

ВЕСТНИК НОВОСИБИРСКОГО ГОСУДАРСТВЕННОГО УНИВЕРСИТЕТА

Научный журнал
Основан в ноябре 1999 года

Серия: Информационные технологии

2026. Том 24, № 1

СОДЕРЖАНИЕ

<i>Анищенко А. Г., Бульонков М. А.</i> Система ПРОСТОР: существующие возможности и новые аспекты развития	5
<i>Арабов М. К., Шумихина А. С.</i> Применение AutoML-технологий для автоматизированного обнаружения веб-атак в сетевом трафике на основе AutoGluon Tabular	18
<i>Браер П. С.</i> Применение нейросетевой модели для классификации и оценки инвестиционных проектов в рамках экономического анализа в ГИС	30
<i>Тумбинская М. В., Гатауллин Б. И., Хаерова Э. И.</i> Виртуальная лаборатория по защите информации на объекте информатизации	40
<i>Диаб Хайдар.</i> Методы генерации критических вопросов для анализа аргументации на русском языке	55
Информация для авторов	67

VESTNIK

NOVOSIBIRSK STATE UNIVERSITY

Scientific Journal
Since 1999, November
In Russian

Series: Information Technologies

2026. Volume 24, № 1

CONTENTS

<i>Anischenko A. G., Bulyonkov M. A.</i> System PROSTOR: Existing Features and New Aspects of Development.....	5
<i>Arabov M. K., Shumikhina A. S.</i> Application of Automl Technologies for Automated Detection of Web Attacks in Network Traffic Based on AutoGluon Tabular	18
<i>Braer P. S.</i> The Use of a Neural Network Model for the Classification and Evaluation of Investment Projects in the Framework of Economic Analysis in GIS.....	30
<i>Tumbinskaya M. V., Gataullin B. I., Haerova E. I.</i> Virtual Laboratory for Information Protection at the Informatization Facility	40
<i>Diab Haidar.</i> Methods for Generating Critical Questions for Analyzing Russian-Language Arguments.....	55
Instructions for Contributors.....	67

Editor in Chief M. M. Lavrentiev

Vice-Editor A. V. Avdeev

Executive Secretary D. P. Iksanova

Editorial Board of the Series

- I. V. Bychkov*, professor, academician (Irkutsk), *B. M. Glinsky*, professor (Novosibirsk)
A. N. Gorban, professor (Lester, GB), *E. P. Gordov*, professor (Tomsk)
B. S. Dobronets, professor (Krasnoyarsk), *A. M. Elizarov*, professor (Kazan)
G. N. Erokhin, professor (Kaliningrad), *A. I. Kamyshnikov*, professor (Khanty-Mansiysk)
G. P. Karev, professor (Maryland, USA), *N. A. Kolchanov*, professor, academician (Novosibirsk)
M. M. Lavrentjev, professor (Novosibirsk), *V. E. Malyshkin*, professor (Novosibirsk)
N. N. Mirenikov, professor (Aizu, Japan), *N. M. Oskorbin*, professor (Barnaul)
D. E. Palchunov, professor (Novosibirsk), *T. Pizansky*, professor (Ljubljana, Slovenia)
V. P. Potapov, professor (Kemerovo), *O. I. Potaturkin*, professor (Novosibirsk)
V. A. Serebryakov, professor (Moscow), *A. V. Starchenko*, professor (Tomsk)
S. I. Smagin, professor, corresponding member of RAS (Khabarovsk)
D. A. Tusupov, professor (Astana, Kazakhstan)
V. V. Shajdurov, professor, corresponding member of RAS (Krasnoyarsk)
Yu. I. Shokin, professor, academician (Novosibirsk)

*The journal is published quarterly in Russian since 1999
by Novosibirsk State University Press*

*The address for correspondence
Faculty of Information Technologies, Novosibirsk State University
1 Pirogov Street, Novosibirsk, 630090, Russia
Tel. +7 (383) 363 42 46*

E-mail address: inftech@vestnik.nsu.ru

On-line version: <http://elibrary.ru>

Научная статья

УДК 332.1, 004.94

DOI 10.25205/1818-7900-2026-24-1-5-17

Система ПРОСТОР: существующие возможности и новые аспекты развития

Александр Георгиевич Анищенко¹

Михаил Алексеевич Бульонков²

¹Институт экономики и организации промышленного производства СО РАН
Новосибирск, Россия

²Институт систем информатики им. А. П. Ершова СО РАН
Новосибирск, Россия

¹a.anishchenko1@g.nsu.ru, <https://orcid.org/0009-0008-6024-7327>

²bulyonkov@gmail.com, <https://orcid.org/0000-0002-2490-4077>

Аннотация

В данной статье описываются как уже существующие возможности, так и новые дополнения в системе ПРОСТОР – программе, которая позволяет рассчитывать оптимальные способы перевозки грузов между пунктами. Целью работы было создать необходимые условия для более удобного для пользователей взаимодействия с программой. Для этого были предложены некоторые нововведения в работе программы. Также в статье для понимания контекста работы с системой ПРОСТОР и вызываемых трудностей и неточностей были описаны существующие возможности программы и особенности работы с ней.

К новым аспектам развития системы были отнесены возможности по автоматическому преобразованию данных в соответствии с территориальным расположением транспортных узлов. Такое дополнение к программе позволяет любым ее пользователям самостоятельно определять размерность эксперимента транспортной задачи и решать его в программе.

Также программа была дополнена возможностью увеличения продуктовой размерности транспортной задачи. Это нововведение является необходимой мерой по увеличению точности создаваемого системой ПРОСТОР решения, поскольку соединяет оптимальные способы перевозки грузов между пунктами с реально сложившейся географической и продуктовой особенностью расположения этих пунктов между собой.

После введения этих возможностей в работу программы взаимодействие пользователей – исследователей, изучающих развитие транспортных систем, с системой ПРОСТОР становится более эффективным. Полученные результаты более точно описывают транспортную отрасль, а соответственно, дают более качественные прогнозы. Также такие изменения работы программы не только создают более точные результаты экспериментов, но и упрощают взаимодействие пользователя с системой.

Ключевые слова

развитие транспортной отрасли, ПРОСТОР, транспорт, транспортная задача, моделирование, ГИС

Финансирование

Исследование выполнено в рамках проекта НИР ИЭОПП СО РАН «Модели и методы исследования социально-экономического развития многоуровневой пространственной системы Российской Федерации в современных условиях».

Для цитирования

Анищенко А. Г., Бульонков М. А. Система ПРОСТОР: существующие возможности и новые аспекты развития // Вестник НГУ. Серия: Информационные технологии. 2026. Т. 24, № 1. С. 5–17. DOI 10.25205/1818-7900-2026-24-1-5-17

© Анищенко А. Г., Бульонков М. А., 2026

System PROSTOR: Existing Features and New Aspects of Development

Alexander G. Anischenko¹, Mikhail A. Bulyonkov²

¹Institute of Economics and Industrial Engineering SB RAS
Novosibirsk, Russian Federation

²The A. P. Ershov Institute of Informatics Systems SB RAS
Novosibirsk, Russian Federation

¹a.anishchenko1@g.nsu.ru, <https://orcid.org/0009-0008-6024-7327>

²bulyonkov@gmail.com, <https://orcid.org/0000-0002-2490-4077>

Abstract

This study aims to describe both existing and new capabilities of the system of PROSTOR and create conditions for more convenient use of the program. PROSTOR is a program designed to find optimal ways of transporting goods between hubs. The author introduces capabilities of the program including innovative features such as automatic data conversation based on location of transport hub. This addition allows each user to determine the dimensions of a transportation problem experiment and solve the problem within PROSTOR. This helps researchers to work with the program easier. Also the ability to increase product dimension was added to the program. This innovation is crucial for improving accuracy of the solutions generated by PROSTOR as it combines optimal ways of transporting goods and actual geographic and product features of hubs. This improvement in the program is necessary because it takes into account the complexity of building relationships in a real transportation system.

The implementation of these features into the PROSTOR system has enhanced its effectiveness for researchers studying the development of transport systems. Now the obtained results provide a more accurate description of the transport sector leading to higher-quality forecasts. These program modifications not only improve the precision of experimental outcomes but also simplify user interaction with the system.

Keywords

development of transport, PROSTOR, transport, transport task, modeling, GIS

Funding

This work was carried out under Research Project IEIE SB RAS No121040100262-7 (0260-2021-0007).

For citation

Anischenko A.G., Bulyonkov M. A. System PROSTOR: existing features and new aspects of development. *Vestnik NSU. Series: Information Technologies*, 2026, vol. 24, no. 1, pp. 5–17 (in Russ.) DOI 10.25205/1818-7900-2026-24-1-5-17

Введение

Развитие транспорта – тема, которая, пожалуй, никогда не потеряет своей актуальности. Для России, как для страны с колоссальными территориями, этот вопрос крайне важен, особенно в контексте развития ее азиатской части. Создание развитой опорной транспортной сети, включающей в себя различные виды транспорта, сможет превратить особенности географического положения России в ее конкурентное преимущество [1, с. 68]. Это позволит более эффективно распределять ресурсы на всей территории страны, что значительно сократит транспортную дискриминацию Азиатского региона, позволит осваивать слаборазвитые части страны, делая их более привлекательными для человека [2]. Актуальность данного вопроса подтверждается значительным увеличением числа публикаций на данную тему [3; 4].

Возникает вопрос о необходимости моделирования транспортной системы России, учитывая все ее многообразие. Это позволит исследовать развитие транспортных систем. Для решения этого вопроса была разработана и продолжает развиваться система ПРОСТОР. В данной статье будут описаны возможности для исследователей, предоставляемые программой, а также будут освещены аспекты нововведений, созданные в последние годы для эффективной работы с программой.

Пользователи системы ПРОСТОР, изучающие развитие транспортной отрасли России, нуждаются в комплексном подходе к проведению своих работ. Именно поэтому возникает

потребность в усовершенствовании системы, которое будет позволять учитывать различные аспекты, связанные с транспортом. Уже сейчас программа представляет собой удобный интерфейс для проведения таких исследований.

Моделирование опорной транспортной сети в системе ПРОСТОР

В данной части необходимо дать описательную характеристику работы системы ПРОСТОР: какие возможности она предлагает исследователям, какие есть нюансы работы с программой, какие преимущества и недостатки. МИКС ПРОСТОР (Модельно-информационная картографическая система, Прогнозирование развития опорной сети транспортной отрасли России) является программной системой для создания решений по формированию и развитию опорной транспортной сети России. Она позволяет определять различные варианты решения задачи оптимальной транспортировки грузов между различными географическими точками. Задача может относиться к разным уровням, начиная с перевозок внутри субъекта Российской Федерации, заканчивая перевозками в стране целиком с включением крупных транспортных узлов других стран [1, с. 70].

Транспортный комплекс модели состоит из двух видов элементов: точечные (узлы) и линейные (плечи). В первую очередь задаются узлы – точки на карте, между которыми происходят перевозки. Транспортные узлы представляются в виде крупных промышленно-транспортных городов, к которым тяготеют смежные районы, не имеющие своих крупных транспортных узлов [1, с. 71], поэтому можно считать, что узел – это не просто точка на карте, а целая географическая местность, объединенная в эту точку. Далее эти точки соединяются транспортными коридорами – плечами. Поскольку перевозку между узлами могут осуществлять различные виды транспорта, то каждый из них представлен своим плечом.

Модель также содержит отраслевую структуру рассматриваемой территории. В каждом узле производится и потребляется несколько видов продукта, задаваемых пользователем или уже определенных в программе. Если в узле разница между потреблением и производством какого-либо продукта положительна, то узел предъявляет спрос на этот продукт извне, иначе продукт из этого узла вывозится.

Таким образом, на основе балансов производства и потребления строится задача по оптимальному распределению продуктов между узлами, при условии расположения этих узлов в географическом пространстве, т. е. с учетом их транспортировки различными видами транспорта.

Также следует отметить, какие выделяются виды работ для каждого транспорта в каждом узле. В первую очередь на некоторый вид транспорт происходит *погрузка* произведенного продукта, которая зарождает грузопоток. Понятно, что грузопоток завершается полностью или частично *выгрузкой* продукта для потребления. Также внутри узла могут происходить промежуточные процессы грузоперевозки: *перегрузка* – смена вида транспорта (либо же перегрузка в рамках одного вида транспорта, связанная с инфраструктурными условиями перевозки¹) и *переработка транзита* – перевоз продукта без смены вида транспорта. На плече же производится только один вид работ – *провоз продукта по плечу*.

На каждый из этих видов работ, транспорта и продукта существует свой определенный тариф. Поэтому, если скомбинировать все условия перевозки из одной точки в другую с учетом использования необходимых видов работ в стоимостном выражении (с использованием тарифов), то можно получить суммарную стоимость этой перевозки.

¹ Так, например, при перевозке груза за границу на железной дороге происходит вынужденная перегрузка продукта, связанная с тем, что в России и в некоторых приграничных странах используется разная ширина колеи железной дороги.

С математической точки зрения решаемую системой ПРОСТОР задачу можно свести к задаче линейного программирования. Минимизируемой целевой функцией в этой задаче является функция издержек транспортировки, отражающая суммарную стоимость перевозки продуктов между узлами [5].

