

УДК 007:681.512.2

ИДЕНТИФИКАЦИЯ ДОСТОВЕРНОСТИ НОВОСТЕЙ С ПОМОЩЬЮ МОДЕЛЕЙ МАШИННОГО ОБУЧЕНИЯ

И. Ю. Каширин, д.т.н., профессор кафедры ВПМ РГРТУ Рязань, Россия;
orcid.org/0000-0003-1694-7410, e-mail: igor-kashirin@mail.ru.

В статье рассматривается новая технология проектирования обучаемых моделей искусственного интеллекта, предназначенная для оценки правдивости электронных новостных материалов. Используются обучаемые модели со знаниями (Machine Learning, ML-модели). Структуры знаний реализуются в форме семантических сетей, описывающих изначально не вычисленные элементы входных данных.

В технологии применены продукционные экспертные правила, позволяющие определить индекс достоверности фактов, изложенных в материалах новостной статьи. Каждое из экспертных правил требует создания соответствующего программного модуля, основанного чаще всего на методологии синтаксического и семантического анализа естественного языка. Вычисленные индексы применяются как S-элементы нейронных сетей или как входные признаки для обучения, тестирования и валидации ML-моделей.

База знаний программных модулей содержит рейтинговые характеристики электронных изданий и рейтинги авторов новостных статей.

Экспериментальная часть исследований проведена для тестового программного обеспечения, реализованного на языке Python v.4 (Anaconda 4). В качестве исходных текстов новостных статей использованы материалы международного репозитория Kaggle и новостная лента российского сервиса электронной почты mail.ru. Выполненная серия экспериментов дает возможность оценить рассматриваемую технологию как технологию оценки достоверности естественно-языковых текстов, не уступающую по эффективности имеющимся на сегодня международным аналогам.

Целью работы является создание оригинальной технологии автоматизированного анализа естественно-языковых текстов новостей, опубликованных в электронных web-ресурсах, на предмет достоверности содержащихся в них сведений.

Ключевые слова: фейк-новости, достоверность сведений, ML-модели, интеллектуальная обработка данных, база знаний, семантические сети, анализ естественного языка, продукционные правила.

DOI:10.21667/1995-4565-2023-83-36-47

Введение

Самым важным экономическим и политическим фактором современной мировой жизни является информационная война, которая ведется наиболее технологически развитыми странами как во время вооруженных конфликтов, так и в относительно мирное время. Самое смертоносное оружие – широко публикуемые информационные материалы с фальсифицированными сообщениями и трактовками событий. Общее количество электронных средств массовой информации и авторов, являющихся специально мотивированной агентурой истеблишмента стран, претендующих на мировое господство, достигло размеров, поставивших человечество на грань третьей мировой войны. Подготовить адекватное количество отечественных специалистов для разоблачения фейк-новостей в рамках одной страны практически невозможно.

В то же время современные технологии искусственного интеллекта в области исследования больших данных (Big Data) [1] предоставляют возможность автоматизированного или даже автоматического анализа электронных новостных лент, социальных сетей и других информационных web-сервисов на предмет достоверности содержащегося в этих сервисах кон-

тента. Мало того, такие автоматизированные средства разоблачения фейков [2] становятся весьма популярными в молодежной среде и вскоре станут далее настолько же востребованными наиболее продвинутой аудиторией, как и специализированное программное обеспечение информационной безопасности, включая антивирусные программы и программы антиплагиата.

Аналогично антивирусным программам средства автоматизированного разоблачения фальсифицирующей пропаганды должны содержать базы данных, а чаще – базы знаний, оперирующие знаниями о конкретных авторах, информационных источниках и электронных изданиях. Этим самым может быть решена еще одна важнейшая проблема – проблема вычисления фейк ранга (Fake Rank) авторов и средств массовой информации. Эта проблема тем более актуальна, что аналогичные разработки могут создаваться недобросовестными международными организациями для фальсификации самих фейк рангов.

В настоящей статье рассматривается построение нового метода проектирования ML-моделей с помощью привлечения теории представления знаний [3] в прикладные разработки программ интеллектуального анализа данных (Data Mining) с целью получения эффективного инструментария для выявления недостоверных новостных материалов.

Теоретическая часть

Существенными недостатками разрабатываемого на сегодня программного обеспечения для вычисления фейк рангов [2] являются:

- заранее подобранные для обучения ML-моделей входные наборы данных, в которых априорные признаки «достоверный материал»/«фейк» выставлены одними и теми же экспертами;
- использование авторами новых подходов одного и того же репозитория, чаще всего Kaggle [4], который не является идеальным с точки зрения базовых признаков, по которым происходит обучение моделей;
- выделенные признаки опубликованных статей, касающиеся главным образом вторичных характеристик статей.

