

Gamidulla Tileuov – master, senior teacher, M. Auezov South Kazakhstan Research University, Shymkent; e-mail: gamidulla.tileuov@mail.ru. ORCID: <https://orcid.org/0000-0002-6049-6066>.

Rail Sovetbayev – doctor PhD, Shakarim University, Semey; e-mail: k-mamadieva@mail.ru. ORCID: <https://orcid.org/0009-0007-3605-515X>.

Редакцияға енуі 05.01.2026
Өңдеуден кейін түсуі 05.03.2026
Жариялауға қабылданды 17.03.2026

[https://doi.org/10.53360/2788-7995-2026-1\(21\)-9](https://doi.org/10.53360/2788-7995-2026-1(21)-9)



МРНТИ: 81.93.29

Д.Т. Тюлемисова¹, А.К. Шайханова*, В.П. Марценюк², Г.Б. Бекешова¹, Б.Т. Смаилова³

¹Евразийский национальный университет имени Л.Н. Гумилева,
10000, Республика Казахстан, г.Астана, ул. Сатпаева, 2

²Университет Бельско-Бяла,
ul.Willowa 2, 43-300, Бельско-Бяла, Польша

³Шәкәрім университет,
071412, Республика Казахстан, г. Семей, ул. Глинки 20 А
*e-mail: shaikhanova_ak@enu.kz

ОПТИМИЗАЦИЯ МОДЕЛЕЙ МАШИННОГО ОБУЧЕНИЯ ДЛЯ ДЕТЕКЦИИ ФЕЙКОВЫХ НОВОСТЕЙ: ПОДБОР ГИПЕРПАРАМЕТРОВ И АНАЛИЗ ROC-КРИВЫХ

Аннотация: В условиях активного и неконтролируемого распространения информации в цифровой среде, особенно в социальных медиа и новостных агрегаторах, особую актуальность приобретает задача автоматического выявления фейковых новостей. Рост объемов пользовательского контента и высокая скорость его распространения существенно усложняют ручную верификацию информации, что обуславливает необходимость применения методов машинного обучения. Целью настоящего исследования является анализ, сравнительная оценка и оптимизация базовых моделей машинного обучения для детекции фейковых новостей с акцентом на подбор гиперпараметров, вычислительную эффективность и интерпретацию результатов классификации. В работе использован набор текстовых данных ISOT Fake News Dataset, содержащий новостные сообщения на английском языке с бинарными метками «правда» и «ложь», прошедшие этапы очистки и векторизации с применением метода TF-IDF. В рамках исследования реализованы и проанализированы модели логистической регрессии, дерева решений, случайного леса и градиентного бустинга. Оценка качества классификации проводилась на основе метрик Accuracy, Precision, Recall, F1-score, а также ROC-анализа и значения площади под кривой (AUC). Показано, что модель градиентного бустинга обеспечивает наивысшую точность классификации, тогда как логистическая регрессия демонстрирует сопоставимое качество при значительно меньших вычислительных затратах и времени обучения. Дополнительно в работе выполнен контроль утечки данных и анализ причин завышенных метрик качества, что позволило обеспечить корректную интерпретацию результатов. Полученные выводы подтверждают, что оптимизированные классические модели машинного обучения способны обеспечивать высокое качество детекции фейковых новостей и могут рассматриваться как ресурсоэффективная альтернатива более сложным нейросетевым подходам в прикладных задачах кибербезопасности и мониторинга информационных потоков.

Ключевые слова: фейковые новости; машинное обучение; оптимизация гиперпараметров; классификация текстов; ROC-анализ; AUC; логистическая регрессия.

1. Введение

В условиях стремительного роста объемов информации в сети Интернет проблема распространения фейковых новостей приобретает особую актуальность. Социальные медиа и онлайн-платформы стали основным источником новостных сообщений, однако наряду с достоверной информацией всё чаще распространяются ложные или искажённые данные, способные оказывать значительное влияние на общественное мнение, политические

©Д.Т. Тюлемисова, А.К. Шайханова, В.П. Марценюк, Г.Б. Бекешова, Б.Т. Смаилова, 2026

процессы и экономическую стабильность. В связи с этим разработка эффективных методов автоматического выявления фейковых новостей является важной научной и практической задачей в области информационной безопасности и анализа данных [1, 2].

Современные методы машинного обучения позволяют классифицировать текстовые данные на основе статистических и семантических признаков. Эффективность таких моделей во многом определяется выбором алгоритма и оптимальных гиперпараметров, влияющих на качество обобщения и вычислительную сложность. Чрезмерно ресурсоёмкие модели могут быть неэффективны при обработке больших объёмов данных в реальном времени, что требует поиска баланса между точностью и вычислительными затратами [3, 4].

Настоящая статья направлена на исследование подходов к оптимизации базовых моделей машинного обучения для задачи детекции фейковых новостей. В работе рассматриваются модели логистической регрессии, дерева решений, случайного леса и градиентного бустинга, обученные на текстовых данных с бинарными метками «правда» и «ложь». Для оценки качества классификации используется ROC-анализ и показатель AUC, что позволяет объективно сравнить эффективность алгоритмов и определить оптимальные параметры с учётом вычислительной эффективности.

