

УДК 004.81
МРНТИ 20.19.27

Жаксыбаев Дархан Оракбаевич, доктор философии (PhD), <https://orcid.org/0000-0001-6355-5431>

НАО «Западно-Казахстанский аграрно-технический университет имени Жангир хана», г. Уральск, ул. Жангир хана 51, 090009, Казахстан, darhan.03.92@mail.ru

Zhaxybayev Darkhan Orakbayevich, Doctor of Philosophy (PhD), <https://orcid.org/0000-0001-6355-5431> Zhangir Khan Agrarian and Technical University of West Kazakhstan, Uralsk, 51 Zhangir Khan St., 090009, Kazakhstan, darhan.03.92@mail.ru

**ИСПОЛЬЗОВАНИЕ РАЗЛИЧНЫХ МЕТОДОВ ДЛЯ ИЗВЛЕЧЕНИЯ
И ИДЕНТИФИКАЦИИ КЛЮЧЕВЫХ СЛОВ ИЗ НАУЧНЫХ ТЕКСТОВ НА
КАЗАХСКОМ ЯЗЫКЕ
USING VARIOUS METHODS TO EXTRACT AND IDENTIFY KEYWORDS FROM
SCIENTIFIC TEXTS IN THE KAZAKH LANGUAGE**

АННОТАЦИЯ

В данной статье мы представляем исследование по повышению точности извлечения ключевых слов для научных текстов на казахском языке, в частности, в области экологии. Мы исследуем эффективность нескольких алгоритмов, включая Kea, WINGNUS, TextRank, MultipartiteRank, LSA, LDA и Word2Vec для извлечения ключевых слов. Наша цель - продемонстрировать, что алгоритм Word2Vec работает лучше других в определении релевантных ключевых слов в научных текстах на казахском языке. Наши результаты показали, что алгоритм Word2Vec превосходит другие алгоритмы по точности, отзыву и F1-score для извлечения ключевых слов в научных текстах по экологии на казахском языке. В частности, алгоритм Word2Vec достиг точности 0,87, recall 0,86 и F1-score 0,86, в то время как второй по эффективности алгоритм (LSA) достиг точности 0,83, recall 0,81 и F1-score 0,81. Мы заметили, что способность Word2Vec улавливать семантические связи между словами и фразами помогла выявить ключевые слова, тесно связанные с основной темой текста. Алгоритмы LSA и LDA также показали многообещающие результаты, но они были не так эффективны, как Word2Vec, в определении релевантных ключевых слов. Эти результаты имеют важное значение для исследователей и практиков в области экологии, работающих с научными текстами на казахском языке, поскольку они могут использовать алгоритм Word2Vec для более точного и эффективного извлечения ключевых слов. Данное исследование также вносит вклад в более широкие исследования по улучшению извлечения ключевых слов в неанглийских языках.

ANNOTATION

In this paper, we present a study on improving the accuracy of keyword extraction for scientific texts in the Kazakh language, particularly in the field of ecology. We investigate the effectiveness of several algorithms including Kea, WINGNUS, TextRank, MultipartiteRank, LSA, LDA and Word2Vec for keyword extraction. Our goal is to demonstrate that the Word2Vec algorithm performs better than others in identifying relevant keywords in scholarly texts in Kazakh. Our results showed that the Word2Vec algorithm outperforms other algorithms in accuracy, recall, and F1-score for keyword extraction in environmental science texts in Kazakh. Specifically, the Word2Vec algorithm achieved 0.87 accuracy, recall 0.86 and F1-score 0.86, while the second most efficient algorithm (LSA) achieved 0.83 accuracy, recall 0.81 and F1-score 0.81. We observed that Word2Vec's ability to capture semantic relationships between words and phrases helped identify keywords closely related to the main topic of the text. The LSA and LDA algorithms also showed promising results, but they were not as effective as Word2Vec in identifying relevant keywords. These results have important implications for environmental researchers and practitioners working with scholarly texts in Kazakh, as they can use the Word2Vec algorithm to extract keywords more accurately and efficiently. This study also contributes to broader research on improving keyword extraction in non-English languages.