При этой целевой функции имеются ограничения, наложенные на переменные погрузки, разгрузки, перегрузки, транзита и перевоза по плечу с учетом спроса и предложения продукции в узле, пропускной способности транспорта на конкретном плече и в узле.

Стоит отметить, что тарифы в данной задаче определяются экзогенно: пользователь самостоятельно может варьировать цены на каждый из видов работ. Так, оптимальные решения по перевозке грузов определяются для каждого из наборов заданных тарифных планов и возможностей сети – пропускной способности дорог.

Также ПРОСТОР дает пользователю возможности варьировать условия моделирования для рассмотрения различных условий и параметров задачи, т. е. задавать не конкретные значения параметров, а диапазон их изменения. В задании эксперимента исследователь, работая с ПРОСТОРом, получает множество свобод.

В первую очередь, конечно, пользователь самостоятельно может задавать эксперимент: он не ограничен определенным набором узлов и плеч, поэтому может сам создавать новые и убирать старые элементы сети. Создание новых узлов позволит детализировать (или при удалении, наоборот, укрупнить) эксперимент. Например, ставя цель изучения развития транспортной системы Азиатской России, можно сконцентрировать узлы в Азиатской России, а не в Европейской, представив последнюю 2–3 узлами [6]. Такое расположение узлов можно увидеть на карте, представленной на рис. 1. Как видно, именно определение узлов формирует глубину и точность эксперимента.



Рис. 1. Карта экспериментальной сети программы
Fig. 1. Map of the experimental network of the program

Также пользователь свободен в создании новых плеч. Например, он может «проложить» железную дорогу Кызыл – Курагино и провести исследование о необходимости создания этой дороги в реальности и последствиях введения ее в эксплуатацию. На рис. 1 также можно увидеть плечи, формирующие Северо-Сибирскую железнодорожную магистраль и Баренцкомур.

В данном эксперименте можно рассматривать, как повлияет на экономику страны введение в эксплуатацию этих двух магистралей.

Выбор транспорта в эксперименте определяется исследователем. Так, пользователь может как включать в модель, так и исключать из нее различные виды транспорта. Конечно, наиболее интересным экспериментом является тот, который учитывает взаимодействие всех видов транспорта между собой. Однако существуют ограничения в информационном наполнении: объемы перевозки, тарифы, что затрудняет использование всех видов транспорта. Это относится и к выбору транспортируемых грузов. Пользователь самостоятельно определяет отраслевую структуру модели. Он может ограничиться небольшим, наиболее значимым с его точки зрения набором перевозимых грузов в модели.

Все эти возможности дают основу для проведения эксперимента, на которой будет строиться исследование, состоящее в варьировании параметров транспортной сети и сравнительном анализе получаемых решений.

Одним из ключевых видов исследований, связанных с изучением транспортной отрасли, является определение тарифов. Как отмечалось ранее, тарифы задаются в модель экзогенно. Это связано с необходимостью уменьшения размерности задачи математического программирования [7]. Тарифы задаются как коэффициенты удорожания для всех видов работ для каждого транспорта и продукта. Таким образом, пользователь самостоятельно задает в программе тарифы на разные виды деятельности для различных видов транспорта и продукта.

Поскольку тариф – это цена на услуги разных видов деятельности в транспортной отрасли, то варьирование тарифа – ключевой вопрос изменения оптимальных решений транспортировки. Например, через тарифы удобнее всего вводить в модель такого участника рынка, как государство и оценивать необходимость его вмешательства в функционирование системы. Также изменения тарифов могут отражать неопределенность на рынке транспортировки грузов, которая присутствует в рыночной экономике. Именно для этих целей в системе программы ПРОСТОР было введено понятие *вариатора*. Сначала пользователь определяет базовые тарифы – цены, которые установились в текущий момент времени на рынке. Далее задается эксперимент, заключающийся в варьировании тарифов относительно базового значения (рис. 2).

Вариатор	Диапазон
Железнодорожный - удорожание для плечей	1.0000 - 2.0000
Автомобильный - удорожание для узлов	1.0000 - 1.5000

Рис. 2. Вариаторы в системе ПРОСТОР

Fig. 2. Variators in the PROSTOR system

Аналогично вариаторы могут быть применены не только к тарифам, но и другим параметрам транспортной сети, например, пропускной способности узлов и плеч. Например, это позволяет поставить эксперимент, при котором изменяются свойства дорог: появились электрифицированные участки, увеличилось количество колеи и т. д.

Далее система ПРОСТОР рассчитывает издержки оптимальных способов перевозки для каждого значения параметра в заданном диапазоне. Если вариаторов несколько, то рассчитываются все комбинации варьируемых значений. Таким образом, формируется (возможно, весьма большое) количество решений. Любое решение система ПРОСТОР может визуализировать (рис. 3). При этом пользователь может управлять визуализацией, как выбирая интересующий его в данный момент перечень продуктов, так и указывая характеристики решения, такие как загруженность плеч, объем перевозимых продуктов и т. п.

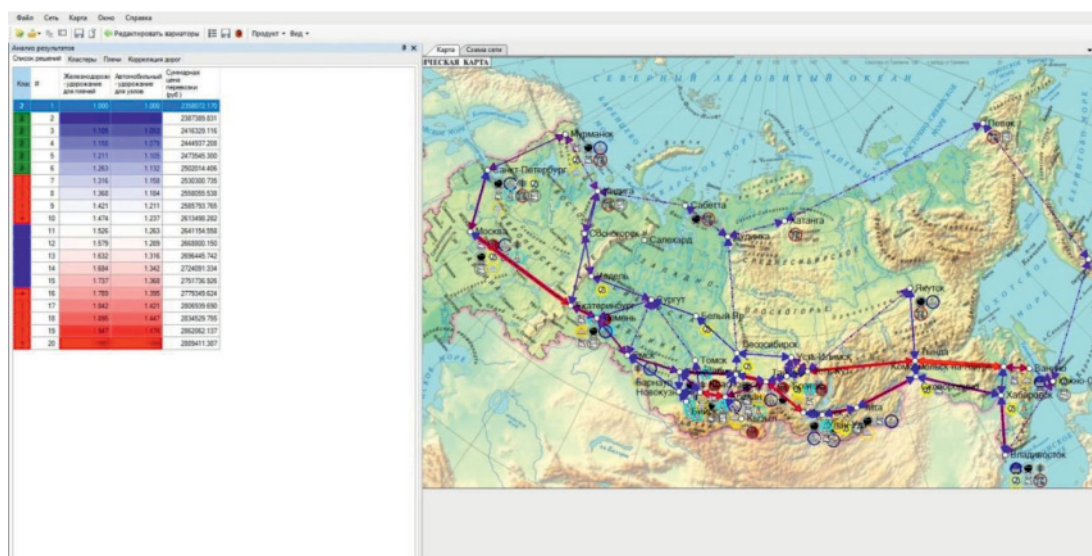


Рис. 3. Визуализация решения в системе ПРОСТОР
Fig. 3. Visualization of the solution in the PROSTOR system

Следующий шаг в проведении эксперимента – сравнительный анализ множества решений. Система ПРОСТОР предлагает пользователю широкий набор инструментов визуализации, основанный на методах математической статистики. Так, например, множество всех решений может быть разбито на кластеры похожих друг на друга решений. Из каждого кластера может быть выбран наиболее представительный, содержательно отражающий типичный способ перевозки. Таким образом, пользователь может сузить рассмотрение от сотен полученных решений до нескольких основных.

Получение входных данных

Ключевой проблемой использования подобной системы является получение реальных данных. Это касается как параметров транспортной сети, так и объемов производства/потребления в конкретном узле. Сбор этих данных – весьма непростая задача. Данные зачастую разрознены или засекречены, могут иметь разный формат и т. п. Эти данные не всегда точно соответствуют определенным выше параметрам, и поэтому извлечение нужной информации может потребовать предварительной обработки. Более того, поскольку мы решаем задачу прогнозирования развития транспортной сети, то надо иметь и обоснованный прогноз производства и потребления отдельных продуктов.

Транспорт является долго развивающейся отраслью [8], поэтому для того, чтобы прогнозировать развитие транспортной системы, необходимо уже сегодня понимать, какие условия перевозки грузов сложатся к прогнозируемому периоду. На базе прогнозных перевозок будут строиться соответствующие модели.

В проведенных нами экспериментах исходными данными о темпах роста являются результаты работы ОМММ (Оптимизационная межрегиональная межотраслевая модель) [9]. Целью работы этой системы является прогнозирование объемов перевозки грузов между субъектами Российской Федерации. Описание способов получения входных данных и применяемых ОМММ методов прогнозирования выходит за рамки данной статьи.

Имея данные о текущих перевозках и темпах их прироста, нетрудно получить прогнозные объемы перевозок в будущем. Поэтому результатом работы ОМММ можно считать опреде-

ление для каждого продукта p матриц T^p , где $T_{r,r'}^p$ – прогнозные объем поставок продукта p из региона r в регион r' . Наша задача состоит в том, чтобы каждый из этих наборов преобразовать во входные данные системы ПРОСТОР и потом сравнить полученные решения.

Проблема состоит в том, чтобы на основании этих объемов определить производство/потребление продукта p в конкретном узле транспортной сети. Для этого необходима информация о том, через какие узлы конкретный регион осуществляет поставки. Естественно, что здесь опять требуется дополнительная статистическая или экспертная информация.

Казалось бы, что можно просто опираться на территориальную принадлежность того или иного транспортного узла, но на практике это не так. Во-первых, как говорилось ранее, несколько регионов могут с точки зрения транспортной сети представляться одним узлом, например, вся европейская часть при анализе перевозок в Сибири и на Дальнем Востоке. Во-вторых, в одном регионе может быть несколько транспортных узлов, по-разному используемых для разных продуктов, как, например, Норильск, Лесосибирск и Красноярск в Красноярском крае. Наконец, возможны и более «хитрые» ситуации, когда некоторый регион использует для поставок как «собственные» узлы, так и те, которые территориально находятся в смежных регионах. Общий метод, учитывающий все эти аспекты, состоит в задании так называемых *матриц участия* $In_{r,j}^p$ и $Out_{j,r}^p$, определяющих долю использования регионом r транспортного узла j для ввоза и вывоза продукта p . Эти матрицы должны удовлетворять следующим свойствам для любого продукта p :

$$\sum_j In_{r,j}^p = \sum_j Out_{j,r}^p = 1.$$

Содержательно они означают, что всё, что регион поставляет/потребляет, так или иначе вывозится/ввозится через некоторый транспортный узел.

Далее, домножив матрицу T^p слева на матрицу In^p , а справа на матрицу Out^p , мы получим матрицу перевозок продукта между узлами:

$$S^p = In^p T^p Out^p.$$

Используя свойства матриц участия, несложно показать, что суммарный объем перевозок остается неизменным.

Заметим, что сумма элементов любой строки в этой матрице означает суммарный ввозимый объем продукта p , а для столбца – вывозимый. Таким образом, мы можем получить объем продукта p , произведенного в узле j :

$$A_j^p = \sum_{j'} S_{j,j'}^p - \sum_j S_{j,j}^p.$$

Если это значение отрицательное, то его следует рассматривать, как потребляемый объем. Отметим, что значение $S_{j,j}^p$ при таком суммировании появляется дважды и в результате теряется. Содержательно это означает, что перевозки внутри одного региона игнорируются.

Таким образом, были получены значения, необходимые для системы ПРОСТОР, но они оказываются не вполне корректными по следующей причине. Просуммировав все объемы данного продукта, ввозимые/вывозимые в данный узел, была потеряна привязка к тому, в каком регионе этот продукт был произведен. Поскольку система ПРОСТОР не учитывает место производства продукта, а только издержки на перевозку, то она может «посчитать», что гораздо выгоднее везти контейнеры в Томск из Красноярска, нежели из Москвы, поскольку не учитывается то, что находится в контейнерах. Другим примером могут служить автомобили, произведен-

ные в Москве и в Нижнем Новгороде – это совершенно разные продукты, и разные регионы предлагают свой спрос на эти продукты.

С другой стороны, вполне возможна и ситуация, когда разные регионы производят один и тот же продукт, и то, что результаты ОМММ распределили его поставки именно так, а не иначе, определяется другими факторами.

Таким образом, перед нами возникает выбор: либо смириться с этой некорректностью, либо дифференцировать продукты по их месту производства и решать существенно более сложную задачу. Забегая вперед, скажем, что были проведены эксперименты как для первого, так и для второго случая. Как можно было ожидать, при дифференциации продуктов загруженность транспортной сети возрастает, хотя и результаты для первого случая могут служить для предварительных оценок.

Для первого случая, когда мы смиряемся с некорректностью дифференциации продукта (рис. 4), представлено решение из программы ПРОСТОР. В этом эксперименте мы не задавали вариаторы показателей, а только посмотрели результат для базовых значений.

В левой части окна показан список решений с указанием суммарной стоимости всех перевозок в эксперименте. В нашем случае только одно решение, поскольку мы не варьировали ни один показатель. В правой же части окна располагается схема транспортных узлов и плеч. Такая схема – аналог физической карты, только при большом количестве узлов и плеч она более наглядная.