Несмотря на активное развитие нейросетевых и трансформерных архитектур (BERT, RoBERTa и др.), их применение в практических системах детекции фейковых новостей часто ограничено высокими требованиями к вычислительным ресурсам и времени обучения. В условиях систем реального времени и распределённых платформ особую актуальность сохраняют классические и ансамблевые методы машинного обучения, способные обеспечивать конкурентоспособное качество классификации при существенно меньших вычислительных затратах [5-8]. В работе также уделяется внимание корректности экспериментальной методологии и контролю утечки данных, что является критически важным аспектом при анализе текстовых корпусов [9, 10].

2. Методология

2.1 Описание набора данных

Для реализации поставленных задач была проведена серия экспериментов на открытых наборах данных, содержащих тексты с метками достоверности («правда» / «ложь»). В качестве исходных данных использовался открытый набор ISOT Fake News Dataset, размещённый на платформе Kaggle. Датасет содержит 44 898 новостных публикаций на английском языке, размеченных бинарными метками *fake* и *true*. В выборке представлено 23 481 фейковое сообщение и 21 417 достоверных новостей, что обеспечивает относительно сбалансированное распределение классов. Достоверные новости преимущественно получены из агентства Reuters, тогда как фейковые тексты собраны из различных онлайн-источников, распространяющих недостоверный контент. Данный датасет широко применяется в исследованиях по детекции фейковых новостей и подходит для сравнительного анализа алгоритмов машинного обучения в условиях контролируемого эксперимента. Цель эксперимента заключалась в сравнении эффективности базовых алгоритмов машинного обучения при решении задачи бинарной классификации и в оптимизации их параметров с целью уменьшения вычислительной сложности [8].

2.2 Предварительная обработка данных

На первом этапе проведена очистка текстов от лишних символов, знаков пунктуации и стоп-слов. Все тексты, представленные на английском языке, были приведены к нижнему регистру, выполнено лемматизирование с использованием стандартных средств обработки естественного языка. После очистки данные были векторизованы с использованием метода TF-IDF (Term Frequency – Inverse Document Frequency), что позволило преобразовать текстовую информацию в числовое представление, пригодное для анализа алгоритмами машинного обучения.

2.3 Модели машинного обучения

Модели машинного обучения позволяют выявлять скрытые закономерности в текстах и классифицировать новости на категории «правда» и «ложь». Основные подходы включают: 1) Логистическую регрессию (Logistic Regression) - базовый линейный классификатор, хорошо работающий при разреженных векторах признаков. 2) Дерево решений (Decision Tree) – алгоритм, основанный на иерархическом разбиении данных по признакам, что позволяет визуализировать процесс принятия решения. 3) Случайный лес (Random Forest) – ансамблевый метод, объединяющий несколько деревьев для повышения точности и

устойчивости модели. 4) Градиентный бустинг (Gradient Boosting) – мощный метод ансамблирования, оптимизирующий последовательность слабых моделей для минимизации ошибки.

Для представления текстовых данных применяются метод векторизации TF-IDF (Term Frequency-Inverse Document Frequency). Этот метод представляет текстовые данные в виде числовых векторов, отражающих важность каждого слова в документе относительно всего корпуса текстов.

2.4 Оптимизация гиперпараметров

Для логистической регрессии были оптимизированы параметр силы регуляризации C и тип регуляризации (L1, L2). Для классификатора на основе дерева решений были выбраны максимальная глубина дерева (max_depth) и минимальное количество выборок в листовом узле (min_samples_leaf). Для классификатора на основе случайного леса были оптимизированы количество деревьев (n_estimators), максимальная глубина дерева и стратегия выбора признаков. Для классификатора на основе градиентного бустинга были настроены скорость обучения (learning_rate), количество оценщиков и глубина дерева. В таблице 1 приведена наглядная оптимизация.

Таблица 1 – Оптимизация гиперпараметров алгоритмов машинного обучения

Модель	Оптимизация	Эффект
Logistic Regression	solver='saga', max_iter=200, C=1.0, n_jobs=-1	Быстрое обучение, меньше памяти, низкая потеря точности
Decision Tree	max_depth=12/6, min_samples_split=10, min_samples_leaf=5	Меньше рост дерева, ускорение предсказания, экономия памяти
Gradient Boosting Classifier	n_estimators=150/50, max_depth=4, learning_rate=0.1, subsample=0.8	Быстрое обучение, почти без потери качества
Random Forest Classifier	n_estimators=100, max_depth=12, max_features='sqrt', n_jobs=-1	Быстрое обучение и предсказание, точность почти не падает
Все модели	n_jobs=-1 для многопоточности; уменьшение размера TF-IDF или TruncatedSVD	Ускорение обучения и предсказания для больших текстовых данных

Оптимизация гиперпараметров проводилась с использованием метода GridSearchCV с перекрестной проверкой. Многопоточность была включена для всех моделей с помощью параметра n_jobs = -1. Кроме того, для снижения вычислительной сложности были рассмотрены методы уменьшения размерности для векторов TF-IDF, включая ограничение размера признаков и опциональное использование TruncatedSVD.