К вторичным признакам публикаций можно отнести следующие:

- количество символов в названии статьи;
- количество слов текстового материала;
- эмоциональная форма содержания статьи;
- соответствие текста синтаксическим и шаблонным нормам языка, на котором написана статья;
- наличие шокирующих фраз, привлекающих внимание;
- наличие псевдонима, скрывающего настоящие данные автора;
- отсутствие информации о месте и дате публикации;
- дублирование некоторых фрагментов статьи;
- выделение словосочетаний, ограниченных строгой последовательностью из одного, двух или трех слов;
- определение стиля статьи из заранее определенных стилей;
- множественные обращения к мнению зарекомендовавших себя авторов;
- признаки непрофессионализма, такие как использование неинформативных слов «действительно», «в том числе», «по сути дела», «как бы», «в данном случае», «непосредственно».

Ограничение исследования данных такими признаками может быть не только неэффективно, но и контрпродуктивно. Это связано с тем, что первоначальными источниками сведений могут оказаться вовсе не дипломированные и/или специально подготовленные журналисты, а непосредственные участники актуальных событий, часто попросту неграмотные, но честные люди.

Излагаемый здесь подход ориентирован на использование для обучения ML-моделей, кроме уже рассмотренных характеристик текстов, онтологических моделей знаний [5]. В

этом случае признаки входного набора данных вначале исследуются инженером по знаниям как оцифрованные свойства концептов, отношений и элементов динамики предметной области. При этом некоторые закономерности, например причинно-следственные отношения, могут быть установлены априори.

Кроме того, для предметной области проектируется соответствующая ей онтологическая схема, которая является фундаментом для установления неявных отношений семантической близости на множестве синтограмм (словосочетаний) из заранее составленных шаблонов словарных конструкций. Для оцифровки отношений семантической близости удобно использовать теорию иерархических чисел [6]. Онтологические модели позволяют также выявлять существование дополнительных признаков, которые могут вычисляться локализованными программными модулями. Вычисление некоторых входных параметров требует использования объемных баз данных и баз знаний, внутренняя организация которых опирается на ранее разработанную онтологическую модель.

Архитектура программного комплекса (автоматизированной системы идентификации), реализующая взаимодействие основных элементов системы идентификации достоверности новостей, представлена рисунком 1. В соответствии с архитектурой входными данными программы являются:

- электронные новостные материалы;
- обучающий набор данных;
- база синтограмм;
- онтологическая база знаний;
- база данных, состоящая из идентификатора изданий и идентификатора авторов.

Первые два компонента входных данных формируются администратором программного комплекса в соответствии с классической последовательностью работы с обучением, настройкой и эксплуатацией ML-моделей. Оставшиеся три компонента представляют собой базы данных и знаний, формируемые и администрируемые инженером по знаниям. Инженер составляет синтограммы для каждого нового фактографического материала, разрабатывает и регулярно актуализирует онтологическую базу знаний, а также следит за коррекцией рейтинговых идентификаторов авторов и изданий. Все перечисленные работы в развитых версиях автоматизированной системы идентификации должны быть полностью автоматизированы. Кроме того, вычисление идентификаторов происходит автоматически на основе рейтингов достоверности, полученных в результате работы ML-моделей технологии Data Minig. Здесь могут быть применены алгоритмы обобщения значений, например среднее арифметическое или медианное значение всех ретроспективных рейтингов автора или электронного издания.

Основными модулями предложенной архитектуры являются:

- NLP-процессор;
- matching-модуль вычисления семантической близости;
- синтезатор дополнительных признаков;
- ансамбль ML-моделей;
- метамодель.

Промежуточными данными в соответствии с архитектурой являются семантические образы материалов, представляющие собой унифицированную схему новостного материала, созданную matching-модулем в результате предварительной обработки естественно-языкового текста электронной новости и определения ключевых слов, идентифицирующих новость с вариантами утверждения фактографии (true fact) или ее опровержения (false fact). Полный набор семантических образов, полученных за все время работы системы идентификации, составляет главную часть базы синтограмм.

Рассмотрим теперь работу модулей идентифицирующей системы более подробно.