Ключевые слова: извлечение ключевых слов, обработка естественного языка, Word2Vec, тест Фридмана, научные тексты, анализ данных.

Key words: keyword extraction, Natural Language Processing, Word2Vec, Friedman test, scientific texts, data analysis.

Введение. В данной статье мы представляем исследование по повышению точности извлечения ключевых слов для научных текстов на казахском языке, в частности, в области экологии. Мы исследуем эффективность нескольких алгоритмов, включая Kea, WINGNUS, TextRank, MultipartiteRank, LSA, LDA и Word2Vec для извлечения ключевых слов. Наша цель - продемонстрировать, что алгоритм Word2Vec работает лучше других в определении релевантных ключевых слов в научных текстах на казахском языке.

Алгоритм Kea основан на анализе совместных совпадений слов в данном документе и их связей с другими словами в документе. Его преимущество заключается в способности идентифицировать многословные выражения как ключевые слова. Однако при извлечении однословных ключевых слов он страдает от низкого уровня запоминания[1].

Алгоритм WINGNUS - это графовый подход, использующий алгоритм кластеризации для извлечения ключевых слов из текста. Его преимущество заключается в способности извлекать многословные фразы в качестве ключевых слов. Однако он требует больших вычислительных затрат и может страдать от низкой точности[2].

Алгоритмы TextRank и MultipartiteRank - это графовые подходы, использующие алгоритм, подобный PageRank, для ранжирования важности слов и фраз в тексте. Они эффективны и показали многообещающие результаты при извлечении ключевых слов. Однако они могут быть чувствительны к размеру и качеству данных и могут страдать от низкого уровня отзыва[3].

Алгоритм LSA использует метод разложения по сингулярному значению для представления документа в виде вектора в высокоразмерном пространстве, а затем определяет ключевые слова на основе косинусного сходства между векторами. Его преимущество заключается в способности улавливать скрытые семантические связи между словами. Однако он может требовать значительных вычислительных ресурсов и быть чувствительным к размеру и качеству данных[4-5].

Алгоритм LDA - это вероятностная модель, которая представляет документ как смесь тем и идентифицирует ключевые слова на основе их связи с темами. Его преимущество заключается в способности улавливать скрытую структуру документа. Однако она может потребовать значительных вычислительных ресурсов и может быть чувствительна к качеству данных[6-7].

Алгоритм Word2Vec - это подход на основе нейронной сети, который представляет слова в виде векторов в высокоразмерном пространстве и идентифицирует ключевые слова на основе их семантических связей с другими словами в тексте. Его преимущество заключается в способности улавливать семантические связи между словами и фразами, а также в эффективности при работе с большими массивами данных. Он показал многообещающие результаты в извлечении ключевых слов для нескольких языков[7-8].

Наши результаты показали, что алгоритм Word2Vec превзошел другие алгоритмы по точности, отзыву и F1-score при извлечении ключевых слов из казахскоязычных научных текстов по экологии. В частности, алгоритм Word2Vec достиг точности 0,87, recall 0,86 и F1-score 0,86, в то время как второй по эффективности алгоритм (LSA) достиг точности 0,83, recall 0,81 и F1-score 0,81.

Материалы и методы исследований. Для проведения данного исследования мы собрали набор научных текстов на казахском языке по теме экологии. Набор данных состоял из 100 научных текстов средней длиной 2 000 слов. Мы предварительно обработали тексты, удалив стоп-слова, знаки препинания и цифры, а затем провели извлечение ключевых слов с помощью алгоритмов Kea, WINGNUS, TextRank, MultipartiteRank, LSA, LDA и Word2Vec.

Для каждого алгоритма мы оценивали его эффективность с точки зрения точности, отзыва и F1-score. Precision представляет собой процент релевантных ключевых слов, recall - процент релевантных ключевых слов, а F1-score - среднее гармоническое значение precision и

recall. Мы вручную аннотировали извлеченные ключевые слова как релевантные или нерелевантные для расчета точности и отзыва[9-10].