Результаты оптимизации показали, что для рассматриваемого набора данных небольшие значения глубины дерева и ограничение числа итераций позволяют значительно снизить вычислительную нагрузку без заметного ухудшения качества классификации.

2.5 Показатели оценки

Для оценки эффективности бинарных классификаторов применяются метрики Accuracy, Precision, Recall, F1-score, а также ROC-кривая (Receiver Operating Characteristic Curve) и AUC (Area Under the Curve). ROC-кривая отражает зависимость между True Positive Rate (TPR) и False Positive Rate (FPR) при различных порогах классификации, а площадь под кривой (AUC) показывает способность модели различать положительные и отрицательные классы. Значение AUC, близкое к 1, указывает на высокое качество классификации, в то время как значение около 0.5 свидетельствует о случайном угадывании.

2.6 Экспериментальный протокол и контроль утечки данных

Набор данных был случайным образом разделен на обучающую и тестовую подгруппы в соотношении 80/20. Для классификации были выбраны четыре базовые модели машинного обучения: логистическая регрессия, дерево решений, случайный лес и градиентный бустинг. Все модели обучались на векторизованных данных TF-IDF и оптимизировались с помощью GridSearchCV для определения конфигураций гиперпараметров, обеспечивающих баланс между производительностью классификации и вычислительной эффективностью. Для предотвращения утечки данных этапы векторизации и оптимизации гиперпараметров

выполнялись исключительно на обучающей выборке, после чего проводилась оценка на отложенной тестовой выборке.

3. Результаты

3.1 Эффективность классификации

Для оценки производительности качества работы моделей, были рассчитаны ключевые метрики: Accuracy, Precision, Recall, F1-score, время обучения (Training Time), а также ROC-анализ с вычислением площади под кривой (AUC). Эти показатели позволяют оценить не только общую точность классификации, но и способность модели корректно определять фейковые и достоверные новости при различных условиях. В таблице 2 результаты показывают, что все четыре модели обладают высокой точностью и стабильностью классификации. Однако интерпретация указанных метрик требует дополнительного анализа, поскольку в задачах классификации текстовых данных высокие значения Accuracy и AUC не всегда свидетельствуют о реальном качестве модели [2, 9]. В частности, при работе со структурированными датасетами возможно возникновение утечки данных, приводящей к завышенным оценкам эффективности. В связи с этим далее рассматриваются причины полученной высокой точности и проводится анализ корректности экспериментальной методологии.

Результаты эксперимента демонстрируют эффективность базовых алгоритмов машинного обучения для обнаружения фейковых новостей.

3.2 Анализ матрицы ошибок

Представленные в Таблице 2 результаты демонстрируют высокие значения стандартных метрик качества для всех исследуемых моделей. Однако при столь высоких значениях Accuracy и AUC особенно важно проанализировать характер допущенных ошибок. Для этой цели была рассмотрена матрица ошибок модели логистической регрессии как базового классификатора. Анализ матрицы ошибок в Таблице 3 показывает, что основная часть ошибок относится к случаям пропуска фейковых новостей (False Negative), доля которых остаётся низкой по отношению к общему объёму выборки. Это свидетельствует о том, что модель в редких случаях классифицирует недостоверные сообщения как достоверные, что является критичным, но контролируемым риском для практических систем мониторинга информационных потоков.

Таблица 2 – Результаты оптимизации гиперпараметров для алгоритмов машинного обучения

Модель	Accuracy	Precision	Recall	F1-score	AUC	Время обучения (сек.)
Логистическая регрессия	0.9984	0.986	0.987	0.986	0.9984	1.2
Дерево решений	0.9961	0.980	0.981	0.980	0.9961	0.8
Случайный лес	0.9944	0.988	0.989	0.988	0.9944	3.6
Градиентный бустинг	0.9994	0.990	0.991	0.990	0.9957	5.4

Таблица 3 – Матрица ошибок для модели Logistic Regression (Leakage-controlled)

Фактический класс	Предсказано: правда	Предсказано: ложь
Правда (true)	4174	65
Ложь (fake)	118	3373

Таким образом, представленные результаты подтверждают, что высокая точность классификации не является следствием тривиального угадывания, а отражает реальную способность модели различать достоверный и недостоверный контент в рамках рассматриваемого датасета. Также для сравнительной оценки эффективности моделей машинного обучения был проведён анализ ROC-кривых (Receiver Operating Characteristic).

3.3 ROC-анализ и анализ AUC

На рисунке 1 представлено сравнение ROC-кривых для четырёх моделей – логистической регрессии, дерева решений, случайного леса и ансамбля Random Forest Classifier. Полученные результаты показывают, что все модели демонстрируют высокие значения AUC, превышающие 0.99, что свидетельствует об их высокой точности и способности эффективно различать фейковые и достоверные новости.