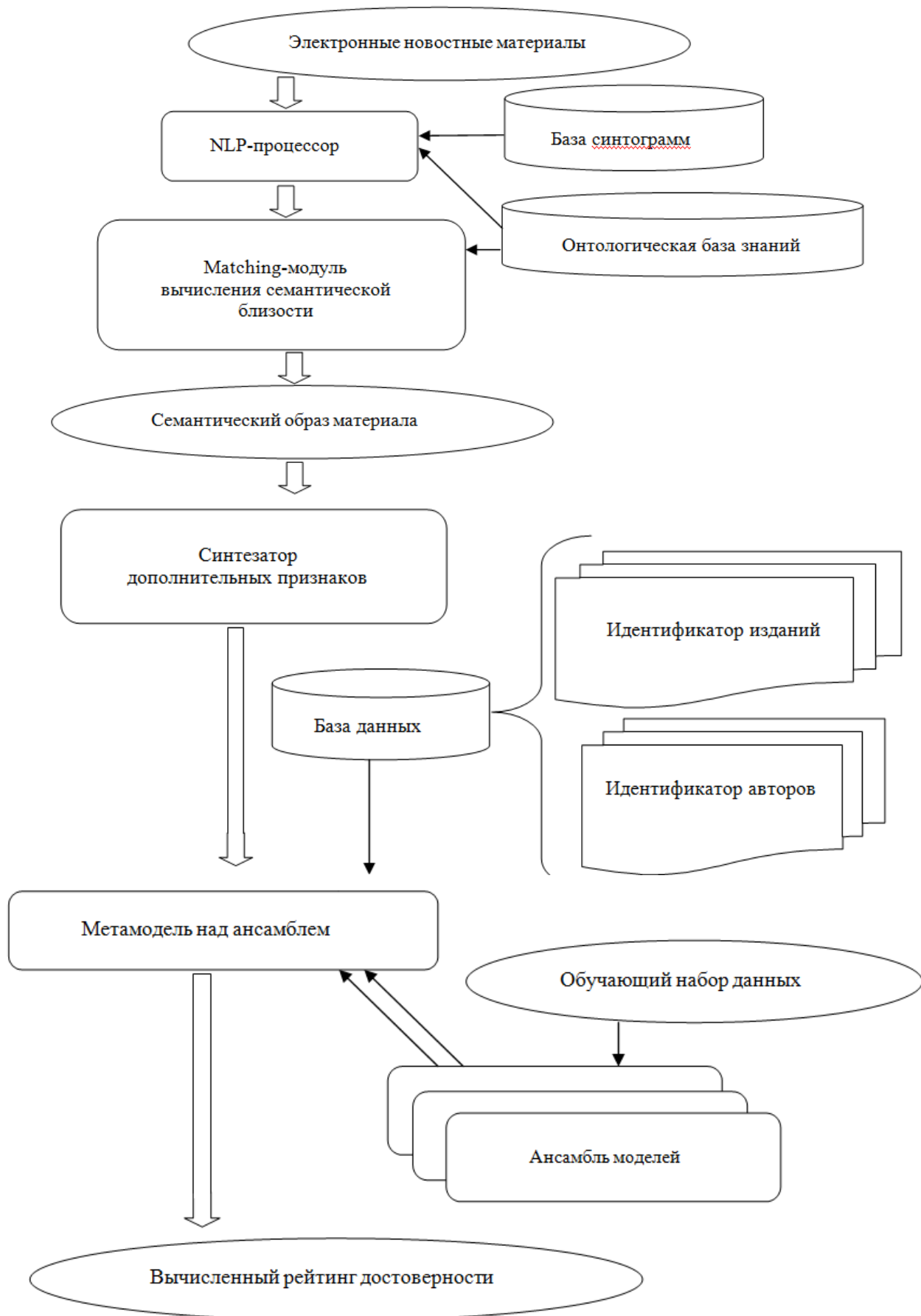


Рисунок 1 – Архитектура системы идентификации достоверности
 Figure 1 – Fake rank calculation system architecture

Для пояснения работы системы идентификации будут использоваться примеры, взятые из известного новостного web-сервиса mail.ru (рисунок 2).