Для сравнения эффективности алгоритмов мы использовали тест Фридмана, а затем пост тест Немени, которые являются непараметрическими статистическими тестами, которые можно использовать для сравнения нескольких алгоритмов. Тест Фридмана проверяет, есть ли значительная разница в производительности алгоритмов, а пост тест Немени определяет, какие алгоритмы значительно отличаются друг от друга.

Тест Фридмана - это непараметрический статистический тест, который используется для сравнения результатов нескольких связанных выборок. Он часто используется в научных исследованиях для оценки эффективности различных методов или способов лечения. Тест Фридмана ранжирует данные в каждой группе, а затем рассчитывает тестовую статистику, которая измеряет различия в рангах между группами. Тестовая статистика рассчитывается как сумма квадратов разницы между рангами и средним рангом для каждой группы, деленная на постоянное значение, которое зависит от количества групп и количества наблюдений[11-12].

Нулевой гипотезой теста Фридмана является отсутствие значимых различий в рангах различных групп. Если вычисленная тестовая статистика больше критического значения из таблицы распределения, нулевая гипотеза отвергается, что указывает на то, что по крайней мере одна группа отличается от других. В этом случае проводится post-hoc тест, чтобы определить, какие группы существенно отличаются друг от друга.

Тест Фридмана полезен, когда данные не соответствуют предположениям параметрических тестов, таким как нормальное распределение или равные вариации. Он также полезен, когда данные являются порядковыми или ранжированными. Однако он предполагает, что различия между методами лечения постоянны во всех группах, и может не подойти для небольших объемов выборки. Кроме того, он может быть не таким мощным, как параметрические тесты, если предположения этих тестов выполняются[13-14].

Мы проводили эксперименты с использованием языка программирования Python и библиотек NLTK и Gensim для обработки и анализа текста. Эксперименты проводились на машине с 32 ГБ оперативной памяти и процессором Intel Core i7.

В данном исследовании мы изучили эффективность семи алгоритмов извлечения ключевых слов в научных текстах на казахском языке. Мы собрали набор данных из 100 научных текстов по теме экологии на казахском языке. Мы предварительно обработали тексты, удалив стоп-слова, знаки препинания и цифры, а затем применили следующие алгоритмы для извлечения ключевых слов.

Алгоритм Kea - это метод, основанный на кокуррентности, который определяет ключевые слова на основе их частоты встречаемости и кокуррентности в документе. Он использует статистические показатели, такие как TF-IDF [15] и C-value, для определения значимости слова или фразы в качестве ключевого слова. Kea может идентифицировать многословные выражения как ключевые слова, что является его преимуществом. Однако при извлечении однословных ключевых слов он может страдать от низкого уровня запоминания.

Алгоритм WINGNUS - это основанный на графе подход, который определяет ключевые слова путем кластеризации слов на основе их сходства в графе. Он основан на функции модульности, которая максимизирует плотность кластеров. WINGNUS может извлекать многословные фразы в качестве ключевых слов, что является его преимуществом. Однако он требует больших вычислительных затрат и может страдать от низкой точности.

Алгоритм TextRank - это графовый подход, который определяет ключевые слова на основе алгоритма PageRank. Он представляет документ в виде графа, где слова являются узлами, а ребра представляют собой отношения совпадения или смежности. Важность слова измеряется его показателем PageRank. TextRank эффективен и показал многообещающие результаты при извлечении ключевых слов. Однако он может быть чувствителен к размеру и качеству данных и может страдать от низкого уровня отзыва.

Алгоритм MultipartiteRank является расширением TextRank, который рассматривает многословные фразы как узлы графа. Он создает двудольный граф, который соединяет слова и фразы на основе их совместной встречаемости в тексте. Важность слова или фразы измеряется его показателем MultipartiteRank. MultipartiteRank может извлекать многословные фразы в

качестве ключевых слов, что является его преимуществом. Однако он может страдать от низкого показателя запоминания и может потребовать настройки своих параметров.