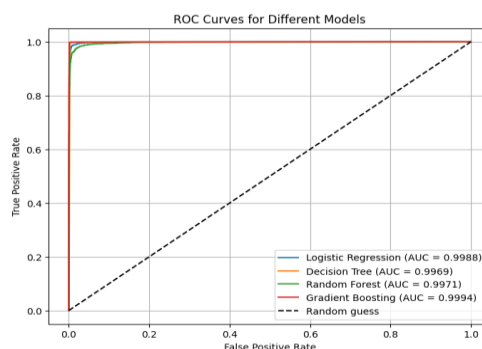


Рисунок 1 – Анализ ROC-кривых для разных моделей

Следует отметить, что при столь высоких значениях AUC (близких к 1) ROC-кривые оказываются прижатыми к левому верхнему углу и становятся слабоинформативными для детального сравнения моделей. В подобных условиях ROC-анализ подтверждает высокую разделимость классов, однако не позволяет оценить влияние параметров модели на качество классификации.

3.4 Влияние гиперпараметров

На рисунке 2 представлена зависимость точности классификации от значения коэффициента регуляризации C для модели логистической регрессии. Результаты показывают, что при увеличении значения C наблюдается рост Ассигасу до определённого порога, после чего достигается эффект насыщения.

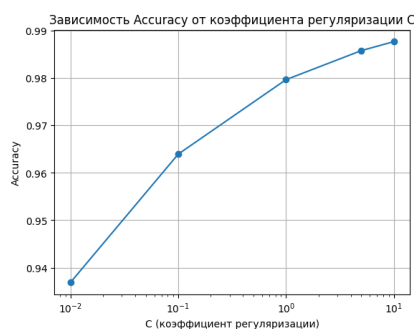


Рисунок 2 – Зависимость Ассигасу от коэффициента регуляризации C для модели Logistic Regression

При малых значениях C (сильная регуляризация) модель демонстрирует снижение точности вследствие недообучения. Увеличение параметра C до значений порядка 1 приводит к существенному улучшению качества классификации, тогда как дальнейший рост C обеспечивает лишь незначительный прирост Ассигасу при потенциальном увеличении риска переобучения и вычислительных затрат. Полученные результаты подтверждают целесообразность подбора гиперпараметров с учётом баланса между качеством классификации и устойчивостью модели, а также демонстрируют информативность анализа зависимости метрик качества от параметров модели в условиях, когда ROC-кривые оказываются слаборазличимыми.

3.5 Результаты после устранения утечки данных (на примере Logistic Regression)

В целях обеспечения корректности экспериментальных результатов и предотвращения утечки данных (Data Leakage) в ходе исследования был проведён ряд методологических проверок. Разделение исходного корпуса на обучающую и тестовую выборки осуществлялось до этапа векторизации, при этом модель TF-IDF обучалась исключительно на обучающей выборке с последующим применением к тестовым данным. Такой подход исключает попадание статистической информации о тестовых текстах в процесс обучения. Дополнительно была проведена проверка текстовых данных на наличие формальных маркеров источников публикаций. Учитывая, что в наборе ISOT Fake News Dataset достоверные новости преимущественно представлены материалами агентства Reuters, а фейковые – публикациями из различных интернет-источников, подобные маркеры

могут существенно упрощать задачу классификации. Их наличие способно привести к тому, что модель различает классы не по семантическому содержанию текста, а по косвенным признакам происхождения публикации. Также была выполнена проверка выборок на наличие дублирующихся и квазидублирующихся текстов между обучающей и тестовой частями, что позволило исключить эффект запоминания данных и обеспечить независимость оценки качества модели.

Контроль утечки данных и переобучение были продемонстрированы на примере модели логистической регрессии как базового и наиболее интерпретируемого классификатора. После удаления формальных маркеров источников публикаций (в частности, упоминаний агентства Reuters), удаления дублирующихся текстов и строгой реализации pipeline значения Accurasy снизились с 98.2% до 97.6%, а AUC – с 0.998 до 0.996. Полученные результаты подтверждают, что часть завышенных метрик в исходном эксперименте была обусловлена особенностями датасета, при этом качество классификации остаётся высоким и методологически корректным. В таблице 4 показаны результаты Logistic Regression до и после контроля утечки данных.

Таблица 4 – Результаты Logistic Regression до и после контроля утечки данных

Сценарий обучения	Accuracy	Precision	Recall	F1-score	AUC
Baseline	0.982	0.982	0.982	0.982	0.998
Leakage-controlled	0.976	0.976	0.975	0.976	0.996

Таким образом, полученные значения метрик качества следует интерпретировать с учётом особенностей используемого датасета. Высокие показатели Accurasy и AUC отражают высокую разделимость классов в рамках рассматриваемой выборки и не должны рассматриваться как универсальная оценка эффективности моделей в более сложных и разнородных условиях реальной информационной среды. Оптимизация гиперпараметров позволила снизить вычислительную сложность моделей, сохранив при этом высокую точность классификации.