Плечи на этой схеме отражают оптимальные способы перевозки грузов. С одной стороны, тип линии соответствует виду транспорта. В данном эксперименте сплошная линия – это железнодорожный транспорт, пунктирная – автомобильный, пунктирная с точкой – водный (речной, морской и ледоходный транспорт). С другой стороны, цвет и толщина линии отражает объем перевозимых грузов по плечу. Так, красные – самые толстые линии указывают на больший объем, тогда как синие и тонкие – на меньший.

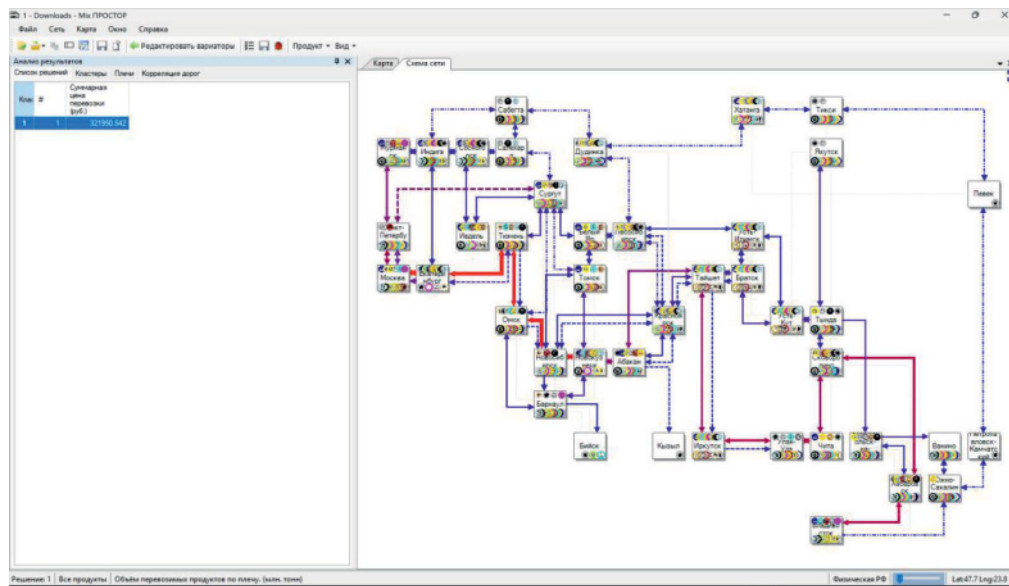


Рис. 4. Решение для первого случая

Fig. 4. The solution of the first case

Видно, что если мы не разделяем продукты одной группы по территориальному расположению производства (считаем, что продукты однородные), то при заданных условиях суммарная стоимость всех перевозок в эксперименте составляет примерно 321 950 у. е.

Теперь, если мы дифференцируем продукты, то получим другое решение (рис. 5). Как видно, стоимость перевозки возрастет до 475 175 у. е., т. е. примерно в полтора раза. Также, не разбирая подробно структуру перевозок, можно увидеть отличие в этой самой структуре для двух типов экспериментов. Некоторые линии поменяли толщину и цвет, т. е. объем перевозки грузов по этому плечу изменился, причем значительно.

Можно сделать вывод, что с экономической точки зрения переход от случая однородных продуктов к случаю дифференцированных территориально продуктов важен. Эта важность особенно остро проявляется для экспериментов, связанных с большим числом узлов и продуктов.

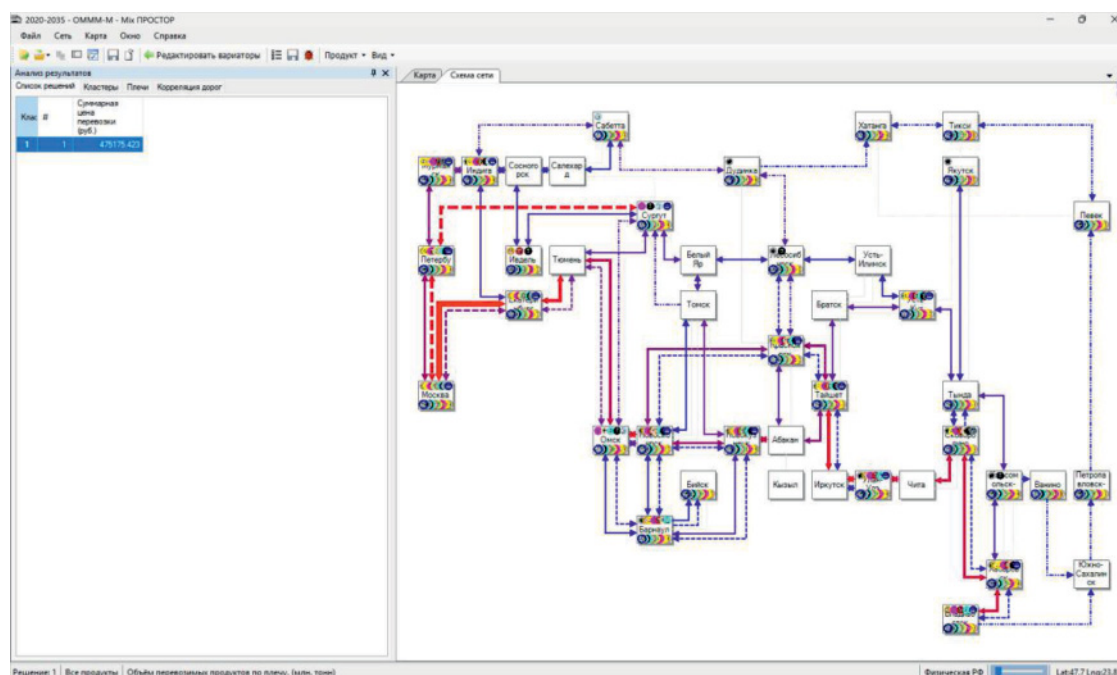


Рис. 5. Решение для второго случая

Fig. 5. The solution of the second case

Подводя итог, технологически процесс исследования состоит в выполнении следующих шагов:

1. *Пользователь* задает узлы, плечи, тарифы и пропускные способности используя систему ПРОСТОР.
2. *Пользователь* добавляет к таблицам MS Excel, хранящим результаты работы системы ОМММ, информацию о матрицах участия, используя коды узлов, заданные в системе ПРОСТОР. Поскольку эти матрицы обычно сильно разрежены, то для их представления разработан компактный формат.
3. Система ПРОСТОР импортирует из заполненных таблиц MS Excel данные как о поставках продуктов между регионами, так и матрицы участия и на их основе заполняет параметры производства/потребления для всех узлов.
4. Система ПРОСТОР просчитывает два варианта – для текущего момента и для прогноза.
5. *Пользователь* анализирует полученные решения либо средствами системы ПРОСТОР, либо сторонними инструментами, предварительно экспортировав данные.

Более подробно метод связи систем ОМММ и ПРОСТОР рассмотрен в работе [10].

Эффективность

Дифференциация продуктов по месту производства приводит к существенному увеличению размерности решаемой задачи. Как уже говорилось, мы сводим транспортную задачу к задаче линейного программирования, а затем используем один из многих доступных решателей для получения объемов перевозки. Каждое равенство или ограничение в решаемой системе неравенств описывает локальные требования для конкретного узла или плеча. Например, «сумма объемов ввозимых/вывозимых в данном узле продуктов должна соответствовать производству/потреблению этого продукта». Понятно, при увеличении количества продуктов количество неравенств в системе увеличивается кратно, что демонстрируется в таблице на примере сети, с которой мы проводили эксперименты.

Размеры решаемых задач

The size of the tasks to be solved

Элементы задачи	Без дифференциации	С дифференциацией
Продукты	19	514 ²
Виды транспорта	5	5
Узлы	42	42
Плечи	99	99
Переменные	9108	223716
Ограничения	4453	105318

Поскольку для решения задачи линейного программирования система ПРОСТОР использует сторонние решатели, то мы можем только предполагать, как вычислительная сложность зависит от количества переменных и ограничений. Кроме того, сложность вычислений в конкретном решателе может зависеть не только от этих значений, но и от вида матрицы, например, от ее разреженности, блочности и т. п.

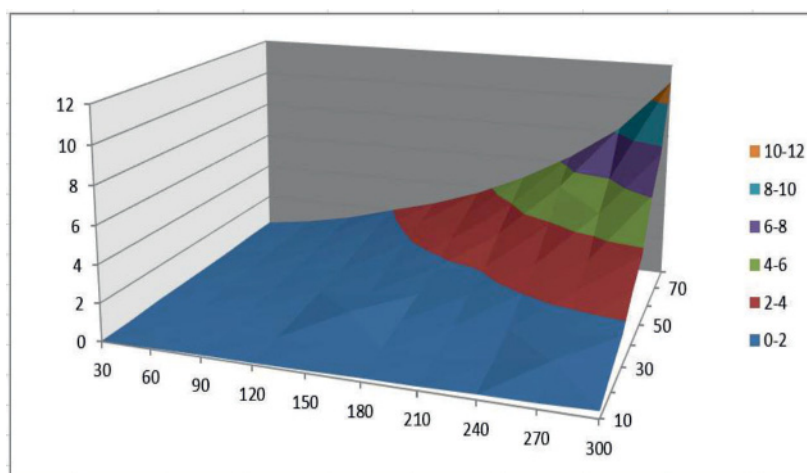


Рис. 6. Зависимость времени работы решателя от параметров задачи

Fig. 6. Dependence of the solver's runtime on the problem parameters

² Заметим, что указанное значение меньше, чем произведение количества продуктов на количество узлов ($19 \cdot 42 = 798$), поскольку на этапе переноса данных из ОМММ мы игнорировали нулевые или пренебрежимо малые поставки продуктов из одного региона в другой.

Помимо чисто теоретической оценки порядка сложности, для нас важны и конкретные значения. Представим, что один решатель справляется с задачей за 10 сек., а другой за 100. С точки зрения пользователя, это может стать критичным, если задачу с вариациями необходимо просчитать 100 раз, т. е. ждать решения либо 16 минут, либо почти 3 часа.

В работе [10] был проведен сравнительный анализ двух доступных решателей – GPLK и GoogleOR Tools. Был создан стенд, который позволял генерировать случайные транспортные сети с заданными параметрами – количеством узлов, плеч и продуктов. Во-первых, выяснилось, что GoogleOR в несколько раз быстрее, чем GPLK. Во-вторых, как и ожидалось, основное влияние на сложность оказывает количество продуктов, как показано на рис. 6, где количество продуктов меняется от 30 до 300, а количество вершин от 10 до 70, а время исчисляется в секундах. Построенные графики позволяют предположить, что зависимость от количества продуктов нелинейна.

Таким образом, используемый в системе ПРОСТОР решатель пока справляется с поставленной задачей, но масштабируемость его в данном применении ограничена.

Заключение

Введение в ПРОСТОР возможности различать продукт по расположению его производства – важный этап модернизации программы. В данной статье было показано, что ситуация с грузопотоками в стране сильно меняется при введении такого различия. Такое изменение приводит нашу работу к более точным результатам, которые правильнее описывают реальную картину мира.

Таким образом, исследователи, пользующиеся системой ПРОСТОР, смогут строить более реалистичные эксперименты и, соответственно, получать более точные результаты своей работы, что, безусловно, является преимуществом системы. Однако важно понимать, что за точность приходится платить увеличением времени работы решателя программы.

Так или иначе, точностью исследования пренебрегать не стоит. ПРОСТОР способен ставить задачи прогнозирования развития транспортной системы России и отвечать на вопрос, какие условия необходимы для достижения оптимального развития. Как уже упоминалось раньше, транспортную отрасль крайне необходимо исследовать в аспекте будущего состояния. Таким образом, исследователю необходимо строить прогноз, предположение о том, какая экономическая ситуация сложится в будущем. Так мы сможем определить, какие транспортные коридоры требуют модернизации, какие проекты требуется запускать уже сейчас, а какие отложить или вообще забыть. Все эти вопросы можно решить с помощью системы ПРОСТОР, именно поэтому так необходима точность работы программы для создания состоятельных оценок прогноза.

Описанная в статье модификация – одна из ключевых, но не единственная для системы ПРОСТОР. Программа имеет большой потенциал к дальнейшему развитию и поэтому пользуется интересом со стороны исследователей.

Список литературы

1. **Воробьева В. В., Малов В. Ю., Радченко В. В. и др.** Модель прогнозирования развития опорной транспортной сети России // Моделирование производственных и региональных систем на основе ГИС и информационных технологий: сб. науч. тр.; Институт экономики и организации промышленного производства СО РАН. Новосибирск: ИЭОПП СО РАН, 2011. С. 68–96.