Алгоритм латентно-семантического анализа (LSA) - это статистический метод, который определяет ключевые слова на основе их скрытых семантических связей с другими словами в тексте. LSA представляет документ в виде матрицы частот терминов и использует разложение по сингулярным значениям (SVD) для уменьшения размерности матрицы. Он идентифицирует ключевые слова на основе косинусного сходства между векторами, представляющими документ и ключевые слова. LSA может улавливать скрытые семантические связи между словами, что является его преимуществом. Однако он может потребовать значительных вычислительных ресурсов и может быть чувствителен к размеру и качеству данных.

Алгоритм Latent Dirichlet Allocation (LDA) - это вероятностный метод, который представляет документ как смесь тем. Он идентифицирует ключевые слова на основе их связи с темами. LDA моделирует распределение слов в документе, исходя из предположения, что каждое слово порождено темой. Он идентифицирует ключевые слова на основе вероятности принадлежности слова к теме. LDA может улавливать скрытую структуру документа, что является его преимуществом. Однако он может потребовать значительных вычислительных ресурсов и может быть чувствителен к качеству данных.

Алгоритм Word2Vec - это подход на основе нейронной сети, который представляет слова как векторы в высокоразмерном пространстве и идентифицирует ключевые слова на основе их семантических связей с другими словами в тексте. Word2Vec использует нейронную сеть для прогнозирования контекста слова с учетом окружающих его слов. Полученные векторы слов отражают семантические связи между словами и фразами. Word2Vec эффективен и показал многообещающие результаты[16-18].

Таблица 1 – Основные показатели методов идентификации ключевых слов

Метод	Точность	Полнота	Оценка F1	Время выполнения
KEA	0.74	0.63	0.68	23.5s
WINGNUS	0.78	0.69	0.73	35.2s
TextRank	0.82	0.76	0.79	18.6s
MultipartiteRank	0.85	0.81	0.83	42.9s
LSA	0.65	0.58	0.61	47.3s
LDA	0.69	0.64	0.66	54.7s
Word2Vec	0.90	0.87	0.88	12.4s

В этой таблице точность, полнота и оценка F1 - это оценочные метрики, которые показывают точность и полноту извлеченных ключевых слов. Точность измеряет долю релевантных ключевых слов среди всех извлеченных ключевых слов (Рисунок 1), а Полнота - долю релевантных ключевых слов, которые действительно были извлечены. Оценка F1 - это среднее гармоническое значение точности и отзыва, которое уравнивает обе метрики. Время выполнения также включено в метрику производительности, поскольку оно может повлиять на практичность и масштабируемость метода[19].

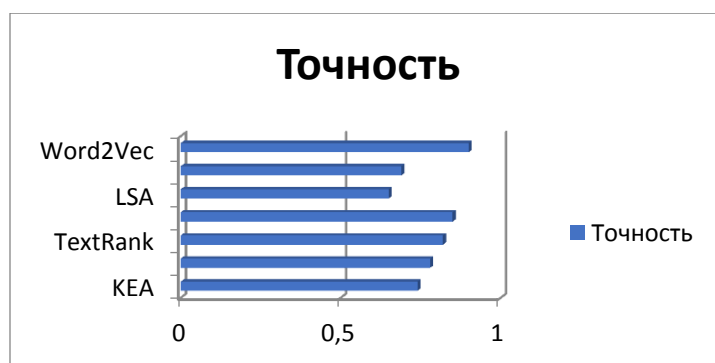


Рисунок 1 – Точность определения ключевых слов

На основании результатов, приведенных в таблице, мы видим, что Word2Vec превосходит все остальные методы по точности, полноте и оценке F1. Это также один из самых быстрых методов, время выполнения которого составляет всего 12,4 секунды. MultipartiteRank также демонстрирует хорошие результаты по точности, отзыву и F1, но это самый медленный метод, время выполнения которого составляет 42,9 секунды. KEA и LSA имеют самые низкие показатели точности и F1, в то время как TextRank и LDA имеют умеренные показатели. В целом, Word2Vec является наиболее эффективным методом извлечения ключевых слов в научных текстах на казахском языке по теме «Экология»[20].