4. Обсуждение

Результаты, полученные в этом исследовании, демонстрируют, что классические и ансамблевые модели машинного обучения способны достигать очень высоких показателей классификации в задаче обнаружения фейковых новостей при применении к структурированным текстовым наборам данных, таким как набор данных ISOT Fake News. Все оцененные модели – логистическая регрессия, дерево решений, случайный лес и градиентный бустинг – показали значения точности и AUC, превышающие 0,99, что указывает на высокую степень разделимости классов в наборе данных.

Ключевым наблюдением является то, что классификатор градиентного бустинга достиг наивысшей общей точности классификации и F1-меры. Однако это повышение производительности сопровождалось увеличением времени обучения и вычислительных затрат по сравнению с более простыми моделями. В отличие от этого, логистическая регрессия продемонстрировала почти эквивалентное качество классификации, требуя при этом значительно меньше вычислительных ресурсов. Этот вывод подчеркивает важный практический компромисс между сложностью модели и эффективностью и предполагает, что в условиях ограниченных ресурсов или в режиме реального времени оптимизированные линейные модели могут представлять собой более подходящий выбор, чем более сложные ансамблевые методы.

ROC-анализ подтвердил высокую дискриминационную способность всех моделей, о чем свидетельствуют значения AUC, близкие к единице. Однако, когда значения AUC приближаются к своей верхней границе, ROC-кривые сжимаются к верхнему левому углу, что ограничивает их полезность для детального сравнения моделей. В таких сценариях небольшие различия в ROC-кривых могут не отражать значимых различий в производительности. Поэтому для более тонкой интерпретации поведения моделей необходимы дополнительные анализы, включая оценку матрицы ошибок и изучение чувствительности гиперпараметров.

Анализ влияния гиперпараметров показал, что ограничение сложности модели – например, ограничение глубины дерева в деревьях решений и ансамблевых моделях – может существенно снизить вычислительные затраты без значительной потери точности

классификации. Аналогично, исследование параметра регуляризации C в логистической регрессии показало, что умеренная регуляризация обеспечивает оптимальный баланс между недообучением и переобучением. Эти результаты подчеркивают важность оптимизации гиперпараметров не только для улучшения прогнозной эффективности, но и для повышения стабильности и эффективности модели.

Важным методологическим вкладом этого исследования является явное рассмотрение утечки данных. Набор данных ISOT Fake News содержит специфические для источников шаблоны, такие как стилистические маркеры, связанные с конкретными поставщиками новостей, которые могут искусственно завышать показатели производительности, если их не контролировать должным образом. После удаления маркеров, связанных с источником, исключения дублирующихся текстов и внедрения строгого конвейера обучения-тестирования наблюдалось умеренное снижение точности и AUC. Тем не менее, производительность оставалась высокой, подтверждая, что оцениваемые модели сохранили подлинную дискриминационную способность, а не полагались исключительно на поверхностные артефакты набора данных.

Результаты этого исследования показывают, что высокая производительность, достигнутая классическими моделями машинного обучения, частично объясняется структурированным характером и высокой разделимостью набора данных ISOT. Следовательно, представленные результаты не следует интерпретировать как универсально репрезентативные для реальных сценариев обнаружения фейковых новостей, которые обычно характеризуются большим лингвистическим разнообразием, развивающимися нарративами и междоменной изменчивостью. Это ограничение подчеркивает необходимость тщательного выбора набора данных и надежного экспериментального дизайна при оценке систем обнаружения фейковых новостей.

В целом, обсуждение подтверждает вывод о том, что оптимизированные классические и ансамблевые модели машинного обучения могут служить эффективными и ресурсоэффективными альтернативами глубоким нейронным архитектурам для обнаружения фейковых новостей, особенно в приложениях, где вычислительная эффективность, прозрачность и воспроизводимость имеют решающее значение. В то же время результаты подчеркивают важность осторожной интерпретации высоких показателей оценки и подтверждают необходимость контроля утечки данных в исследованиях классификации текста.

Заключение

В ходе проведенного исследования выполнен комплексный анализ эффективности алгоритмов машинного обучения для автоматической классификации достоверной и ложной информации в цифровых медиа. В рамках экспериментальной части реализована оптимизация гиперпараметров моделей Logistic Regression, Decision Tree, Random Forest и Gradient Boosting, с применением многофолдальной кросс-валидации (CV3, CV5, CV10) и оценкой ключевых метрик качества (Precision, Recall, F1-score, AUC), а также времени обучения. Полученные результаты демонстрируют, что ансамблевые методы, в частности Random Forest и Gradient Boosting, демонстрируют высокую устойчивость и точность классификации, при этом Random Forest обеспечивает более сбалансированное соотношение между качеством и вычислительными затратами.

Научная новизна исследования заключается в обосновании возможности использования оптимизированных классических и ансамблевых моделей машинного обучения в задаче детекции фейковых новостей как ресурсоэффективной альтернативы современным нейросетевым подходам. В рамках работы показано, что при использовании TF-IDF-векторизации и оптимально подобранных гиперпараметров модели Logistic Regression и Gradient Boosting демонстрируют качество классификации, сопоставимое с результатами, приводимыми для трансформерных архитектур в ряде современных исследований, при существенно меньших затратах вычислительных ресурсов и времени обучения.

