceedings. – 2021. – vol. 2830.
20. J. Devlin Open Sourcing BERT: State-of-the-Art Pre-training for Natural Language Processing. Google AI Blog. Retrieved 2019-11-27. (Дата обращения: 10.01.2023).
21. Jeff P., Socher R., Manning. 2014. GloVe: Global Vectors for Word Representation. Computer Science Department, Stanford University, Stanford. IEEE Access. 2014, p.14.
22. Khalid A., Härmäläinen M., ¡Qué maravilla! Multimodal Sarcasm Detection in Spanish: A Dataset and a Baseline. Computation and Language arXiv preprint :2105.05542.
23. Koltsova O.Y., Alexeeva S.V., Kolcov S.N. An Opinion Word Lexicon and a Training Dataset for Russian Sentiment Analysis of Social Media. Компьютерная лингвистика и интеллектуальные технологии. 2016. С. 277–287.
24. Korobov M. Morphological Analyzer and Generator for Russian and Ukrainian Languages. Analysis of Images, Social Networks and Texts, 2015, p. 320–332.
25. Kuratov. Y., Arkhipov. M. (2019). Adaptation of Deep Bidirectional Multilingual Transformers for Russian Language. arXiv preprint:1905.07213.
26. Mishra A., Dey K., Bhattacharyya P. Learning cognitive features from gaze data for sentiment and sarcasm classification using convolutional neural network. Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics. vol. 1, 2017. p. 377–387.
27. Mohammed M., Eltyeb S., A. Elsamani., Automatic Sarcasm Detection in Arabic Text:A Supervised Classification Approach., International Journal of New Technology and Research (IJNTR) ISSN: 2454-4116, vol.7, Issue 8, August 2021 p. 32-42.
28. Radu T., FreSaDa: A French Satire Data Set for Cross-Domain Satire Detection., arXiv:2104.04828v1
29. Rajmakers D., Dirk R., Emojis, Emoticons and Sarcasm Detection for English and Korean Twitter Data., Thesis submitted in partial fulfillment of the requirements for the degree of master of science in data science and society department of cognitive science and artificial intelligence school of humanities and digital sciences Tilburg university.
30. Reynier O.,Francisco R.,Delia I.,Paolo R.,Manuel M.,Jos´e E., Overview of the Task on Irony Detection in Spanish Variants., CEUR-WS, Vol. 2421.
31. Riloff E., Qadir A., Surve A., De Silva L., Gilbert N., Huang R. Sarcasm as contrast between a positive sentiment and negative situation. Proceedings of the Conference on Empirical Methods in Natural Language Processing, 2013, p. 704–714.
32. Sayed S., Agrawal N., Darkunde M. Deep LSTM-RNN with Word Embedding for Sarcasm Detection on Twitter. 2020 International Conference for Emerging Technology (INCET).
33. Son L., Kumar A., Sangwan S., Arora A, Nayyar A. Sarcasm detection using soft attention-based bidirectional long short-term memory model with convolution network. IEEE Access. vol. 7. 2019. p. 23319–23328.

34. Tommaso C., Nicole N., Viviana P., Paolo R., EVALITA Evaluation of NLP and Speech Tools for Italian., Proceedings of the Final Workshop 12-13 December 2018, Naples.
35. Zhu Y., Aligning Books and Movies. Towards Story-like Visual Explanations by Watching Movies and Reading Books. Aligning Books and Movies. Towards Story-like Visual Explanations by Watching Movies and Reading Books. 2015.

References in Cyrillics

1. 1. Dolbin, A. V. Raspoznavanie sarkazma v zadache opredeleniya tonal'nosti teksta na este-stvennom yazyke . Molodoj uchenyj. 2018. № 49 (235). — S. 13-17.
2. 2. Evdokimova I.S., Evdokimov D.A. Model' identifikacii sarkasticheskikh predlozhenij v estestvenno-yazykovom tekste . Internauka: elektron. nauchn. zhurn. 2022. № 22 (245).
3. 3. Kulagin D. I. Otkrytyj tonal'nyj slovar' russkogo yazyka KartaSlovSent URL: https://github.com/dkulagin/kartaslov/tree/master/dataset/emo_dict (Data obrashcheniya: 10.01.2023).
4. 4. Lapina, I. K. Bol'shoj enciklopedicheskij slovar'. / I. K. Lapina, E. N. Matalina, R. G. Seka-chev. — M.: AST, Astrel', 2008. — 553 с.
5. 5. Rubcova YU.V. Metody avtomaticheskogo izvlecheniya terminov v dinamicheski obnovlyaemykh kollekcijah dlya postroeniya slovarya emocional'noj leksiki na osnove mikroblogoj platformy Twitter. Doklady Tomskogo gosudarstvennogo universiteta sistem upravleniya i radioelektroniki. 2014, № 3 (33). — S.140-144.
6. 6. Somov V. P. Po-latyni mezhdu prochim. Slovar' latinskih vyrazhenij. — M.: Gitis, 1992. — 230 s.

Ключевые слова

определение тональности текста, определение сарказма, модель Word2Vec, языковая модель BERT, обучение с учителем.

Гурин Анатолий Анатольевич,

лаборант-исследователь учебно-научной лаборатории искусственного интеллекта, нейротехнологий и бизнес аналитики РЭУ им. Г.В. Плеханова,

anatoly196674@gmail.com

Жуков Тимур Алекперович,

лаборант-исследователь учебно-научной лаборатории искусственного интеллекта, нейротехнологий и бизнес аналитики РЭУ им. Г.В. Плеханова,

tim29093@gmail.com

Keywords

Sarcasm detection in Russian, sentiment analysis, machine learning, BERT, Word2Vec, supervised learning

DOI: ???

JELclassification C45 – Нейронные сети и связанные темы; M15 Управление информационными технологиями

Abstract

We discuss approaches forwards automatic detection of sarcastic expressions in Russians. A solution for sarcasm detection in Twitter posts is exposed.