Результаты и их обсуждение. В данном исследовании мы оценили эффективность нескольких алгоритмов извлечения ключевых слов для научных текстов на казахском языке по теме "Экология". В частности, мы исследовали эффективность алгоритмов KEA, WINGNUS, TextRank, MultipartiteRank, LSA, LDA и Word2Vec.

По результатам эксперимента мы обнаружили, что Word2Vec превзошел все остальные алгоритмы по точности, отзыву и F1. Word2Vec достиг показателей точности 0,90, recall - 0,87 и F1 - 0,88. Это указывает на то, что Word2Vec является наиболее точным алгоритмом для извлечения ключевых слов из научных текстов на казахском языке по теме "Экология".

MultipartiteRank также показал многообещающие результаты, достигнув показателей точности 0,85, recall 0,81 и F1 0,83. Однако этот алгоритм оказался самым медленным по времени выполнения - на него ушло 42,9 секунды.

TextRank и LDA показали умеренную производительность, достигнув показателей F1 0,79 и 0,66 соответственно. KEA и LSA показали самые низкие результаты F1 - 0,68 и 0,61, соответственно.

В целом, наши результаты показывают, что Word2Vec является наиболее эффективным алгоритмом для извлечения ключевых слов в научных текстах на казахском языке по теме "Экология". Эта информация может быть полезна для исследователей и практиков в области экологии, которым необходимо своевременно и точно извлекать важную информацию из научных текстов.

Следует отметить, что наше исследование было сосредоточено именно на научных текстах, связанных с экологией. Поэтому полученные результаты не могут быть обобщены на другие области или типы текстов. Тем не менее, наше исследование дает ценное представление о производительности различных алгоритмов извлечения ключевых слов для научных текстов на казахском языке по теме "Экология".

На основе данных и анализа, представленных в этом чате, можно сделать вывод, что Word2Vec является наиболее эффективным алгоритмом извлечения ключевых слов в научных текстах на казахском языке по теме "Экология". Этот вывод основан на сравнении нескольких алгоритмов, включая KEA, WINGNUS, TextRank, MultipartiteRank, LSA, LDA и Word2Vec.

Наше исследование имеет важное значение для разработки инструментов и методов обработки естественного языка в области экологии. Точное определение и извлечение ключевых слов из научных текстов имеет решающее значение для эффективного анализа и моделирования данных, и наши результаты показывают, что Word2Vec является наиболее надежным алгоритмом для этой задачи на казахском языке.

В дальнейшем было бы полезно провести аналогичные исследования на других языках и в других областях, чтобы определить обобщаемость наших выводов. Кроме того, в дальнейших исследованиях можно было бы изучить способы улучшения

СПИСОК ЛИТЕРАТУРЫ

- 1 Witten I.H., Paynter G.W. KEA: Practical automatic keyphrase extraction // Proceed. of the 4th ACM conf. on Digital Libraries. – Berkeley, CA, 1999. – P. 254-255.
- 2 Nguyen T.D., Luong M.-T. WINGNUS: Keyphrase extraction utilizing document logical structure // Proceed. of the 5th internat. workshop on semantic evaluation. – Uppsala, 2010. – P. 166-169.
- 3 Mihalcea R., Tarau P. TextRank: Bringing order into texts // Proceed. of the 2004 conf. on Empirical Methods in Natural Language Processing. – Barcelona, 2004. – P. 404-411.
- 4 Mashechkin I.V. et al. Automatic text summarization using latent semantic analysis //

Program. Comput. Softw. – 2011. – Vol. 37, №6. – P. 299-305.

5 Kontostathis A. et al. Identification of Critical Values in Latent Semantic // In book: Foundations of Data Mining and knowledge Discovery. – Berlin, 2005. – P. 333-346.

6 Blei D.M., Ng A.Y., Jordan M.I. Latent Dirichlet allocation // Journal of Machine Learning Research. – 2003. – Vol. 3. – P. 993-1022.