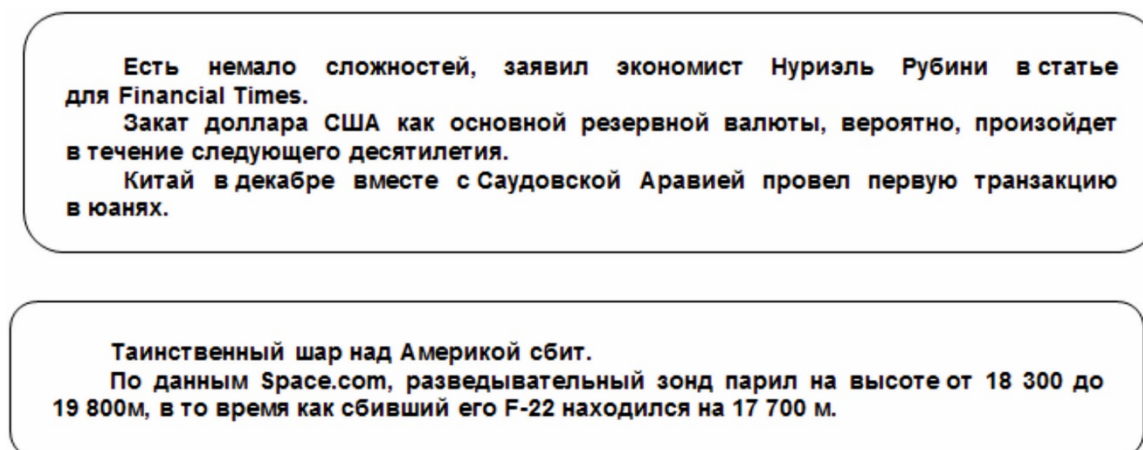


Рисунок 2 – Пример электронных новостей из mail.ru
Figure 2 – Example of electronic news from mail.ru

Первым этапом анализа приведенных новостей является NLP (Natural Language Processing), т.е. обработка входных предложений материала средствами естественно-языковых лингвистических процессоров, таких как Natasha [7] или Spacy 3.3 [8]. Этими процессорами осуществляется базовый синтаксический разбор предложений, классически состоящий из этапов морфологического и синтаксического анализов. Адаптированные тексты новостных статей, получаемые в результате автоматизированной обработки препроцессором нормализации текстов или инженером по знаниям (Matching-модуль), выглядят следующим образом (рисунок 3).

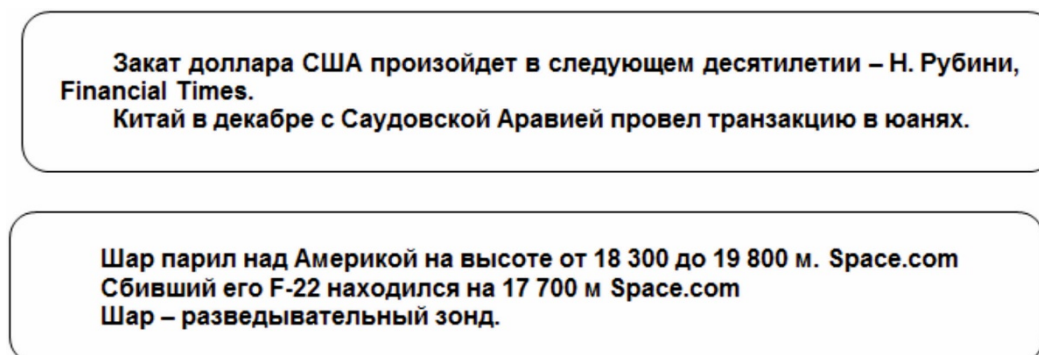


Рисунок 3 – Пример нормализованного текста электронных новостей
Figure 3 – Example of normalized e-news text

Естественно-языковая обработка в инструментальной среде Python 3.8 Google Colab для инструментария Spacy 3.3 имеет следующий программный код:

```
!pip install natasha spacy==3.3
import spacy
!python -m spacy download ru_core_news_sm
nlp = spacy.load('ru_core_news_sm')
text1 = 'Закат доллара США произойдет в следующем десятилетии'
text2 = 'Китай с Саудовской Аравией
        провел транзакцию в юанях.'

Doc1 = nlp(text1)
Doc2 = nlp(text2)
for token in Doc1:
    print (token.text1, token.pos_, token.dep_)
```

Результатом работы станет следующая выходная форма:

```
«Закат NOUN nsubj
доллара NOUN nmod
США PROPN nmod
произойдет VERB ROOT
в ADP case
следующем ADJ amod
десятилетия NOUN obl
. PUNCT punct»
```

```
for token in Doc2:
    print (token.text2, token.pos_, token.dep_)
```

Для второго текста результат будет следующий:

```
«Китай PROPN nsubj
с ADP case
Саудовской ADJ amod
Аравией PROPN nmod
провел VERB ROOT
транзакцию NOUN obj
в ADP case
юнях NOUN obl»
```

```
displacy.render(Doc1, style="dep", jupyter=True)
```

Деревом синтаксического подчинения для первого предложения будет рисунок 4.