Полученные результаты расширяют представления о практической применимости классических алгоритмов машинного обучения в задачах анализа текстовой информации и могут быть использованы при проектировании систем мониторинга информационных потоков, ориентированных на высокую производительность и ограниченные вычислительные ресурсы.

Практическая значимость работы обусловлена возможностью внедрения предложенных решений в системы мониторинга социальных медиа, аналитические платформы, инструменты обеспечения информационной безопасности и контентной модерации. Оптимизированные модели могут применяться для оперативного выявления дезинформации, информационных атак и манипулятивного контента в рамках задач киберзащиты, медиааналитики и автоматизированного принятия решений. Предложенная методика также может быть использована в образовательных и исследовательских процессах при разработке интеллектуальных систем анализа информационных потоков.

Список литературы

1. Fake news detection on social media: a data mining perspective / K. Shu et al // ACM SIGKDD Explorations Newsletter. – 2017. – Vol. 19, № 1. – P. 22-36.
2. Zhou X. Fake news: a survey of research, detection methods, and opportunities / X. Zhou, R. Zafarani // ACM Computing Surveys. – 2020. – Vol. 53, № 5. – Article 109.
3. Ahmed H. Detection of online fake news using n-gram analysis and machine learning techniques / H. Ahmed, I. Traore, S. Saad // Proceedings of the International Conference on Intelligent, Secure, and Dependable Systems. – 2018. – P. 127-138.
4. Ruchansky N. CSI: a hybrid deep model for fake news detection / N. Ruchansky, S. Seo, Y. Liu // Proceedings of the ACM International Conference on Information and Knowledge Management (CIKM). – 2017. – P. 797-806.
5. BERT: pre-training of deep bidirectional transformers for language understanding / J. Devlin et al // Proceedings of the Conference of the North American Chapter of the Association for Computational Linguistics (NAACL-HLT). – 2019. – P. 4171-4186.
6. exBAKE: automatic fake news detection model based on bidirectional encoder representations from transformers / H. Jwa et al // Applied Sciences. – 2019. – Vol. 9, № 19. – P. 1-14.
7. Kaliyar R.K. FakeBERT: fake news detection in social media with a BERT-based deep learning approach / R.K. Kaliyar, A. Goswami, P. Narang // Multimedia Tools and Applications. – 2021. – Vol. 80, № 8. – P. 11765-11788.
8. Strubell E., Ganesh A., McCallum A. Energy and policy considerations for deep learning in NLP / Strubell E., Ganesh A., McCallum A. // Proceedings of the Annual Meeting of the Association for Computational Linguistics (ACL). – 2019. – P. 3645-3650.
9. Horne B.D. This just in: fake news packs a lot in title, uses simpler, repetitive content / B.D. Horne, S. Adalı // Proceedings of the International AAAI Conference on Web and Social Media (ICWSM). – 2017. – P. 759-766.
10. FakeNewsNet: a data repository with news content, social context and dynamic information / K. Shu et al // Big Data. – 2020. – Vol. 8, № 3. – P. 171-188.

Д.Т. Тюлемисова¹, А.К. Шайханова*, В.П. Марценюк², Г.Б. Бекешова¹, Б.Т. Смаилова³

¹Л.Н. Гумилев атындағы Евразия ұлттық университеті,
10000, Қазақстан Республикасы, Астана қ., Сәтбаев к-сі, 2

²Бельско-Биала университеті,
43-300, Польша, Бельско-Биала, к-сі Виллова 2,

³Шәкәрім университеті,
071412, Қазақстан Республикасы, Семей қ., Глинка к-сі 20 А

*e-mail: shaikhanova_ak@enu.kz

ЖАЛҒАН ЖАҢАЛЫҚТАРДЫ АНЫҚТАУ ҮШІН МАШИНАЛЫҚ ОҚЫТУ МОДЕЛЬДЕРІН ОҢТАЙЛАНДЫРУ: ГИПЕРПАРАМЕТРЛІ ТАҢДАУ ЖӘНЕ РОС ҚИСЫҚТАРЫН ТАЛДАУ

Сандық ортада, әсіресе әлеуметтік желілерде және жаңалықтар агрегаторларында ақпараттың белсенді және бақылаусыз таралуы жағдайында жалған жаңалықтарды автоматты түрде анықтау міндеті ерекше өзекті болып келеді. Пайдаланушы жасаған мазмұн көлемінің артуы және оны таратудың жоғары жылдамдығы ақпаратты қолмен тексеруді айтарлықтай қиындатады, бұл машиналық оқыту әдістерін қолдануды қажет етеді. Бұл зерттеудің мақсаты – гиперпараметрді таңдауға, есептеу тиімділігіне және жіктеу нәтижелерін түсіндіруге назар аударатырып, жалған жаңалықтарды анықтау үшін базалық машиналық оқыту модельдерін талдау, салыстырмалы түрде бағалау және оңтайландыру. Бұл зерттеуде TF-IDF әдісін қолдана отырып тазартылған және векторланған, «шын» және «жалған» деп белгіленген екілік белгілері бар ағылшын тіліндегі жаңалықтар элементтерін қамтитын мәтіндік деректер жиынтығы - ISOT