2. **Крюков В. А., Коломак Е. А., Сулов Н. И., Костин А. В.** Азиатская Россия: от проблем к росту // Науч. тр. Вольного эконо. общества России. 2023. Т. 241, № 3. С. 110–128. DOI 10.38197/2072-2060-2023-241-3-110-128.
3. Анализ и оценка процессов создания и развития в Азиатской России транспортной магистральной сети различного назначения. Новосибирск: ИЭОПП СО РАН, 2024. 484 с. DOI 10.36264/978-5-89665-385-1-2024-021-484
4. **Щербанин Ю. А.** Транспорт Азиатской России: вызовы и возможности // ЭКО. 2024. № 3 (597). С. 134–156. DOI 10.30680/ECO0131-7652-2024-3-134-156
5. **Бульонков М. А., Воронов Ю. П., Ершов Ю. С.** и др. Ситуационная комната как элемент организации экспертного сообщества: задачи планирования и прогнозирования / Ин-т экономики и организации промышленного производства СО РАН. Новосибирск: ИЭОПП СО РАН, 2018. 259 с.
6. **Анищенко А. Г.** Сопряжение результатов расчетов по народнохозяйственной модели (ОМММ) и модели опорной транспортной сети: специфика Азиатской России // Молодые ученые – экономике региона: Материалы XXIV Всерос. науч.-практ. конф. с междунар. участием, Вологда, 6 дек. 2024 г. Вологда: Вологод. науч. центр РАН, 2025. С. 35–40.
7. **Бульонков М. А., Нестеренко Т. В.** Автоматизация исследований развития опорной транспортной сети // Вестник СибГУТИ. 2019. № 3. С. 45–54.
8. **Гасанов М. А., Бабаева Д. Р.** Формирование долгосрочной стратегии развития производственно-отраслевой инфраструктуры региона // Региональные проблемы преобразования экономики. 2017. № 5 (79). С. 76–85.
9. **Гранберг А. Г., Селиверстов В. Е., Сулов В. И.** и др. Оптимизационные межрегиональные межотраслевые модели. Новосибирск: Наука, 1989. 257 с.
10. **Бекчанов Х. М.** Сравнительный анализ производительности решателей задачи линейного программирования в применении к транспортным сетям: выпускная квалификационная работа магистра. Новосибирск: НГУ, ММФ, 2024. 30 с.

References

1. **Vorobyeva V. V., Malov V. Y., Radchenko V. V. et al.** Model prognozirovaniya razvitiya opornoy transportnoy seti Rossii [A model for forecasting the development of Russia's core transport network]. In *Modelirovaniye proizvodstvennykh i regional'nykh sistem na osnove GIS i informatsionnykh tekhnologiy: sbornik nauchnykh trudov*. Novosibirsk: IEIE SB RAS, 2011, pp. 68–96. (in Russ.)
2. **Kryukov V. A., Kolomak E. A., Suslov N. I., Kostin, A. V.** Asian Russia: from Problems to Growth. *The Free Economy Journal*, 2023, vol. 241, no. 3, pp. 110–128. DOI 10.38197/2072-2060-2023-241-3-110-128 (in Russ.)
3. Analysis and evaluation of the processes of creation and development of a transport backbone network for various purpose in Asian Russia. Novosibirsk: IEIE SB RAS publ., 2024. 484 p. DOI 10.36264/978-5-89665-385-1-2024-021-484. (in Russ.)
4. **Shcherbanin Y. A.** Transportation in Asian Russia: challenges and opportunities. *ECO*, 2024, no. 3(597), pp. 134–156. DOI 10.30680/ECO0131-7652-2024-3-134-156 (in Russ.)
5. **Bulyonkov M. A., Voronov Y. P., Yershov Y. S. et al.** Situation room as an element of expert community organization: planning and forecasting tasks. Novosibirsk: IEOPP SO RAN publ., 2018. 259 p. (in Russ.)
6. **Anishchenko A. G.** Interfacing the Results of the Economic Model (Ommm) and the Model of the Backbone Transportation Network: Specifics of Azian Russia. In *Young Scientists for the Regional Economy*. Vologda, 2025, pp. 35–40. (in Russ.)

7. **Bulyonkov M. A., Nesterenko T. V.** Research Automation of the Transport Network Development. *Vestnik SibGUTI*, 2019, no. 3, pp. 45–54. (in Russ.)
8. **Gasanov M. A., Babayeva D. R.** Forming a Long-Term Strategy Of Development of the Production-Branch Infrastructure of the Region. *Regional Problems of Economic Transformation*, 2017, no. 5(79), pp. 76–85. (in Russ.)
9. **Granberg A. G., Seliverstov V. Ye., Suslov V. I. et al.** Optimizatsionnyye mezhregional'nyye mezhotraslevyye modeli [Optimization interregional interindustry models]. Novosibirsk: Nauka publ., 1989. 257 p. (in Russ.)
10. **Bekchanov Kh. M.** Sravnitel'nyy analiz proizvoditel'nosti reshatelay zadachi lineynogo programmirovaniya v primenении k transportnym setyam [Comparative analysis of the performance of linear programming solvers applied to transport networks]. Novosibirsk: NSU publ., 2024. 30 p. (in Russ.)

Информация об авторах

Анищенко Александр Георгиевич, лаборант

Бульонков Михаил Алексеевич, кандидат физико-математических наук, старший научный сотрудник, заведующий лабораторией ИСИ СО РАН

Information about the Authors

Alexander G. Anischenko, Laboratorian

Mikhail A. Bulyonkov, Candidate of Physical and Mathematical Sciences, Senior Researcher, Head of the Laboratory IIS SB RAS

Статья поступила в редакцию 25.11.2025;

одобрена после рецензирования 10.02.2026; принята к публикации 10.02.2026

The article was submitted 25.11.2025;

approved after reviewing 10.02.2026; accepted for publication 10.02.2026

Научная статья

УДК 004.8

DOI 10.25205/1818-7900-2026-24-1-18-29

Применение AutoML-технологий для автоматизированного обнаружения веб-атак в сетевом трафике на основе AutoGluon Tabular

**Муллошараф Курбонович Арабов¹
Анна Сергеевна Шумихина²**

Казанский федеральный университет
Казань, Россия

¹MKArabov@kpfu.ru, <https://orcid.org/0000-0003-2525-1183>

² ansh00@mail.ru

Аннотация

Рост сложности и масштабов сетевых атак обуславливает необходимость перехода от традиционных сигнатурных систем обнаружения вторжений (IDS) к более адаптивным подходам, основанным на методах машинного обучения. В настоящем исследовании представлен сравнительный анализ эффективности классических алгоритмов машинного обучения и фреймворка автоматизированного машинного обучения (AutoML) AutoGluon при решении задачи многоклассовой классификации сетевого трафика. В качестве экспериментальной базы использован общедоступный датасет CICIDS2017, включающий как легитимные соединения, так и различные типы атак, моделирующие реальные условия функционирования сетей.

Проведен анализ сетевых характеристик и их дискриминативной способности (Ассигасу), F1-меры и времени обучения. Результаты показали, что ансамблевые алгоритмы, в частности случайный лес (Random Forest) и AutoGluon, достигают наивысших показателей общей точности (более 99,5 %). Вместе с тем выявлена критически низкая эффективность обнаружения миноритарных классов атак, таких как SQL-инъекции (SQL Injection) и межсайтовый скриптинг (XSS). Установлено, что данная проблема обусловлена не только сильным дисбалансом данных, но и особенностями самих атак, которые часто маскируются под легитимный трафик и не формируют выраженных статистических аномалий.

Таким образом, исследование демонстрирует потенциал и ограничения применения AutoML-фреймворков в области кибербезопасности. Практическая значимость работы заключается в возможности сокращения времени разработки систем обнаружения вторжений (IDS) при сохранении высокой точности, а перспективы дальнейших исследований связаны с интеграцией методов балансировки классов и разработкой гибридных моделей для улучшения распознавания редких атак.

Ключевые слова

обнаружение сетевых атак, машинное обучение, автоматизированное машинное обучение, AutoGluon Tabular, кибербезопасность, веб-атаки, анализ сетевого трафика

Для цитирования

Арабов М. К., Шумихина А. С. Применение AutoML-технологий для автоматизированного обнаружения веб-атак в сетевом трафике на основе AutoGluon Tabular // Вестник НГУ. Серия: Информационные технологии. 2026. Т. 24, № 1. С. 18–29. DOI 10.25205/1818-7900-2026-24-1-18-29

© Арабов М. К., Шумихина А. С., 2026

Application of Automl Technologies for Automated Detection of Web Attacks in Network Traffic Based on AutoGluon Tabular

Mullosharaf K. Arabov¹, Anna S. Shumikhina²

Kazan Federal University
Kazan, Russian Federation

¹MKArabov@kpfu.ru, <https://orcid.org/0000-0003-2525-1183>

² ansh00@mail.ru

Abstract

The increasing complexity and scale of network attacks necessitates the transition from traditional signature intrusion detection systems (IDS) to more adaptive approaches based on machine learning methods. This study presents a comparative analysis of the effectiveness of classical machine learning algorithms and the automated machine learning (AutoML) framework AutoGluon in solving the problem of multiclass classification of network traffic. The publicly available CICIDS2017 dataset was used as an experimental base, which includes both legitimate connections and various types of attacks that simulate real-world network conditions.

An analysis of network characteristics and their discriminative ability is carried out, as well as a detailed assessment of model performance using the metrics of general accuracy, F1-measure, and training time. The results showed that ensemble algorithms, in particular Random Forest and AutoGluon, achieve the highest overall accuracy (more than 99.5%). At the same time, the detection efficiency of minor classes of attacks, such as SQL Injection and cross-site scripting (XSS), is critically low. It has been established that this problem is caused not only by a strong data imbalance, but also by the characteristics of the attacks themselves, which are often disguised as legitimate traffic and do not form pronounced statistical anomalies.

Thus, the study demonstrates the potential and limitations of using AutoML frameworks in the field of cybersecurity. The practical significance of the work lies in the possibility of reducing the development time of intrusion detection systems (IDS) while maintaining high accuracy, and the prospects for further research are related to the integration of class balancing methods and the development of hybrid models to improve the recognition of rare attacks.

Keywords

network attack detection, machine learning, automated machine learning, AutoGluon Tabular, cybersecurity, web attacks, network traffic analysis

For citation

Arabov M. K., Shumikhina A. S. Application of AutoML technologies for automated detection of web attacks in network traffic based on AutoGluon Tabular. *Vestnik NSU. Series: Information Technologies*, 2026, vol. 24, no. 1, pp. 18–29 (in Russ.) DOI 10.25205/1818-7900-2026-24-1-18-29

Введение

Рост количества и сложности кибератак, особенно на веб-приложения, требует создания более эффективных и адаптивных систем защиты. Традиционные сигнатурные методы часто не успевают за новыми угрозами, что стимулирует активное внедрение методов машинного обучения (МО). Однако создание моделей машинного обучения требует значительных экспертных знаний и ресурсов для этапов инженерии признаков (feature engineering), выбора алгоритмов и тонкой настройки гиперпараметров.

Автоматизированное машинное обучение (AutoML) предлагает решение этой проблемы, автоматизируя процесс построения полного конвейера машинного обучения. Одним из современных и производительных фреймворков в этой области является AutoGluon Tabular, предназначенный для работы с табличными данными.

Целью данного исследования является оценка эффективности применения фреймворка AutoGluon Tabular для задачи автоматизированного обнаружения веб-атак в сетевом трафике в сравнении с классическими алгоритмами машинного обучения. В задачи входило: проведение детального разведочного анализа данных (EDA), оценка значимости сетевых признаков,

обучение и сравнение моделей, а также анализ их эффективности на сильно несбалансированном датасете.

1. Обзор предметной области

Существует множество работ, посвященных применению машинного обучения для анализа сетевого трафика [1; 2]. Исследования [3; 4] демонстрируют высокую эффективность ансамблевых методов, таких как случайный лес (Random Forest) и XGBoost, для обнаружения вторжений. В работе [5] акцент делается на проблеме дисбаланса классов в сетевых датасетах. В работах [6; 7] исследуется эффективность применения автоматизированного машинного обучения (AutoML) для задач прогнозирования временных рядов, демонстрируются его преимущества в сокращении трудозатрат на построение конвейера и достижении конкурентоспособной точности.

Однако вопросы, связанные с автоматизацией всего процесса построения моделей для обнаружения сетевых атак с помощью фреймворков AutoML (в частности, AutoGluon Tabular) и всесторонним анализом информативности сетевых признаков для конкретных типов веб-атак, освещены недостаточно полно. Настоящее исследование направлено на восполнение этого пробела, расширяя область применения AutoML на задачи кибербезопасности.

2. Методология исследования

2.1. Датасет и предобработка данных

В данной работе для исследования задачи обнаружения сетевых атак использовался публичный датасет CICIDS2017 [1], разработанный Canadian Institute for Cybersecurity. Данные были собраны в контролируемой среде и включают как нормальный (BENIGN), так и вредоносный сетевой трафик, симулирующий реальные атаки, включая веб-атаки: атаку перебором (Brute Force), межсайтовый скриптинг (XSS) и SQL-инъекции (SQL Injection).

Общий объем датасета составляет 164 300 сетевых потоков, описываемых 64 признаками. Каждый поток характеризуется временными, статистическими и протокольными метриками, полученными с помощью инструментов генерации сетевого трафика, таких как CICFlowMeter.