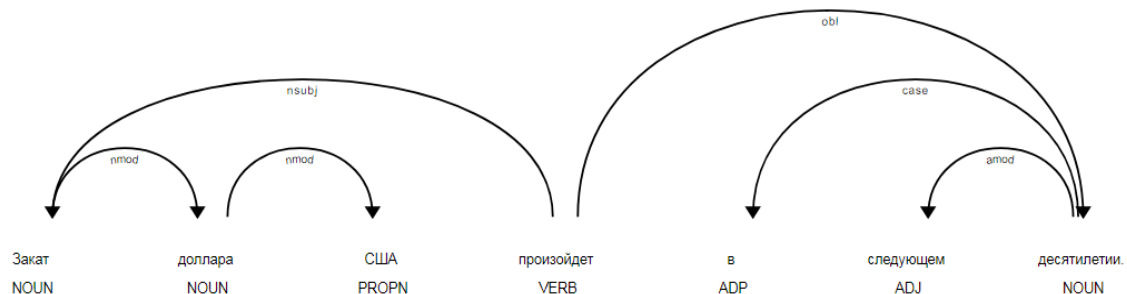


Рисунок 4 – Дерево синтаксического подчинения для первого предложения

Figure 4 – Syntactic subordination tree for the first sentence

```
displacy.render(Doc2, style="dep", jupyter=True)
```

Деревом синтаксического подчинения для второго предложения будет рисунок 5.

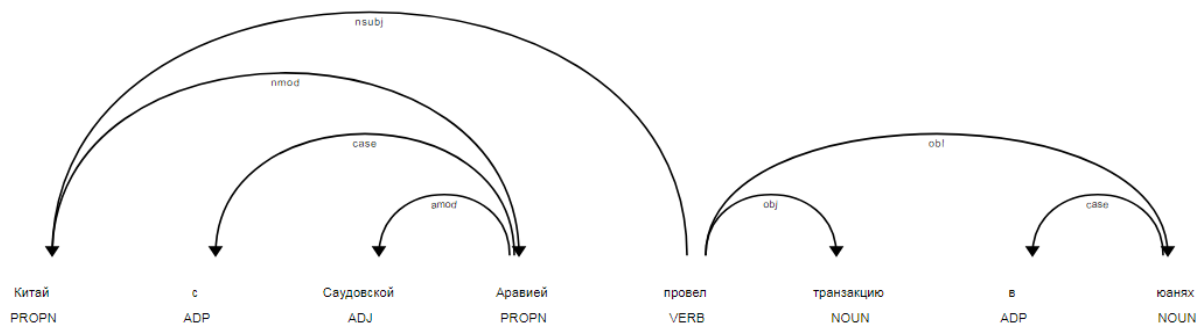


Рисунок 5 – Дерево синтаксического подчинения для второго предложения

Figure 5 – Syntactic subordination tree for the second sentence

Работа с онтологической базой знаний (рисунок 8) сводится к выявлению фактов новости, которые могут иметь ложную или истинную интерпретацию в предметной области. Факты новостной статьи являются экземплярами фактов в родовидовой таксономии базы знаний. То же самое касается сущностей «Автор» и «Издание». Их экземпляры имеют ключевые слоты «Ссылки» и «Достоверность», предназначенные для хранения информации о взаимных упоминаниях в одних источниках информации на другие, а также числовой рейтинг достоверности (Fake Rank) соответственно.



Рисунок 8 – Фрагмент онтологической базы знаний
Figure 8 – Fragment of ontological knowledge base

В дальнейшем элементы онтологической базы знаний участвуют в формировании дополнительных признаков в ML-моделях.

Следующим этапом предлагаемой технологии в соответствии с общей архитектурой является использование синтезатора дополнительных признаков (рисунок 1). Его задача – подвергнуть анализу полученные деревья синтаксического подчинения и семантические описания, выделив существенный факт новости, например использовав такой код:

```

docS.indication_1(segmenter)
for indication in docS:
    print (token_bold.texts, token.pos_, token.dep_)
  
```

Результатом применения метода индикации будет следующая структура (рисунок 9):

Закат	<u>nsubj</u>	Закат	<u>nsubj</u>
доллара		доллара	
США	<u>nmod</u>	США	<u>nmod</u>
произойдет		произойдет	
в	<u>case</u>	в	<u>case</u>
следующем	<u>amod</u>	следующем	<u>amod</u>
десятилетия	<u>obl</u>	десятилетия	<u>obl</u>
–	<u>punct</u>	–	<u>punct</u>
Н	<u>parataxis</u>	Н	<u>parataxis</u>
.	<u>punct</u>	.	<u>punct</u>
Рубини	<u>flat:name</u>	Рубини	<u>flat:name</u>
Financial	<u>punct</u>	Financial	<u>punct</u>
Times	<u>conj</u>	Times	<u>conj</u>
	<u>flat:foreign</u>		<u>flat:foreign</u>

Рисунок 9 – Структура факта новости
Figure 9 – News fact structure

Полученный факт «закат доллара» становится информационной единицей соответствующего справочника базы данных и получает собственный идентификатор `id_fact`, включае-

мый в дополнительные вводные данные ML-моделей. В качестве дополнительных признаков включаются также имя автора и наименование издания с собственными идентификаторами `id_avtor` и `id_edition` соответственно.