жалған жаңалықтар деректер жиынтығы қолданылады. Зерттеуде логистикалық регрессия, шешім ағашы, кездейсоқ орман және градиент күшейту модельдері енгізіліп, талданды. Жіктеу сапасы келесі көрсеткіштерді қолдану арқылы бағаланды: Дәлдік, Дәлдік, Еске түсіру, F1-ұпай, сондай-ақ ROC талдауы және қисық астындағы аудан (AUC). Градиент күшейту моделі ең жоғары жіктеу дәлдігін қамтамасыз ететіні, ал логистикалық регрессия есептеу шығындары мен оқыту уақытын айтарлықтай төмен деңгейде салыстырмалы сапаны көрсететіні көрсетілді. Сонымен қатар, зерттеуге деректердің ағып кетуін бақылау және нәтижелердің дұрыс түсіндірілуін қамтамасыз ететін шамадан тыс сапа көрсеткіштерінің себептерін талдау кірді. Зерттеу нәтижелері оңтайландырылған классикалық машиналық оқыту модельдері жоғары сапалы жалған жаңалықтарды анықтауды қамтамасыз ете алатынын және қолданбалы киберқауіпсіздік пен ақпарат ағынын бақылауда күрделі нейрондық желілік тәсілдерге ресурстарды тиімді пайдаланатын балама ретінде қарастырылуы мүмкін екенін растайды.

Түйін сөздер: жалған жаңалықтар, машиналық оқыту, гиперпараметрді оңтайландыру, мәтінді жіктеу, ROC талдауы, AUC, логистикалық регрессия.

D. Tyulemissova¹, A. Shaikhanova*, V. Martsenyuk², G. Bekeshova¹, B. Smailova³

¹L.N. Gumilyov Eurasian National University,
10000, Republic of Kazakhstan, Astana, Satpayev Street, 2

²University of Bielsko-Biala,
ul.Willowa 2, 43-300, Bielsko-Biala, Poland

³Shakarym University,
071412, Republic of Kazakhstan, Semey, Glinka Street 20 A

*e-mail: shaikhanova_ak@enu.kz

OPTIMIZING MACHINE LEARNING MODELS FOR FAKE NEWS DETECTION: HYPERPARAMETER SELECTION AND ROC CURVES ANALYSIS

In the context of the active and uncontrolled dissemination of information in the digital environment, especially on social media and news aggregators, the task of automatically identifying fake news has become especially relevant. The growing volume of user-generated content and the high speed of its distribution significantly complicate manual information verification, necessitating the use of machine learning methods. The aim of this study is to analyze, comparatively evaluate, and optimize baseline machine learning models for fake news detection, focusing on hyperparameter selection, computational efficiency, and interpretation of classification results. This study utilizes the ISOT Fake News Dataset, a text dataset containing English-language news items with binary labels labeled «true» and «false», cleaned and vectorized using the TF-IDF method. The study implemented and analyzed logistic regression, decision tree, random forest, and gradient boosting models. Classification quality was assessed using the following metrics: Accuracy, Precision, Recall, F1-score, as well as ROC analysis and the area under the curve (AUC). It was shown that the gradient boosting model provides the highest classification accuracy, while logistic regression demonstrates comparable quality with significantly lower computational costs and training time. Additionally, the study included data leakage monitoring and an analysis of the causes of inflated quality metrics, ensuring the correct interpretation of the results. The findings confirm that optimized classical machine learning models are capable of providing high-quality fake news detection and can be considered a resource-efficient alternative to more complex neural network approaches in applied cybersecurity and information flow monitoring.

Key words: fake news; machine learning; hyperparameter optimization; text classification; ROC analysis; AUC; logistic regression.

Сведения об авторах

Дана Болатовна Тюлемисова – магистр технических наук, докторант специальности 8D06306 – «Системы информационной безопасности» Евразийского национального университета им. Л.Н. Гумилева, г. Астана, Казахстан; e-mail: tyulemissova_db_3@enu.kz. ORCID: <https://orcid.org/0009-0006-5319-7742>.

Айгуль Кайрулаевна Шайханова* – доктор философии, профессор кафедры информационной безопасности Евразийского национального университета им. Л.Н. Гумилева, г. Астана, Республика Казахстан; e-mail: shaikhanova_ak@enu.kz. ORCID: <https://orcid.org/0000-0001-6006-4813>.

Васыл Марценюк – доктор технических наук, профессор кафедры информатики и автоматизации Бельско-Бяльского университета, г. Бельско-Бяла, Польша; e-mail: vmartsenyuk@ath.bielsko.pl. ORCID: <https://orcid.org/0000-0001-5622-1038>.

Гульвира Бауржановна Бекешова – магистр технических наук, ст. преподаватель кафедры информационной безопасности Евразийского национального университета им. Л.Н. Гумилева, г. Астана, Республика Казахстан; e-mail: bekeshova_gb@enu.kz. ORCID: <https://orcid.org/0000-0002-1635-4693>.