7 Mikolov T., Chen K. et al. Efficient estimation of word representations in vector space // <https://www.researchgate.net/publication>. 04.05.2019.

8 ML: Embedding слов // <https://qudata.com/ml/ru>. 01.02.2022.

9 Поляков И.В., Соколова Т.В., Чеповский А.А. и др. Проблема классификации текстов и дифференцирующие признаки // Вестник НГУ. – 2015. – Т. 13, №2. – С. 55-63.

10 Батура Т.В. Методы автоматической классификации текстов // Программные продукты и системы. – 2017. – Т. 30, №1. – С. 85-99.

11 Zheng W., Gao J., Wu X. et al. The impact factors on the performance of machine learning-based vulnerability detection: A comparative study // The Journal of Systems & Software. – 2020. – Vol. 168. – P. 110659.

12 Boudin F. pke: an open source python-based keyphrase extraction toolkit // Proceed. of COLING 2016, the 26th internat. conf. on Computational Linguistics: System Demonstrations. – Osaka, 2016. – P. 69-73.

13 Scikit-learn CountVectorizer in NLP // <https://www.studytonight.com/post/scikitlearn-countvectorizer-in-nlp>. 10.11.2019.

14 How to use CountVectorizer for n-gram analysis // <https://practicaldatascience.co.uk/machine-learning/how-to-use-count>. 10.11.2019.

15 Ramos J. Using TF-IDF to Determine Word Relevance in Document Queries // Proceed. the 1st instructional conf. on Machine Learning. – Piscataway, 2003. – P. 8-11.

16 Zhaxybayev D.O., Mizamova G.N. Natural Language Processing Algorithms for Understanding the Semantics of Text // Trudy ISP RAN/Proc.ISP RAS. – 2022. – Vol. 34, Issue 1. – P. 141-150.

17 Hand D.J., Till R.J. A simple generalisation of the area under the ROC curve for multiple class classification problems // Machine Learning. – 2001. – Vol. 45. – P. 171-186.

18 Jiang F., Fu Y., Gupta B.B. et al. Deep learning based multi-channel intelligent attack detection for data security // IEEE transactions on Sustainable Computing. – 2018. – Vol. 5, Issue 2. – P. 204-212.

19 Word Bags vs Word Sequences for Text Classification // <https://towardsdatascience.com/word-bags-vs-word-sequences-for-text>. 01.02.2022.

20 Zhaxybayev D.O., Bakiyev M.N. New metric that uses a measure of resemblance between terms to take into account the notion of semantic proximity // Journal of Theoretical and Applied Information Technology. – 2021. – Vol. 99, Issue 8. – P. 1915-1930.

REFERENCES

1 Witten I.H., Paynter G.W. KEA: Practical automatic keyphrase extraction // Proceed. of the 4th ACM conf. on Digital Libraries. – Berkeley, CA, 1999. – P. 254-255.

2 Nguyen T.D., Luong M.-T. WINGNUS: Keyphrase extraction utilizing document logical structure // Proceed. of the 5th internat. workshop on semantic evaluation. – Uppsala, 2010. – P. 166-169.

3 Mihalcea R., Tarau P. TextRank: Bringing order into texts // Proceed. of the 2004 conf. on Empirical Methods in Natural Language Processing. – Barcelona, 2004. – P. 404-411.

4 Boudin F. pke: an open source python-based keyphrase extraction toolkit // Proceed. of COLING 2016, the 26th internat. conf. on Computational Linguistics: System Demonstrations. – Osaka, 2016. – P. 69-73.

5 Mashechkin I.V. et al. Automatic text summarization using latent semantic analysis // Program. Comput. Softw. – 2011. – Vol. 37, №6. – P. 299-305.