Синтезатор семантических признаков в развитых версиях технологии идентификации достоверности новостей включает в себя библиотеку программных модулей, определяющих на основе базы синтограмм (рисунок 1) и, возможно, интеллектуальных web-агентов сложные признаки вычисления индексов достоверности, некоторые из которых собраны в таблицу 1. В таблице приведены кроме прочего базовые значения достоверности со знаками «+» или «-», что указывает на повышения рейтинговой оценки достоверности материала или понижения этого рейтинга. Эти значения могут уточняться при обучении ML-моделей [9-11].

Существует множество прямых и косвенных признаков, позволяющих определять рейтинг достоверности для авторов материалов и опубликовавшего материалы web-сервиса. Эти признаки требуют программирования специализированных интеллектуальных web-агентов, использующих информационные хранилища и прибегающих к знаниям о владельцах информационных ресурсов, принадлежности автора к элите какой-либо страны или сообществу с явно выраженными политическими и/или экономическими интересами. При проектировании интеллектуальных агентов обычно используются продукционные экспертные системы, содержащие, например, следующие правила:

- 1) лицо, отвергающее заведомую истину, может лгать;
- 2) лицо, искажающее исторические факты, лжет;
- 3) лицо, заинтересованное в фейке, может лгать;
- 4) лицо, заинтересованное в фейке, но опровергающее его, может говорить правду.

Особенности вычисления таких признаков оставлены за рамками настоящей статьи.

Таблица 1 – Таблица словарных категорий индикаторов

Table 1 – Table of indicator vocabulary categories

№ п/п	Признаки достоверной/недостоверной новости	Формула синтограммы	Базовый коэффициент достоверности
1	Совершено преступное действие, и известно лицо, назначившее виновника до расследования	«это сделали» X, «вина ложится на» X, «удар нанесен» X, X «должны взять ответственность»	-0,75
2	Автор, осведомленный о важной информации, сообщает, что не знает этой информации	«нам неизвестно», «мы не располагаем данными», «нет информации о»	-0,60
3	Истинная информация опирается на исторические документы, однако фальсифицирована в материале	«как известно», X «начали» Y, X «являются причиной»	+0,95
4	Автор сообщает о лживости другого автора	X «лживо утверждает», X «неправду», X «замалчивает факт»	-0,55
5	Источник лживой информации признает сам, что опубликованная им информация – ложь	«мы ошибались», «оказалось не так», «приносим извинения»	-0,3
6	Расследование началось, но не закончилось в разумные сроки	«расследование продолжается», «расследование» X «пока неизвестно»	-0,7

7	Лицо, сообщающее о факте и говорящее, что есть подтверждающая информация, но не дающее этой информации	«есть сведения», «есть информация», «нам стало известно», «доподлинно известно», «уже ясно, что»	-0,75
8	Публикация анонимна	Поиск подписи	-0,2
9	Утверждения крайностей	«никогда не лжет», «всегда говорит», «всегда ошибаются»	-0,3
10	Опубликованная новость быстро удалена с ресурса	Отслеживание длительности публикации	-0,75
11	Признаки боязни (неуверенности в своей правоте)	«скорее всего», «некоторые страны», «почти»	-0,25

Заключительным этапом идентификации достоверности новостной статьи является работа ансамбля ML-моделей, полученных на ранних этапах разработки с помощью обучения с учителем. Как правило, используются наиболее известные методы Data Minig:

- деревья решений;
- градиентный спуск;
- бустинг;
- метод опорных векторов;
- нейронные сети.

Метамоделю, обобщающей результаты этих методов, и обучаемой на их результатах, становится, например, логистическая регрессия.

Экспериментальные исследования

Эксперименты производились с микс-датасетом, состоящим из данных репозитория Kaggle [4] и накопленных в течение полугода данными новостных лент электронных русскоязычных ресурсов mail.ru, lenta.ru.

Качественные характеристики ML-модели, оценивающей новостные материалы с помощью описанной в статье технологии, представлены на рисунке 10.