На этапе предварительной обработки данных были выявлены и исключены признаки, не несущие предиктивной ценности:

- Признак Bwd PSH Flags был удален из-за нулевой дисперсии (все значения идентичны), что делает его бесполезным для классификации.
- Признак Avg Bwd Segment Size был исключен, так как он не был использован ни одной моделью в процессе генерации признаков, что указывает на его избыточность или низкую информативность.

Таблица 1

Распределение сетевых потоков по типам трафика

Table 1

Distribution of Network Flows by Traffic Types

Тип трафика	Количество	Доля, %
BENIGN (нормальный трафик)	162 157	98,696
Веб-атака – атака перебором (Brute Force)	1 470	0,895
Веб-атака – межсайтовый скриптинг (XSS)	652	0,397
Веб-атака – SQL-инъекции (SQL Injection)	21	0,013
Всего	164 300	100

В результате предобработки для дальнейшего анализа было отобрано 62 признака, которые были использованы для обучения моделей машинного обучения.

В распределении по классам (табл. 1) наблюдается выраженный дисбаланс: нормальный трафик (BENIGN) составляет 98,696 % от общего объема данных, в то время как веб-атаки – атака перебором (Brute Force), межсайтовый скриптинг (XSS) и SQL-инъекции (SQL Injection) – составляют 0,895; 0,397 и 0,013 % соответственно. Такое распределение типично для реальных сетевых сред и требует применения метрик, устойчивых к дисбалансу.

Для оценки производительности моделей датасет был разделен на обучающую (131 440 образцов) и тестовую (32 860 образцов) выборки с сохранением стратификации по классам, что обеспечивает репрезентативность распределения атак в обеих выборках и позволяет корректно оценивать обобщающую способность моделей.

2.2. Сетевые признаки и анализ аномалий

Все 62 признака были классифицированы по семантическим категориям, включающим: `port_related`, `packet_stats`, `packet_length`, `flow_timing`, `forward_timing`, `backward_timing`, `protocol_flags`, `header_info`, `packet_rates`, `packet_size`, `window_size`, `subflow_stats`, `connection_stats`, `ratios`.

Статистический анализ выявил значительное количество выбросов (аномалий) в ключевых признаках, определенных по методу межквартильного размаха (IQR). Наибольшее число аномалий зафиксировано в следующих признаках:

- Flow Duration – 23,58 % аномалий ($\mu = 12,89$ млн мкс, $\sigma = 32,41$ млн мкс),
- Total Backward Packets – 22,76 %,
- Flow Bytes/s – 20,41 %.

Эти признаки демонстрируют высокую вариативность и могут быть критичны для выявления аномального поведения, например, при длительных сессиях или интенсивных атаках.

Анализ целевых портов показал, что основной трафик приходится на стандартные сервисы:

- порт 53 (DNS) – 74 370 подключений,
- порт 443 (HTTPS) – 34 124 подключения,
- порт 80 (HTTP) – 19 746 подключений.

Такое распределение соответствует реальной практике сетевого взаимодействия и указывает на то, что злоумышленники часто эксплуатируют наиболее распространенные протоколы, включая DNS и HTTPS.

2.3. Методы машинного обучения и оценка моделей

Для решения задачи многоклассовой классификации сетевого трафика был применен современный фреймворк AutoGluon Tabular (версия 1.4.0), реализующий подход AutoML (автоматическое машинное обучение). AutoGluon автоматически выполняет предобработку данных, генерацию признаков, подбор моделей, настройку гиперпараметров и построение ансамблей. В процессе обучения использовалась метрика качества `f1_weighted`, что особенно важно при сильном дисбалансе классов.

Были заданы пользовательские гиперпараметры для следующих моделей:

- XGBoost: `n_estimators=100`,
- RandomForest: `n_estimators=100`,
- LightGBM: обучение по умолчанию.

В результате AutoGluon обучил три базовые модели (LightGBM, Random Forest, XGBoost) и построил ансамбль второго уровня `WeightedEnsemble_L2`, в котором XGBoost получил вес 1.0, что указывает на его доминирующую роль в финальных предсказаниях.

Для сравнения с AutoGluon были протестированы традиционные алгоритмы машинного обучения:

- случайный лес (Random Forest) – ансамблевый метод на основе решающих деревьев,
- метод опорных векторов (Support Vector Machine, SVM) с RBF-ядром,
- логистическая регрессия (Logistic Regression) – линейный классификатор,
- наивный байесовский классификатор (Gaussian Naive Bayes, GNB) – вероятностная модель, основанная на предположении независимости признаков.

Все модели оценивались по следующим метрикам на независимой тестовой выборке:

- Accuracy – общая точность,
- Precision (точность), Recall (полнота), F1-мера (F1-score) – с акцентом на взвешенную макросреднюю (f1_weighted) для учета дисбаланса классов.

Для анализа значимости признаков использовались:

- Дисперсионный анализ (ANOVA) с F-критерием – для оценки статистической значимости различий средних значений признаков между классами. Установлено, что 46 из 62 признаков являются статистически значимыми ($p < 0,001$).
- Перестановочная важность (permutation importance) – для определения влияния каждого признака на качество модели. Топ-5 наиболее важных признаков: Init_Win_bytes_backward, Destination Port, Fwd IAT Min, Bwd Packets/s, Total Fwd Packets.

Финальная оценка моделей проводилась на тестовой выборке, что позволило объективно сравнить их эффективность в условиях реалистичного распределения трафика.

3. Результаты и обсуждение

3.1. Сравнительный анализ моделей

Результаты оценки производительности моделей на тестовой выборке (32 860 образцов) представлены в табл. 2. Оценка проводилась по метрикам Accuracy (общая точность), Precision (точность), Recall (полнота) и F1-мере (F1-score) (взвешенный), что позволяет объективно сравнивать модели в условиях сильного дисбаланса классов.

Таблица 2

Сравнительные результаты моделей машинного обучения

Table 2

Comparative Results of Machine Learning Models

Модель	Accuracy	Accuracy	Время обучения, с
Random Forest	0,9953	0,9952	19,57
AutoGluon	0,9960	0,9946	33,96
SVM (RBF)	0,9874	0,9816	2888,41
Logistic Regression	0,9509	0,9678	73,69
Gaussian Naive Bayes	0,7638	0,8603	1,50

Примечание. F1-score рассчитан в взвешенном варианте (weighted), учитывающем дисбаланс классов.

Анализ показывает, что случайный лес (Random Forest) и AutoGluon демонстрируют сопоставимо высокую производительность, но с различными акцентами:

- AutoGluon достигает наивысшей общей точности (Accuracy = 0,9960) – это лучший результат среди всех моделей.

- Random Forest показывает наилучшую F1-меру (0,9952), что указывает на более сбалансированное соотношение между точностью (Precision) и полнотой (Recall).

Таким образом, в зависимости от критерия оценки, лучшей моделью признаются:

- по общей точности (Accuracy) – AutoGluon;
- по F1-мере (F1-score) – случайный лес (Random Forest).

Это расхождение объясняется тем, что AutoGluon, несмотря на более высокую общую точность, демонстрирует чуть более низкую F1-меру из-за повышенного количества ложноположительных срабатываний в редких классах (например, межсайтовый скриптинг (XSS)), что снижает его F1-меру при сохранении высокой полноты (Recall).

AutoGluon построил ансамбль второго уровня WeightedEnsemble_L2, в котором XGBoost получил вес 1,0. Это подчеркивает, что XGBoost оказался наиболее эффективной базовой моделью в условиях данного датасета. Более того, фактически AutoGluon выявил, что ансамбль из одной модели XGBoost является оптимальным решением, что упрощает итоговую модель без потери точности.

Метод опорных векторов (SVM) с RBF-ядром показал хорошее качество (общая точность равна 0,9874), но крайне непрактичное время обучения (2888 с), что делает его непригодным для применения в системах реального времени. Линейные и вероятностные модели (логистическая регрессия, наивный байесовский классификатор) оказались неэффективны, особенно в обнаружении атак, что подтверждает необходимость использования ансамблевых или AutoML-подходов для задач кибербезопасности.

3.2. Анализ важности признаков

Для выявления наиболее информативных сетевых признаков был проведен комплексный анализ, включающий дисперсионный анализ (ANOVA) и перестановочную важность (permutation importance).

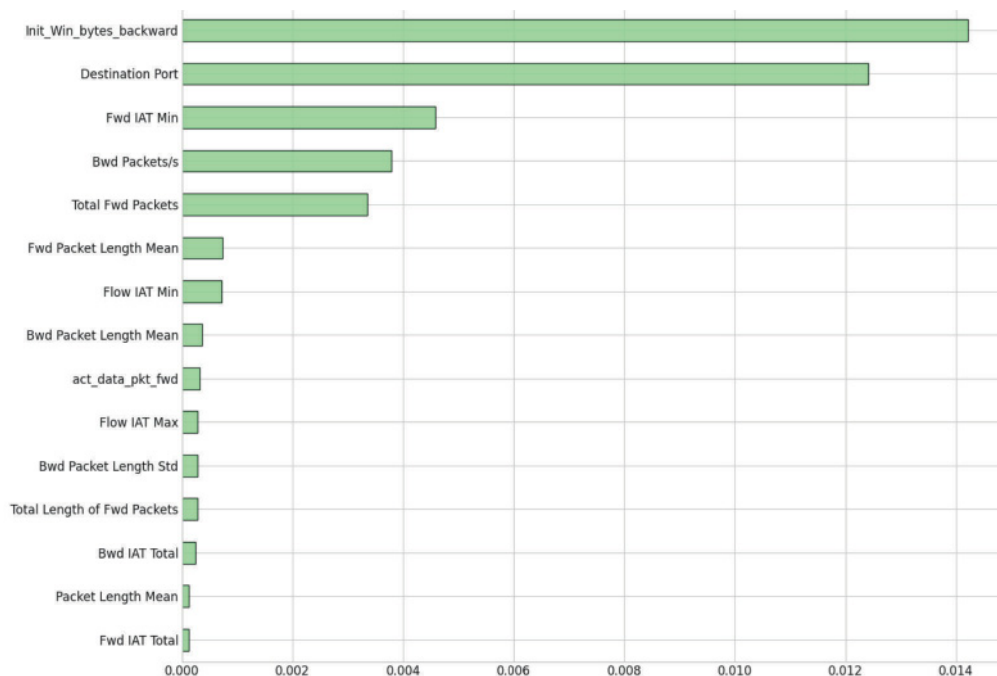


Рис. 1. Важность сетевых признаков для обнаружения атак

Fig. 1. The importance of network features for attack detection

Дисперсионный анализ (ANOVA) с F-критерием показал, что 46 из 62 признаков являются статистически значимыми ($p < 0,001$), т. е. их распределение существенно различается между классами трафика. Это свидетельствует о высокой дискриминативной способности большинства признаков.

Перестановочная важность, вычисленная на модели AutoGluon, выявила следующие топ-5 наиболее важных признаков (рис. 1):

1. Init_Win_bytes_backward (window_size): 0,0142 – размер окна в байтах в обратном направлении при инициализации TCP-сессии. Критичен для выявления аномального поведения на этапе установки соединения.
2. Destination Port (port_related): 0,0124 – целевой порт. Подтверждает, что атаки концентрируются на стандартных сервисах (HTTP, HTTPS, DNS).
3. Fwd IAT Min (forward_timing): 0,0046 – минимальная межпакетная задержка в прямом направлении. Указывает на интенсивность передачи данных, характерную для автоматизированных атак (например, атака перебором).
4. Bwd Packets/s (packet_rates): 0,0038 – количество пакетов в секунду в обратном направлении. Отражает реакцию сервера на атакующий трафик.
5. Total Fwd Packets (packet_stats): 0,0033 – общее количество пакетов в прямом направлении. Часто используется для сканирования и атак с высокой нагрузкой.

Высокая важность временных (Fwd IAT Min), объемных (Init_Win_bytes_backward) и протокольных (Destination Port) признаков указывает на то, что веб-атаки существенно изменяют поведенческий профиль сетевого соединения, что позволяет их выявлять без анализа содержимого пакетов.

3.3. Детальный анализ по классам

Детальный отчет классификации выявляет существенные различия в эффективности обнаружения по классам, что напрямую связано с дисбалансом данных и природой атак (рис. 2).

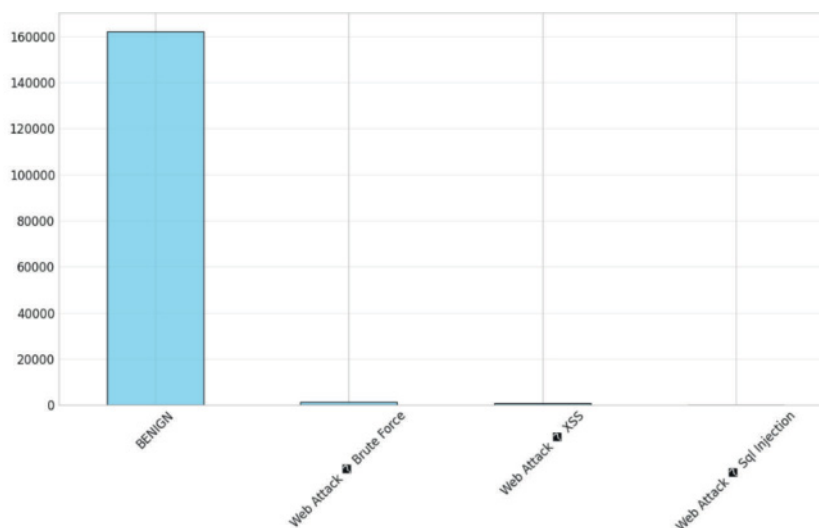


Рис. 2. Распределение типов сетевого трафика

Fig. 2. Distribution of network traffic types

- Класс BENIGN (нормальный трафик): обнаруживается практически идеально. Точность (Precision) = 1,00, полнота (Recall) = 1,00, F1-мера = 1,00. Это ожидаемо, так как модель

обучается в основном на легитимном трафике, и его корректная классификация формирует высокие агрегированные метрики.