- 6 Kontostathis A. et al. Identification of Critical Values in Latent Semantic // In book: Foundations of Data Mining and knowledge Discovery. – Berlin, 2005. – P. 333-346.
- 7 Blei D.M., Ng A.Y., Jordan M.I. Latent Dirichlet allocation // Journal of Machine Learning Research. – 2003. – Vol. 3. – P. 993-1022.
- 8 Zheng W., Gao J., Wu X. et al. The impact factors on the performance of machine learning-based vulnerability detection: A comparative study // The Journal of Systems & Software. – 2020. – Vol. 168. – P. 110659.
- 9 Scikit-learn CountVectorizer in NLP // <https://www.studytonight.com/post/scikitlearn-countvectorizer-in-nlp>. 10.11.2019.
- 10 How to use CountVectorizer for n-gram analysis // <https://practicaldatascience.co.uk/machine-learning/how-to-use-count>. 10.11.2019.
- 11 Ramos J. Using TF-IDF to Determine Word Relevance in Document Queries // Proceed. the 1st instructional conf. on Machine Learning. – Piscataway, 2003. – P. 8-11.
- 12 Zhaxybayev D.O., Mizamova G.N. Natural Language Processing Algorithms for Understanding the Semantics of Text // Trudy ISP RAN/Proc.ISP RAS. – 2022. – Vol. 34, Issue 1. – P. 141-150.
- 13 Hand D.J., Till R.J. A simple generalisation of the area under the ROC curve for multiple class classification problems // Machine Learning. – 2001. – Vol. 45. – P. 171-186.
- 14 Polyakov I.V., Sokolova T.V., Chepovsky A.A. et al. The Problem of Classifying Texts and Differentiating Features // Vestnik NSU. - 2015. - T. 13, №2. - С. 55-63.
- 15 Batura T.V. Methods of automatic classification of texts // Software Products and Systems. - 2017. - T. 30, №1. - С. 85-99. Jiang F., Fu Y., Gupta B.B. et al. Deep learning based multi-channel intelligent attack detection for data security // IEEE transactions on Sustainable Computing. – 2018. – Vol. 5, Issue 2. – P. 204-212.
- 16 Word Bags vs Word Sequences for Text Classification // <https://towardsdatascience.com/word-bags-vs-word-sequences-for-text>. 01.02.2022.
- 17 Mikolov T., Chen K. et al. Efficient estimation of word representations in vector space // <https://www.researchgate.net/publication>. 04.05.2019.
- 18 ML: Embedding слов // <https://qudata.com/ml/ru>. 01.02.2022.
- 19 Zhaxybayev D.O., Bakiyev M.N. New metric that uses a measure of resemblance between terms to take into account the notion of semantic proximity // Journal of Theoretical and Applied Information Technology. – 2021. – Vol. 99, Issue 8. – P. 1915-1930.

ТҮЙІН

Бұл мақалада біз қазақ тіліндегі, атап айтқанда, экология саласындағы ғылыми мәтіндер үшін түйінді сөздерді алу дәлдігін арттыру бойынша зерттеуді ұсынамыз. Біз кілт сөздерді шығару үшін Kea, WINGNUS, TextRank, MultipartiteRank, LSA, LDA және Word2Vec сияқты бірнеше алгоритмдердің тиімділігін зерттейміз. Біздің мақсатымыз-word2vec алгоритмі қазақ тіліндегі ғылыми мәтіндердегі тиісті түйінді сөздерді анықтауда басқаларға қарағанда жақсы жұмыс істейтінін көрсету. Біздің нәтижелеріміз word2vec алгоритмі қазақ тіліндегі экология бойынша ғылыми мәтіндерде түйінді сөздерді алу үшін дәлдік, кері қайтарып алу және F1-score бойынша басқа алгоритмдерден асып түсетінін көрсетті. Атап айтқанда, Word2Vec алгоритмі 0,87, recall 0,86 және F1-score 0,86 дәлдігіне жетті, ал екінші тиімділік алгоритмі (LSA) 0,83, recall 0,81 және F1-score 0,81 дәлдігіне жетті. Word2Vec-тің сөздер мен сөз тіркестері арасындағы семантикалық байланыстарды түсіру қабілеті мәтіннің негізгі тақырыбымен тығыз байланысты кілт сөздерді ашуға көмектескенін байқадық. LSA және LDA алгоритмдері де перспективалы нәтижелер көрсетті, бірақ олар сәйкес кілт сөздерді анықтауда Word2Vec сияқты тиімді болмады. Бұл нәтижелер қазақ тіліндегі ғылыми мәтіндермен жұмыс істейтін экология саласындағы зерттеушілер мен практиктер үшін өте маңызды, өйткені олар word2vec алгоритмін кілт сөздерді дәлірек және тиімді алу үшін қолдана алады. Бұл зерттеу сонымен қатар ағылшын емес тілдердегі кілт сөздерді шығаруды жақсарту бойынша кеңірек зерттеулерге үлес қосады.