Балжан Темірболатқызы Смаилова – магистр, заведующая кафедрой математики, Шәкәрім университет, Семей, Республика Казакстан; e-mail: st.balzhan@gmail.com. ORCID: <https://orcid.org/0009-0005-4932-1774>.

Авторлар туралы мәліметтер

Дана Болатқызы Тюлемисова – 8D06306 – Ақпараттық қауіпсіздік жүйелері мамандығы бойынша инженерия ғылымдарының магистрі, PhD докторанты, Л.Н. Гумилев атындағы Еуразия ұлттық университеті, Астана, Қазақстан; e-mail: tyulemissova_db_3@enu.kz. ORCID: <https://orcid.org/0009-0006-5319-7742>.

Айгүл Кайрулақызы Шайханова* – PhD докторы, Л.Н. Гумилев атындағы Еуразия ұлттық университетінің ақпараттық қауіпсіздік кафедрасының профессоры, Астана, Қазақстан Республикасы; e-mail: shaikhanova_ak@enu.kz. ORCID: <https://orcid.org/0000-0001-6006-4813>.

Васыл Марценюк – инженерия ғылымдарының докторы, Бельско-Бяла университетінің есептеу техникасы және автоматика кафедрасының профессоры, Бельско-Бяла, Польша; e-mail: vmartsenyuk@ath.bielsko.pl. ORCID: <https://orcid.org/0000-0001-5622-1038>.

Гүлвира Бауржанқызы Бекешова – техника ғылымдарының магистрі, Л.Н. Гумилев атындағы Еуразия ұлттық университетінің Ақпараттық қауіпсіздік кафедрасының аға оқытушысы, Астана қ., Қазақстан Республикасы; e-mail: bekeshova_gb@enu.kz. ORCID: <https://orcid.org/0000-0002-1635-4693>.

Балжан Темірболатқызы Смаилова – магистр, математика кафедрасының меңгерушісі, Шәкәрім университеті; Семей, Қазақстан Республикасы; e-mail: st.balzhan@gmail.com. ORCID: <https://orcid.org/0009-0005-4932-1774>.

Information about the authors

Dana Tyulemissova – Master of Science (Engineering), PhD candidate in the specialty 8D06306 – Information Security Systems, L.N. Gumilyov Eurasian National University, Astana, Kazakhstan; e-mail: tyulemissova_db_3@enu.kz. ORCID: <https://orcid.org/0009-0006-5319-7742>.

Aigul Shaikhanova* – PhD, Professor, Department of Information Security, L.N. Gumilyov Eurasian National University, Astana, Republic of Kazakhstan; e-mail: shaikhanova_ak@enu.kz. ORCID: <https://orcid.org/0000-0001-6006-4813>.

Vasyl Martsenyuk – Doctor of Engineering Sciences, Professor, Department of Computer Science and Automation, Bielsko-Biala University, Bielsko-Biala, Poland; e-mail: vmartsenyuk@ath.bielsko.pl. ORCID: <https://orcid.org/0000-0001-5622-1038>.

Gulvira Baurzhanovna Bekeshova – Master of Technical Sciences, Senior Lecturer at the Department of Information Security, L.N. Gumilyov Eurasian National University; Astana, Republic of Kazakhstan; e-mail: bekeshova_gb@enu.kz. ORCID: <https://orcid.org/0000-0002-1635-4693>.

Balzhan Smailova – Master of Science, Head of Mathematics Department, Shakarim University; Semey, Republic of Kazakhstan; e-mail: st.balzhan@gmail.com. ORCID: <https://orcid.org/0009-0005-4932-1774>.

Поступила в редакцию 09.01.2026

Поступила после доработки 19.02.2026

Принята к публикации 20.02.2026

[https://doi.org/10.53360/2788-7995-2026-1\(21\)-10](https://doi.org/10.53360/2788-7995-2026-1(21)-10)

МРНТИ: 55.30.03



Т.С. Жылқыбаев*, Е.Р. Советбаев, А.Д. Золотов, Д.В. Мясоєдов, К.У. Зенкович

Шәкәрім университет,
071412, Республика Казакстан, г. Семей, ул. Глилки, 20 А
*e-mail: tursynhan.jylqybaev@shakarim.kz

ПРИМЕНЕНИЕ ЦИФРОВЫХ ДВОЙНИКОВ В УЧЕБНЫХ ЛАБОРАТОРИЯХ ПО АВТОМАТИЗАЦИИ

Аннотация: В статье представлены результаты разработки и исследования цифрового двойника учебного манипулятора в двумя звеньями длинами 230 и 170 мм и рабочим диапазоном 170-400 мм. Целью работы является создание виртуальной модели, позволяющая адекватно воспроизводить кинематическое и функциональное поведение реального устройства. Дополнительно проводить оценку возможностей ее применения в образовательном процессе. В

©Т.С. Жылқыбаев, Е.Р. Советбаев, А.Д. Золотов, Д.В. Мясоєдов, К.У. Зенкович, 2026