- Класс веб-атака – атака перебором (Web Attack – Brute Force): обнаруживается очень эффективно (полнота (Recall) = 0,99), что означает, что почти все такие атаки были выявлены. Однако точность (Precision) = 0,70, что указывает на высокий уровень ложноположительных срабатываний. Это может быть связано с тем, что поведение атаки перебором (множественные запросы к одному порту) иногда имитируется легитимными скриптами или сканерами.
- Класс веб-атака – межсайтовый скриптинг (Web Attack – XSS): обнаруживается крайне плохо. Полнота (Recall) = 0,053 (5,3 %). Несмотря на 652 примера в исходном датасете, модель почти полностью их пропускает. При этом точность (Precision) = 0,88, что говорит о высокой достоверности тех немногих предсказаний, которые были сделаны. Низкая полнота объясняется малым количеством обучающих примеров (521) и способностью XSS-атак маскироваться под нормальный HTTP-трафик.
- Класс веб-атака – SQL-инъекции (Web Attack – SQL Injection): не обнаружен вообще (полнота (Recall) = 0,000). В обучающей выборке было лишь 17 примеров, а в тестовой – 4, что недостаточно для обучения любой модели. Это яркий пример проблемы «редких атак», когда даже высокая общая точность маскирует полную неспособность обнаруживать критически важные угрозы.

3.4. Практические рекомендации

На основе полученных результатов можно сформулировать следующие рекомендации по исследованию датасета CICIDS2017:

1. Для обнаружения атак перебором (Brute Force) – можно использовать AutoGluon или случайный лес (Random Forest) без дополнительной настройки, так как эти атаки обнаруживаются эффективно.
2. Для обнаружения межсайтового скриптинга (XSS) и SQL-инъекций (SQLi) требуется:
 - увеличение обучающей выборки по редким классам (например, через синтез или сбор дополнительных данных);
 - применение методов увеличения выборки (oversampling) (SMOTE, ADASYN);
 - взвешивание классов в процессе обучения;
 - переход к бинарной классификации для каждой редкой атаки отдельно.
3. Для повышения общей эффективности:
 - использовать AutoGluon для автоматического построения высокоточных моделей;
 - применять случайный лес (Random Forest) как интерпретируемую и стабильную альтернативу;
 - учитывать временные и объемные признаки как ключевые для обнаружения аномалий.
4. Для инженерии признаков (feature engineering) – сфокусироваться на разработке и сборе признаков, связанных с поведенческими паттернами (временные задержки, последовательности пакетов, размеры окон), которые оказались наиболее дискриминативными, что не требует анализа полезной нагрузки пакетов и потому не нарушает конфиденциальности передаваемых данных.

3.5. Выводы по разделу

- AutoGluon достиг наивысшей общей точности (Accuracy = 0,9960), превосходя случайный лес (Random Forest).
- Случайный лес (Random Forest) показал наилучшую F1-меру (F1-score = 0,9952) и является оптимальным выбором для сбалансированного обнаружения.

- 46 из 62 признаков являются статистически значимыми, а топ-5 наиболее важных связаны с временными задержками и размером окон.
- Атака перебором (Brute Force) обнаруживается эффективно, межсайтовый скриптинг (XSS) – плохо, SQL-инъекции (SQL Injection) не обнаруживаются вовсе.
- Для построения надежной системы обнаружения вторжений (IDS) необходимо специально оптимизировать модель под редкие атаки, а не полагаться на агрегированные метрики.
- Таким образом, хотя AutoML-фреймворки, такие как AutoGluon, предлагают мощный инструмент для быстрого прототипирования, их эффективность в значительной степени ограничена качеством и сбалансированностью исходных данных.

Заключение

Проведенное исследование было направлено на сравнительную оценку эффективности различных методов машинного обучения в задаче обнаружения сетевых атак на основе датасета CICIDS2017. Экспериментальные результаты, полученные на репрезентативной выборке из 164 300 сетевых потоков, позволяют сформулировать следующие ключевые выводы:

1. Случайный лес (Random Forest) продемонстрировал наилучшее значение F1-меры ($F1\text{-score} = 0,9952$), что свидетельствует о его высокой способности к сбалансированному обнаружению как нормального, так и аномального трафика при сохранении высокой точности, это дополнительно подтверждается в других статьях [8]. При этом время обучения модели составило 19,57 с, что делает ее привлекательной для применения в условиях ограниченных вычислительных ресурсов. Высокая эффективность случайного леса (Random Forest) обусловлена его устойчивостью к шуму, способностью к обработке высокоразмерных данных и минимальной чувствительностью к переобучению.

2. Применение парадигмы автоматизированного машинного обучения (AutoML) с использованием фреймворка AutoGluon Tabular позволило достичь наивысшей общей точности ($\text{Accuracy} = 0,9960$) при полной автоматизации процессов предобработки данных, генерации признаков, подбора моделей и построения ансамблей. Лучшей оказалась композитная модель `WeightedEnsemble_L2`, в которой XGBoost получил вес 1,0, что подчеркивает его доминирующую эффективность на данном датасете. Таким образом, AutoGluon подтвердил свою целесообразность как инструмент для быстрого прототипирования и построения высокоточных моделей без необходимости глубокой настройки.

3. Комплексный анализ важности признаков, включающий дисперсионный анализ (ANOVA) с F-критерием и перестановочную важность (permutation importance), выявил 46 статистически значимых признаков ($p < 0,001$). Наиболее дискриминативными оказались:

- `Init_Win_bytes_backward` (`window_size`),
- `Destination Port` (`port_related`),
- `Fwd IAT Min` (`forward_timing`),
- `Bwd Packets/s` (`packet_rates`),
- `Total Fwd Packets` (`packet_stats`).

Эти признаки, связанные с параметрами инициализации TCP-сессии, целевыми портами и временными задержками, оказались ключевыми для определения легитимного и атакующего поведения, что открывает возможности для построения поведенческих моделей обнаружения аномалий без анализа содержимого пакетов.

4. Несмотря на высокие агрегированные метрики, все исследуемые модели проявили критическую несостоятельность в обнаружении редких атак. Так, эффективность детектирования атак типа межсайтового скриптинга (XSS) составила полноту ($\text{Recall} = 0,053$), а SQL-инъекций (SQL Injection) – полноту ($\text{Recall} = 0$), что обусловлено экстремальным дисбалансом классов

(в обучающей выборке – 521 и 17 примеров соответственно). Это указывает на то, что высокая общая точность может маскировать полную неспособность системы обнаруживать наиболее опасные, но редкие угрозы [9].

Практическая значимость результатов состоит в разработке методики быстрого развертывания эффективных систем обнаружения вторжений (IDS) на основе AutoML, требующей минимальных экспертных знаний и позволяющей достигать высоких показателей производительности даже без ручной настройки моделей.

Перспективы дальнейших исследований связаны:

- с интеграцией методов балансировки классов (SMOTE, ADASYN, GANs) непосредственно в AutoML-конвейеры [10],
- разработкой специализированных признаков и эвристик для детектирования имитационных атак (mimicry-атак) (XSS, SQL-инъекции),
- созданием гибридных систем, сочетающих автоматизированное машинное обучение с сигнатурным анализом и поведенческим мониторингом [11].

Полученные результаты вносят вклад в развитие интеллектуальных систем кибербезопасности и демонстрируют, что автоматизированное машинное обучение является перспективным направлением, но его эффективность напрямую зависит от решения проблемы дисбаланса классов и адекватного представления редких угроз в обучающих выборках.

Список литературы

1. **Sharafaldin I., Lashkari A. H., Ghorbani A. A.** Toward generating a new intrusion detection dataset and intrusion traffic characterization // *Proceedings of the 4th International Conference on Information Systems Security and Privacy (ICISSP)*. 2018. P. 108–116. DOI: 10.5220/0006639801080116. URL: <https://www.unb.ca/cic/datasets/ids-2017.html> (дата обращения: 20.09.2025).
2. **Ahmad Z., Shahid Khan A., Wai Shiang C. et al.** Network intrusion detection system: A systematic study of machine learning and deep learning approaches // *Transactions on Emerging Telecommunications Technologies*, 2021. Vol. 32. No. 1. P. E4150. DOI: 10.1002/ett.4150
3. **Binbusayyis A., Vaiyapuri T.** Identifying and benchmarking key features for cyber intrusion detection: An ensemble approach // *IEEE Access*, 2019. Vol. 7. P. 106495–106513. DOI: 10.1109/ACCESS.2019.2929487
4. **Johnson J., Khoshgoftaar T.** Survey on deep learning with class imbalance // *Journal of Big Data*, 2019, Vol. 6, No. 1, P. 27. DOI: 10.1186/s40537-019-0192-5
5. **Арабов М. К., Седых В. В.** Прогнозирование динамики финансовых временных рядов с помощью методов автоматизированного машинного обучения: случай российского рубля // *Научно-технический вестник Поволжья*. 2025. № 6. С. 199–201.
6. **Arabov M. K., Nazipova A. F., Burnashev R. A.** Algorithm Application of Machine Learning Algorithms to Predict Energy Demand // *Proceedings – 2024 International Conference on Industrial Engineering, Applications and Manufacturing (ICIEAM 2024)*. DOI: 10.1109/ICIEAM60818.2024.10553672
7. **Erickson N. et al.** AutoGluon-Tabular: Robust and Accurate AutoML for Structured Data. arXiv preprint arXiv:2003.06505, 2020. DOI: 10.48550/arXiv.2003.06505
8. **Бабичева М. В., Третьяков И. А.** Применение методов машинного обучения для автоматизированного обнаружения сетевых вторжений // *Вестник Дагестанского государственного технического университета. Технические науки*. 2023. Т. 50, № 1. С. 53–61. DOI: 10.21822/2073-6185-2023-50-1-53-61
9. **Павлычев А. В., Лебедев И. С.** Использование алгоритма машинного обучения Random Forest для выявления сложных компьютерных инцидентов // *Вестник Санкт-Петербург-*

- ского университета. Серия 10: Прикладная математика, информатика, процессы управления. 2022. Т. 18, № 4. С. 567–580. DOI: 10.21681/2311-3456-2022-5-74-81
10. **Yadav P., Singh R. M.** A Systematic Mapping Study of Generative Adversarial Networks for Addressing Class Imbalance in Cybersecurity // *Journal of Cybersecurity and Privacy*. 2025. Vol. 5, No. 1. P. 1–25. DOI: 10.48550/arXiv.2502.16535
 11. **Ferrag M. A., Babaghayou M., Yazici A.** Generative AI in Cybersecurity: A Comprehensive Review of Applications and Challenges // *Computers & Security*. 2025. Vol. 128. P. 103–120. DOI: 10.1016/j.iotcps.2025.01.001