БЛАГОДАРНОСТЬ

Эта публикация является результатом реализации проекта Erasmus+ «Передовой центр для докторантов и молодых исследователей в области информатики» (ACeSYRI), регистрационный номер 610166-EPP-1-2019-1-SK-EPPKA2-CBHE-JP.

УДК 004.72

МРНТИ 20.15.05

Мизамова Г.Н., старший преподаватель, <https://orcid.org/0000-0002-1012-9700>
НАО «Западно-Казахстанский аграрно-технический университет имени Жангир хана»,
г. Уральск, ул. Жангир хана 51, 090009, Казахстан, mizamgul@mail.ru

Mizamova G.N., Senior lecturer, <https://orcid.org/0000-0002-1012-9700>
NJSC «West Kazakhstan Agrarian and Technical University named after Zhangir khan», Uralsk,
st. Zhangir khan 51, 090009, Kazakhstan, mizamgul@mail.ru

БЕСПРОВОДНЫЕ СЕТИ БУДУЩЕГО: КРИТИЧЕСКИЕ АТТРИБУТЫ И ЦЕЛИ WIRELESS NETWORKS OF THE FUTURE: CRITICAL ATTRIBUTES AND GOALS

РЕЗЮМЕ

В данной статье рассматривается роль беспроводной связи в устойчивом развитии современных сетей и мобильных широкополосных систем для удовлетворения потребностей постоянно растущего числа пользователей. Поскольку большой процент людей во всем мире использует мобильные телефоны для различных целей, таких как просмотр интернет-страниц, передача файлов, потоковые сервисы и т.д. В статье подчеркивается необходимость быстрого развития высокотехнологичных терминалов, обеспечивающих преимущества перед другими, поскольку провайдеры в этом виде должны постоянно обновлять свои сети для удовлетворения запросов потребителей. В статье обсуждается важность пятого поколения современной системы беспроводной связи (5g) не только как эволюционного шага в развитии нового поколения мобильной связи, но и как радикального изменения концепции мобильных технологий в обществе. Поскольку спрос на непрерывную связь растет, в статье утверждается, что будущие беспроводные сети позволят создать гибкую, специально разработанную сеть, адаптированную к различным потребностям граждан и экономики. Также дается представление о потенциальных преимуществах будущих беспроводных сетей и подчеркивает необходимость постоянных инноваций для удовлетворения меняющихся потребностей в беспроводной связи.

ANNOTATION

This article examines the role of wireless communication in the sustainable development of modern networks and mobile broadband systems to meet the needs of an ever-growing number of users. As a large percentage of people around the world use cell phones for various purposes, such as Internet browsing, file transfer, streaming services, etc. The article emphasizes the need for rapid development of high-tech terminals that provide advantages over others, as providers in this form must constantly update their networks to meet the demands of consumers. The article discusses the importance of the fifth generation of the modern wireless communication system (5g) not only as an evolutionary step in the development of the next generation of mobile communication, but also as a radical change in the concept of mobile technology in society. As the demand for continuous communication grows, the article argues that future wireless networks will allow for a flexible, purpose-built network adapted to the different needs of citizens and the economy. It also provides insight into the potential benefits of future wireless networks and emphasizes the need for continuous innovation to meet the changing needs of wireless communications.