ОТОБРАЖЕНИЕ ROC - КРИВОЙ

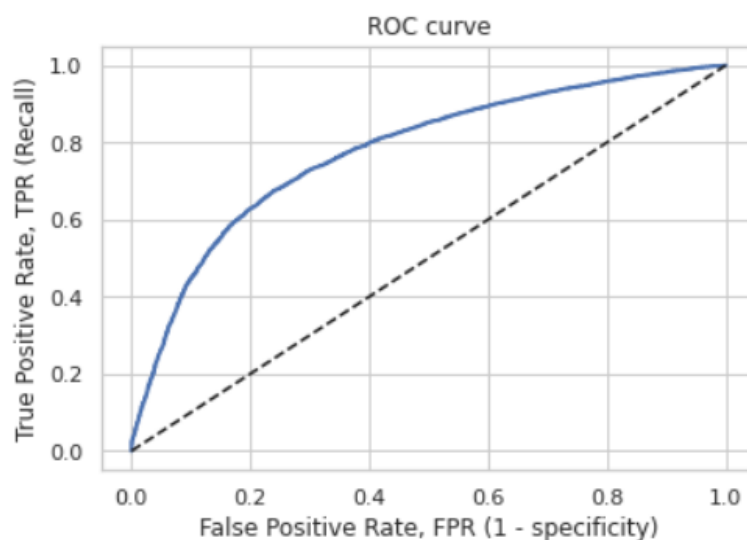


Рисунок 10 – Характеристики качества ML-моделей
Figure 10 – Characteristics of ML-models quality

Оптимизация набора входных S-элементов нейронной сети на основе предложенной технологии повышает на 2-8 % качество идентификации ранга достоверности, что подтверждено результатами экспериментов (более 10).

Эксперименты показали, что испытательный стенд (ЭСМИАД v.18.01.2023), реализованный на языке Python v.3.7 в среде Anaconda 3.3, дает возможность повысить эффективность идентификации новостных статей в различных предметных областях.

В качестве вывода можно констатировать, что для проведенных экспериментов характеристика macro-average ROC AUC score (площадь под ROC кривой, receiver operating characteristic, рабочая характеристика приёмника) после применения новой технологии возросла от 0,69 до 0,73.

Библиографический список

1. **Каширин И. Ю.** Data Mining с использованием иерархических чисел в ретроспективной диагностике // Вестник Рязанского государственного радиотехнического университета. 2022. № 79. С.118-126.
2. **Каширин И. Ю.** Интеллектуальный анализ данных в ретроспективной диагностике // V Международный научно-технический форум СТНО-2022. Сборник трудов. Том 4, Рязань, Book Jet, 2022, С. 11-17.
3. **Каширин И. Ю.** Иерархические числа в проектировании общих онтологий // V Международный научно-технический форум СТНО-2022. Сборник трудов. Том 4, Рязань, Book Jet, 2022, С.17-24.
4. Международный репозиторий для анализа данных и оригинальных технологических решений. [Электронный ресурс]. 2023. Дата обновления: 17.01.2023. URL: <https://www.kaggle.com/> (дата обращения: 17.01.2022).
5. **Каширин И. Ю.** Применение теории иерархических чисел в проектировании ICF-таксономии для оптимизации нейронных сетей // Вестник Рязанского государственного радиотехнического университета. 2022. № 80. С. 36-43.
6. **Утлик Э. П.** Психология личности: Учебник / Э.П. Утлик. М.: Academia, 2018. 448 с.
7. **Kashirin I. Yu., Filatov I. Yu.** Formalized Description Of Intuitive Perception Of Spatial Situations. 2019 8th Mediterranean Conference on Embedded Computing (MECO), Budva, Montenegro, 2019, pp. 1-4.
8. Natasha – качественное компактное решение для извлечения именованных сущностей из новостных статей на русском языке. [Электронный ресурс]. 2023. Дата обновления: 09.02.2023. URL: <https://natasha.github.io/ner/> (дата обращения: 09.02.2023).
9. Можно всё: решение NLP задач при помощи spacy. [Электронный ресурс]. 2023. Дата обновления: 09.02.2023. URL: <https://habr.com/ru/post/531940/> (дата обращения: 09.02.2023).
10. **Каширин И. Ю.** Нейронные сети для идентификации пользователя на основе анализа посещений новостного сайта // Вестник Рязанского государственного радиотехнического университета. 2022. № 82. С.104-111.
11. **Сатыбалдина Д. Ж., Овечкин Г. В., Калымова К. А.** Система распознавания статических жестов рук с использованием камеры глубины // Вестник Рязанского государственного радиотехнического университета. 2020. № 72. С. 93-105.

UDC 007:681.512.2

NEWS RELIABILITY IDENTIFICATION USING ML-MODELS

I. Yu. Kashirin, Dr. Sc. (Tech.), full professor, RSREU, Ryazan, Russia;
orcid.org/0000-0003-1694-7410, e-mail: igor-kashirin@mail.ru

The article discusses a new technology for designing trainable artificial intelligence models, designed to assess the veracity of electronic news materials. Trained models with knowledge (Machine Learning, ML-models) are used. Knowledge structures are implemented in the form of semantic networks that describe initially uncalculated input data elements.