References

1. **Sharafaldin I., Lashkari A. H., Ghorbani A. A.** Toward generating a new intrusion detection dataset and intrusion traffic characterization. *Proceedings of the 4th International Conference on Information Systems Security and Privacy (ICISSP)*, 2018. P. 108–116. DOI: 10.5220/0006639801080116. URL: <https://www.unb.ca/cic/datasets/ids-2017.html> (date of access: 20.09.2025)
2. **Ahmad Z., Shahid Khan A., Wai Shiang C.** et al. Network intrusion detection system: A systematic study of machine learning and deep learning approaches. *Transactions on Emerging Telecommunications Technologies*, 2021, vol. 32, no. 1, pp. e4150. DOI: 10.1002/ett.4150
3. **Binbusayyis A., Vaiyapuri T.** Identifying and benchmarking key features for cyber intrusion detection: An ensemble approach. *IEEE Access*, 2019, vol. 7, pp. 106495–106513. DOI: 10.1109/ACCESS.2019.2929487
4. **Johnson J., Khoshgoftaar T.** Survey on deep learning with class imbalance. *Journal of Big Data*, 2019, vol. 6, no. 1, pp. 27. DOI: 10.1186/s40537-019-0192-5
5. **Arabov M. K., Sedykh V. V.** Prognostirovanie dinamiki finansovykh vremennykh ryadov s pomoshchiu metodov avtomatizirovannogo mashinnogo obucheniya: sluchai rossiiskogo rublya [Forecasting financial time series dynamics using AutoML methods: The case of the Russian ruble]. *Nauchno-tehnicheskii vestnik Povolzh'ya*, 2025, no. 6, pp. 199–201. (in Russ.)
6. **Arabov M. K., Nazipova A. F., Burnashev R. A.** Algorithm Application of Machine Learning Algorithms to Predict Energy Demand. *Proceedings – 2024 International Conference on Industrial Engineering, Applications and Manufacturing (ICIEAM 2024)*. DOI: 10.1109/ICIEAM60818.2024.10553672
7. **Erickson N.** et al. AutoGluon-Tabular: Robust and Accurate AutoML for Structured Data. arXiv preprint arXiv:2003.06505, 2020. DOI: 10.48550/arXiv.2003.06505
8. **Babicheva M. V., Tretyakov I. A.** Primenenie metodov mashinnogo obucheniya dlya avtomatizirovannogo obnaruzheniya setevykh vtorzheniy [Application of machine learning methods for automated detection of network intrusions]. *Herald of Dagestan State Technical University. Technical Sciences*, 2023, vol. 50, no. 1, pp. 53–61. (in Russ.) DOI: 10.21822/2073-6185-2023-50-1-53-61
9. **Pavlychev A. V., Lebedev I. S.** Ispolzovanie algoritma mashinnogo obucheniya Random Forest dlya vyyavleniya slozhnykh kompyuternykh intsidentov [Using the Random Forest Machine Learning Algorithm to Detect Complex Computer Incidents]. *Vestnik Sankt-Peterburgskogo universiteta. Seriya 10: Prikladnaya matematika, informatika, protsessy upravleniya*, 2022, vol. 18, no. 4, pp. 567–580. (in Russ.) DOI: 10.21681/2311-3456-2022-5-74-81
10. **Yadav P., Singh R. M.** A Systematic Mapping Study of Generative Adversarial Networks for Addressing Class Imbalance in Cybersecurity. *Journal of Cybersecurity and Privacy*, 2025, vol. 5, no. 1, pp. 1–25. DOI: 10.48550/arXiv.2502.16535

11. **Ferrag M. A., Babaghayou M., Yazici A.** Generative AI in Cybersecurity: A Comprehensive Review of Applications and Challenges. *Computers & Security*, 2025, vol. 128, pp. 103–120. DOI: 10.1016/j.iotcps.2025.01.001

Информация об авторах

Арабов Муллошараф Курбонович, кандидат физико-математических наук, доцент

Шумихина Анна Сергеевна, магистрант

Information about the Authors

Mullosharaf K. Arabov, Candidate of Physical and Mathematical Sciences, Associate Professor

Anna S. Shumikhina, Master's Student

Статья поступила в редакцию 23.11.2025;

одобрена после рецензирования 11.04.2026; принята к публикации 11.04.2026

The article was submitted 23.11.2025;

approved after reviewing 11.04.2026; accepted for publication 11.04.2026

Научная статья

УДК 004.021

DOI 10.25205/1818-7900-2026-24-1-30-39

Применение нейросетевой модели для классификации и оценки инвестиционных проектов в рамках экономического анализа в ГИС

Полина Сергеевна Браер

Новосибирский государственный университет

Новосибирск, Россия

Braer.a.s@gmail.com

Аннотация

Работа посвящена разработке интеллектуальной геоинформационной системы для анализа инвестиционных проектов на основе автоматизированной обработки текстовой информации. Система интегрирует методы обработки естественного языка, машинного обучения и традиционные методы геопространственного анализа для структурирования информации об экономических объектах из разнородных открытых источников. Основной вклад работы состоит в разработке методологии многоэтапной обработки текстовых данных и реализации нейросетевой архитектуры на основе модели T5, которая обеспечивает одновременное решение задач классификации стадии проекта и предсказания его стоимости. Система демонстрирует высокий потенциал для практического применения в области инвестиционного анализа и управления проектными портфелями. Работа содержит рекомендации по совершенствованию системы, включая применение методов активного обучения, трансферного обучения (transfer learning) и развитие методов интерпретируемости результатов.

Ключевые слова

геоинформационные системы, обработка естественного языка, машинное обучение, инвестиционный анализ, извлечение информации, нейросетевые модели, пространственный анализ

Для цитирования

Браер П. С. Применение нейросетевой модели для классификации и оценки инвестиционных проектов в рамках экономического анализа в ГИС // Вестник НГУ. Серия: Информационные технологии. 2026. Т. 24, № 1. С. 30–39. DOI 10.25205/1818-7900-2026-24-1-30-39

The Use of a Neural Network Model for the Classification and Evaluation of Investment Projects in the Framework of Economic Analysis in GIS

Polina S. Braer

Novosibirsk National Research State University

Novosibirsk, Russian Federation

Braer.a.s@gmail.com

Abstract

This study addresses the challenge of automating information extraction from textual sources to support investment decision-making through geospatial frameworks. An integrated analytical system has been developed that leverages

© Браер П. С., 2026

ISSN 1818-7900 (Print). ISSN 2410-0420 (Online)

Вестник НГУ. Серия: Информационные технологии. 2026. Том 24, № 1

Vestnik NSU. Series: Information Technologies, 2026, vol. 24, no. 1

advances in natural language processing, deep learning, and geographic information science to identify and classify investment projects from diverse open-access documents. The proposed approach employs a sequence-to-sequence neural architecture capable of simultaneously performing categorical analysis (determining project development phases) and numerical prediction (estimating project value). These findings underscore the importance of task-specific optimization strategies in mixed-type information processing. The system demonstrates viable applicability for portfolio assessment, strategic resource allocation, and geospatial economic monitoring. Future development pathways include curriculum-based learning approaches, cross-task knowledge transfer, and enhanced model transparency mechanisms.

Keywords

geographic information systems, natural language processing, machine learning, investment analysis, information extraction, neural network models, spatial analysis

For citation

Braer P. S. The use of a neural network model for the classification and evaluation of investment projects in the framework of economic analysis in GIS. *Vestnik NSU. Series: Information Technologies*, 2026, vol. 24, no. 1, pp. 30–39 (in Russ.) DOI 10.25205/1818-7900-2026-24-1-30-39

Введение

Интеграция технологий искусственного интеллекта (ИИ) в геоинформационные системы (ГИС) представляет собой одно из наиболее перспективных направлений развития методов пространственного анализа экономических явлений [1]. В течение последних пятнадцати лет в России проводится систематизированное изучение геоинформатики и осуществляются целенаправленные разработки по внедрению ГИС в решение разноплановых задач территориального управления, городского планирования, кадастра и анализа инвестиционной активности. Развитие этого направления обусловлено растущей потребностью в инструментах, способных интегрировать разнородные информационные потоки с учетом пространственного измерения хозяйственных процессов.

Традиционные геоинформационные системы предназначены для сбора, анализа и отображения географически привязанных данных, что существенно облегчает поддержку принятия решений в городском планировании, управлении природными ресурсами, управлении инфраструктурой и стратегической бизнес-аналитике. Однако анализ текущих тенденций в области разработки информационных систем указывает на объективную необходимость внедрения методов ИИ, которые способны упростить возможности обнаружения скрытых закономерностей в больших объемах данных, повысить надежность прогнозирования географических явлений и автоматизировать процессы обновления картографических и атрибутивных баз данных [2].

Одной из ключевых проблем анализа инвестиционной активности территорий является то, что данные об инвестиционных проектах трудно получить и систематизировать. Информация о проектах часто разрознена по разным источникам и представлена в неструктурированном виде, из-за чего ее сложно напрямую загружать в ГИС и использовать для пространственного анализа.

В связи с этим задача работы заключается в разработке интеллектуального подхода на основе методов обработки естественного языка и больших языковых моделей, который позволит автоматически извлекать из текстов основные параметры проектов, преобразовывать их в структурированный формат и интегрировать в геоинформационную систему для дальнейшей визуализации и анализа.

Актуальность разработки интеллектуальной системы обусловлена современной выраженной тенденцией к интеллектуализации ГИС и постепенному переходу от описательного анализа к предиктивному и рекомендательному анализу. Несмотря на активное и динамичное развитие технологий ГИС в последние два десятилетия, интеграция методов обработки естественного языка и интеллектуального информационного поиска для автоматического обновления атрибутивных таблиц и актуализации баз данных остается развитой областью исследований [3].

В результате работы разработан интеллектуальный модуль для ГИС, который автоматически извлекает ключевые параметры инвестиционных проектов из неструктурированных текстов и переводит их в структурированный формат для последующей визуализации и анализа. Система решает проблему разрозненности данных и сокращает ручной сбор и обработку информации, повышая актуальность и полноту базы проектов.

Методы применения искусственного интеллекта в геоинформационных системах

Необходимость и целесообразность применения методов ИИ в ГИС обусловлена двумя взаимосвязанными факторами: резким ростом доступности и объемов геопространственных данных благодаря развитию технологий дистанционного зондирования, и одновременным усложнением профессиональных и исследовательских задач, требующих автоматизированного выявления сложных закономерностей в пространстве и времени [4]. При разработке и внедрении ГИС под методами, использующими искусственный интеллект, целесообразно понимать не только традиционные алгоритмы машинного обучения как таковые, но и комплексные аналитические контуры, включающие предварительную подготовку данных, извлечение релевантных признаков, пространственно-статистическое моделирование явлений, формальную оценку неопределенности и встраивание результатов анализа в процессы принятия решений на различных уровнях управления.

Характерной особенностью геоданных является неизбежное наличие пространственной автокорреляции, значительные неоднородности распределений, структурная зависимость качества наблюдений от характеристик датчика и выбранного масштаба анализа, а также проявление эффекта изменяемой территориальной единицы (MAUP). Указанные факторы накладывают специфические и дополнительные требования к выбору и корректной интерпретации ИИ-методов в сравнении с обычными табличными задачами, где пространственная структура данных отсутствует или игнорируется [5].

Наиболее распространенная и хорошо изученная группа методов связана с обучением моделей на структурированных пространственных признаках, представленных векторными слоями и табличными атрибутами объектов. В этой области широко применяются модели контролируемого обучения для выполнения классификации и регрессии: логистическая регрессия, ансамблевые методы (случайный лес, градиентный бустинг с учетом пространственных особенностей), методы опорных векторов (SVM), а также нейронные сети малой и средней глубины. Их практическая ценность в ГИС определяется удачным сочетанием прогностической силы и интерпретируемости результатов, что критически важно для принятия обоснованных управленческих решений.

При применении таких методов классические схемы валидации требуют обязательной адаптации к пространственной природе и структуре данных. Простое случайное разбиение на обучающую и тестовую выборки может неоправданно завышать оценки качества модели из-за пространственной близости объектов и обусловленной этим утечки пространственного контекста. Поэтому в практике ГИС-моделирования целесообразнее использовать пространственную кросс-валидацию, блочные схемы разбиения или проверку на географически отделенных независимых подвыборках.

Вторая значимая группа методов относится к обработке растровых данных дистанционного зондирования Земли, где доминируют подходы компьютерного зрения и глубокого обучения [6]. Сверточные нейронные сети используются для выполнения семантической сегментации, детекции и классификации объектов, инстанс-сегментации, а также для задач обнаружения временных изменений в землепользовании. Встраивание таких моделей в ГИС-про-

цессы решает проблему актуализации картографической основы и уменьшения трудозатрат на ручную дешифровку снимков.

В течение последних пяти лет активно развиваются графовые нейронные сети и подходы к моделированию сетевых структур, позволяющие описывать сложные пространственные зависимости и взаимовлияния через граф, где узлы соответствуют пространственным объектам, а ребра отображают различные типы отношений между ними. Такие модели полезны в задачах оптимизации транспортных систем, анализа связности и надежности инженерной инфраструктуры и оценки уязвимости сложных сетей к техногенным и природным воздействиям.

Особый интерес для данной работы представляют технологии, связанные с обработкой текстовой информации и созданием естественно-языковых интерфейсов к геопространственным данным.

Применение больших языковых моделей в геоинформационных системах

Большие языковые модели (LLM) вроде GPT-4, Claude и их функциональных аналогов открывают принципиально новое направление в развитии и совершенствовании ГИС – возможность взаимодействия с геопространственными данными на естественном языке без необходимости изучения специальных программных интерфейсов. Внедрение LLM в ГИС следует рассматривать не как прямую замену традиционных геоаналитических инструментов, а как эффективное развитие пользовательского слоя и вспомогательных контуров обработки информации, которые ранее требовали либо высокой технической квалификации операторов, либо значительных трудозатрат.

В отличие от классических методов машинного обучения в ГИС, ориентированных преимущественно на обработку числовых признаков и растровых изображений, LLM изначально разработаны и оптимизированы для работы с текстом: для извлечения смысловых единиц и семантических связей, построения связного объяснения результатов, преобразования естественных запросов в формальные языки и поддержания интерактивного диалога между пользователем и системой.

Наиболее очевидный и практически значимый способ применения LLM в ГИС связан с созданием естественно-языкового доступа к пространственным данным и функциям анализа. LLM позволяют преобразовывать формулируемые пользователем запросы на русском или другом естественном языке в последовательность формальных действий: генерацию выражений пространственного SQL для выборки данных, инициирование вызовов геоаналитических алгоритмов, установку параметров инструментов анализа доступности и видимости, разработку сценариев построения тематических карт.