The technology uses production expert rules to determine the index of reliability of the facts presented in the materials of a news article. Each of expert rules requires the creation of an appropriate program module, most often based on the methodology of syntactic and semantic analysis of natural language. The calculated indices are used as S-elements of neural networks or as input features for training, testing and validation of ML models.

The knowledge base of program modules contains rating characteristics of electronic publications and ratings of the authors of news articles.

The experimental part of the research was carried out for test software implemented in Python v.4 (Anaconda 4). As source texts for news articles the materials from international Kaggle repository and news feed of Russian mail.ru e-mail service were used.

The performed series of experiments makes it possible to evaluate the technology under consideration as a technology for assessing the reliability of natural language texts, which is not inferior in efficiency to international analogues available today.

The aim of the work is to create an original technology for automated analysis of natural language news texts published in electronic web resources for the reliability of the information contained in them.

Key words: fake news, reliability of information, ML models, data mining, knowledge base, semantic networks, natural language analysis, production rules.

DOI:10.21667/1995-4565-2023-83-36-47

References

1. **Kashirin I. Yu.** Data Mining s ispol'zovaniyem iyerarkhicheskikh chisel v retrospektivnoy diagnostike. *Vestnik Ryazanskogo gosudarstvennogo radiotekhnicheskogo universiteta*. 2022, no. 79, pp. 118-126. (in Russian).
2. **Kashirin I. Yu.** Intellektual'nyy analiz dannykh v retrospektivnoy diagnostike. *V Mezhdunarodnyy nauchno-tekhnicheskiiy forum STNO-2022. Sbornik trudov*, vol. 4, Ryazan', Book Jet, 2022, pp.11-17. (in Russian).
3. **Kashirin I. Yu.** Iyerarkhicheskiye chisla v proyektirovaniy obshchikh ontologiy. *V Mezhdunarodnyy nauchno-tekhnicheskiiy forum STNO-2022. Sbornik trudov*. Tom 4, Ryazan', Book Jet, 2022, pp.17-24. (in Russian).
4. *Mezhdunarodnyy repozitoriy dlya analiza dannykh i original'nykh tekhnologicheskikh resheniy*. [Elektronnyy resurs]. 2023. Data obnovleniya: 17.01.2023. URL: <https://www.kaggle.com/> (data obrashcheniya: 17.01.2022). (in Russian).
5. **Kashirin I. Yu.** Primeneniye teorii iyerarkhicheskikh chisel v proyektirovaniy ICF-taksonomii dlya optimizatsii neyronnykh setey. *Vestnik Ryazanskogo gosudarstvennogo radiotekhnicheskogo universiteta*. 2022, no. 80, pp.36-43. (in Russian).
6. **Utlik, E. P.** *Psikhologiya lichnosti*: Uchebnik E. P. Utlik. Moscow: Akademiya. 2018. 448 p. (in Russian).
7. **Kashirin I. Yu., Filatov I. Yu.** Formalized Description Of Intuitive Perception Of Spatial Situations. *2019 8th Mediterranean Conference on Embedded Computing (MECO)*, Budva, Montenegro, 2019, pp. 1-4.
8. Natasha – kachestvennoye kompaktnoye resheniye dlya izvlecheniya imenovannykh sushchnostey iz novostnykh statey na russkom yazyke. [Elektronnyy resurs]. 2023. Data obnovleniya: 09.02.2023. URL: <https://natasha.github.io/ner/> (data obrashcheniya: 09.02.2022). (in Russian).
9. *Mozhno vso: resheniye NLP zadach pri pomoshchi spacy*. [Elektronnyy resurs]. 2023. Data obnovleniya: 09.02.2023. URL: <https://habr.com/ru/post/531940/> (data obrashcheniya: 09.02.2022). (in Russian).
10. **Kashirin I. Yu.** Neyronnyye seti dlya identifikatsii pol'zovatelya na osnove analiza poseshcheniy novostnogo sayta. *Vestnik Ryazanskogo gosudarstvennogo radiotekhnicheskogo universiteta*. 2022, no. 82, pp.104-111. (in Russian).
11. **Satybaldina D. Zh., Ovechkin G. V., Kalymova K. A.** Sistema raspoznavaniya staticheskikh zhestov ruk s ispol'zovaniyem kamery glubiny. *Vestnik Ryazanskogo gosudarstvennogo radiotekhnicheskogo universiteta*. 2020, no. 72, pp. 93-105. (in Russian).