Второй вектор применения LLM связан с извлечением и нормализацией текстовых сведений, релевантных геопространственному анализу и поддержке принятия решений. Реальные ГИС-проекты в экономической сфере опираются на массивы разнообразной документации: технико-хозяйственные паспорта объектов, инвестиционные меморандумы, отчеты о социально-экономическом развитии территорий, нормативные акты и нормативно-правовые документы. LLM позволяют автоматически формировать структурированные представления содержания таких документов, выполняя извлечение сущностей (названия проектов, организаций, суммы финансирования, географические объекты) и устанавливая семантические отношения между ними.

В третьей группе сценариев применения LLM позволяют автоматизировать создание семантического описания наборов геоданных, формировать краткие информативные паспорта географических слоев, извлекать из технической документации критичную информацию о качестве данных, ограничениях применимости и рекомендуемых режимах использования.

Практика интеграции языковых моделей в ГИС отражает применимость LLM к типовым задачам пространственного анализа: построению SQL-запросов к геоданным, классификации атрибутов объектов, автоматизированному формированию картографических описаний и генерации программных скриптов для обработки данных в ГИС [7]. В зарубежной практике языковые модели (BERT, RoBERTa, T5) интегрированы в инструментарий платформ для задач классификации текстов, извлечения именованных сущностей и перевода [8; 9]. В части решения задач извлечения информации об инвестиционных проектах из текстовых источников и последующей интеграции результатов в атрибутивные таблицы ГИС-слоев на данный момент наблюдается недостаточность исследований, что определяет актуальность настоящей работы.

Однако практическое применение LLM в ГИС выявляет методологические ограничения, требующие учета при разработке систем. Языковые модели по своей внутренней природе ориентированы на правдоподобную и связную генерацию текста, а не на гарантированную строгую истинность утверждений, что порождает существенный риск «галлюцинаций» – уверенно и грамотно сформулированных, но фактически неверных утверждений и выводов. Для многих геопространственных задач критически важно корректное обращение с системами координат, единицами измерения, топологическими отношениями и выбранным масштабом анализа; эти аспекты не всегда устойчиво и корректно выражаются в естественном языке и требуют строгих формальных проверок.

Практическая реализация LLM в ГИС неизбежно предполагает многоуровневый контроль формата выходных данных, валидацию результатов работы внешними правилами и логическими ограничениями, систематическое сопоставление с первичными источниками информации и верификацию результатов.

Наиболее перспективным и обоснованным представляется инструментальный подход к применению LLM в ГИС, когда языковая модель рассматривается как интеллектуальный оркестратор и интерпретатор, который последовательно переводит намерения и запросы пользователя в формальные действия, опирается на контролируемое извлечение релевантных контекстов из защищенных баз знаний и корпоративных каталогов, а также формирует пояснительные объяснения в соответствии с заданными стандартами и требованиями.

Разработка системы сбора данных и проектирование интерфейса Постановка проблемы

В ходе проведенного исследования была разработана многоуровневая система автоматического сбора и предварительной обработки данных об инвестиционных проектах, хранящихся в базе данных ГИС ИЭОПП СО РАН, различного масштаба и типов финансирования. Система веб-парсинга и мониторинга осуществляет непрерывный контроль целевых веб-ресурсов, включая официальные государственные базы инвестиционных проектов, портал Министерства инвестиций, промышленности и торговли Российской Федерации, региональные информационные системы и порталы развития, поддерживаемые органами власти субъектов Российской Федерации, а также новостные порталы и специализированные экономические издания.

Практическая реализация и тестирование системы веб-парсинга выявили ряд существенных системных проблем и ограничений. Официальные государственные базы данных инвестиционных проектов обновляются нерегулярно, часто с существенными и непредсказуемыми временными задержками, что приводит к отставанию информации о текущем статусе реализуемых проектов.

Новостные издания и медиаресурсы не обеспечивают равномерное и полное покрытие сведений о большинстве инвестиционных проектов. Публикации о проектах часто концентрируются вокруг ключевых событий (объявление о запуске, достижение важных этапов, возник-

новение проблем), между которыми может пройти значительное время без каких-либо новостных упоминаний.

Региональные информационные системы и агентства развития, как правило, ведут свои базы данных с использованием различных стандартов структурирования информации, несовместимыми форматами представления данных и неоднородными схемами обновления и поддержки. Один и тот же инвестиционный проект может быть описан в различных источниках с использованием совершенно разных формулировок, неоднозначных аббревиатур, вариативных названий компаний и объектов, что требует от системы высокой степени обобщения и устойчивости к вариативности текстов.

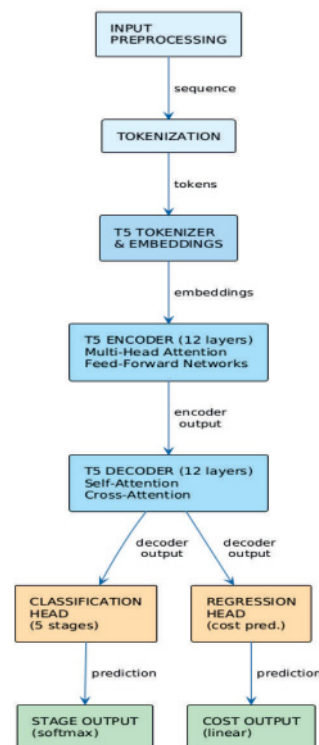
Несмотря на указанные объективные ограничения и вызовы, качество реализованного механизма веб-парсинга было последовательно улучшено и оптимизировано в течение периода практической работы.

Архитектура модели и этапы обработки информации

Разработанная система реализует трехэтапный механизм комплексной обработки информации. На первом этапе осуществляется сбор и консолидация текстовых данных посредством веб-парсинга целевых информационных ресурсов, найденных посредством запроса, состоящего из наименования проекта, отрасли и региона проекта, полученных из ГИС, и сохранение собранной информации в унифицированном текстовом формате с метаданными. На этом этапе производится автоматический краулинг и индексирование целевых веб-ресурсов, выделение релевантных текстовых блоков, содержащих информацию о проектах, фильтрация и очистка от технических элементов HTML-разметки и скриптов, нормализация кодировки символов для обеспечения совместимости с последующими этапами обработки [10].

На втором этапе осуществляется предварительная обработка и подготовка текстов перед подачей их в нейросетевую модель машинного обучения. Процесс предварительной обработки включает нормализацию текста путем преобразования всех символов в нижний регистр, удаление лишних пробельных символов и знаков пунктуации, унификацию форматов представления числовых значений, особенно денежных сумм с разными валютами и масштабами. После нормализации выполняется токенизация текста с использованием специально обученного токенизатора SentencePiece на базе алгоритма Byte Pair Encoding (BPE). Токенизатор содержит словарь из 32 000 подслов (субтокенов), оптимизированный для работы с русскоязычными текстами в предметной области инвестиционного проектирования. Максимальная длина входной последовательности ограничена 512 токенами.

Архитектура нейросетевой модели, представленная на рисунке, основана на T5 с использованием генеративного подхода T5ForConditionalGeneration из библиотеки HuggingFace Transformers. Выбор T5-модели в качестве основной архитектуры нейросетевой системы был обоснован ее универсальностью и адаптируемостью к различным задачам обработки текста в режиме seq2seq. В сравнении с альтернативными подходами (использование BERT для кодирования, отдельные модели классификации и ре-



Архитектура нейросетевой модели
Architecture of the neural network model

грессии), архитектура T5 предоставляет ряд методологических преимуществ, таких как унифицированный интерфейс для различных типов задач, возможность эффективного трансферного обучения (transfer learning), и относительная стабильность при работе с ограниченными объемами размеченных данных.

Модель реализует полносвязную архитектуру encoder-decoder, где энкодер состоит из шести трансформер-слоев размерности 512 элементов, содержащий промежуточный слой с 2048 нейронами для нелинейных преобразований. Энкодер обрабатывает входной текст двупольно, создавая мощные контекстные представления каждого токена с учетом всего входного документа и его смысловой структуры. Декодер также содержит шесть слоев трансформера с аналогичной конфигурацией, но использует однонаправленное внимание для автоматической генерации выходной последовательности в режиме, пригодном для реального времени. Длина генерируемой последовательности ограничивалась 256 токенами с использованием алгоритма decoding с beam search размером луча, равным 2, что обеспечивает баланс между качеством и скоростью генерации.

Результаты работы нейросетевого модуля – извлеченные стадия и стоимость проекта – передаются в атрибутивную таблицу ГИС-слоя. Привязка к пространству осуществляется через геокодирование: топонимы и названия регионов, извлеченные из текста, сопоставляются с геометрией административных единиц из базы данных ГИС.

Процесс обучения и стратегия валидации

Обучение модели проводилось на наборе данных, содержащем 200 000 примеров пар «текст – ответ» в структурированном формате JSONL, с разделением на обучающую выборку (80 %), валидационную выборку (10 %) и тестовую выборку (10 %) без утечки данных. Набор включает как сгенерированные данные на основе шаблонов и с использованием LLM (80 % данных), так и реальные тексты об инвестиционных проектах из официальных источников, размеченные вручную (20 % данных). Выбор преимущественно в пользу синтетических данных был обусловлен отсутствием массово размеченного корпуса реальных публикаций, однако дополнительное использование реальных примеров позволило оценить применимость модели вне контролируемой обучающей среды.

Каждый обучающий пример состоит из входного текста с описанием инвестиционного проекта и ожидаемого выхода в формате JSON с извлекаемыми параметрами: стадия проекта (инициирование, подготовка, проектирование, экспертиза, строительство, введен в эксплуатацию, приостановлен или отменен) и стоимость проекта в рублях. Дополнительно для внешней проверки применимости подхода использовались реальные тексты из открытых источников, что позволило оценить поведение модели вне синтетического корпуса.

Процесс обучения модели осуществлялся с использованием алгоритма оптимизации Adam с параметром весового распада (L2-регуляризация) в размере 0,01 и линейным прогревом обучения (learning rate warmup) в течение первых 3 % шагов от общего количества итераций для стабилизации процесса обучения. Валидация модели проводилась каждые 1000 шагов обучения с фиксированием лучшей модели на основе валидационной метрики.

Для объективной оценки качества разработанной модели использовались три основные метрики: stage_ассигасу (доля примеров из тестового набора, где модель корректно определила стадию реализации проекта), price_ассигасу (доля примеров, где модель правильно извлекла информацию о стоимости проекта в исходном формате или в нормализованном виде) и joint_ассигасу (доля примеров, где оба параметра были извлечены корректно и полностью).

Результаты обучения и анализ производительности

Динамика развития метрик качества в процессе обучения модели демонстрирует стабильное и монотонное повышение точности по всем извлекаемым параметрам на протяжении 50 000 итераций обучения. На этапе полного обучения были достигнуты следующие результаты: stage_accuracy составила 96,2 %, price_accuracy – 89,7 %, и joint_accuracy – 86,4 % на независимой тестовой выборке. При этом на реальных размеченных данных из 20 проектов stage_accuracy составила 96,0 %, price_accuracy – 85,3 %, joint_accuracy – 85,7 %.

Наблюдаемое различие в точности между предсказанием стадии проекта (stage_accuracy > 96 %) и предсказанием стоимости (price_accuracy > 85%) объясняется фундаментальными различиями в характере этих двух задач. Классификация стадии проекта представляет собой задачу многоклассовой классификации с дискретным набором из восьми заданных классов, каждый из которых имеет четкие и легко распознаваемые лингвистические маркеры в тексте проектов. Например, стадия «строительство» часто упоминается в текстах рядом со специфическими ключевыми словами и фразами («ведется строительство», «строительно-монтажные работы», «введены в эксплуатацию первые объекты»), которые модель может относительно легко выучить и распознать при анализе новых текстов.

Извлечение стоимости проекта из текста представляет собой существенно более сложную задачу, требующую не только надежного распознавания числовых значений, но и корректной интерпретации единиц измерения (млн рублей, млрд рублей, доллары), различных масштабных коэффициентов и непредсказуемых форматов представления денежных сумм в исходных текстах, а также глубокого понимания контекстной информации о проекте. Несмотря на эти существенные сложности и вызовы, достигнутая точность совместного предсказания (joint_accuracy > 85 %) демонстрирует полную пригодность модели для практического применения в системах автоматизированной обработки инвестиционных данных.

Проведенный анализ ошибок модели выявил, что в 7,3 % случаев модель неправильно интерпретировала единицы измерения стоимости. В 4,1 % случаев стадию проекта определила неверно из-за наличия в тексте нескольких упоминаний различных этапов; в 2,2 % случаев модель пропустила релевантную информацию при превышении максимальной длины входной последовательности.

Заключение

В результате работы реализован интеллектуальный модуль для ГИС, обеспечивающий автоматическое извлечение ключевых параметров инвестиционных проектов из неструктурированных текстовых открытых источников и преобразование их в единый структурированный формат, пригодный для загрузки в атрибутивные таблицы ГИС и дальнейшего пространственного анализа. Разработанный подход позволяет минимизировать ручной сбор данных, повысить актуальность базы проектов, а также обеспечить их последующую визуализацию на цифровой карте и использование в аналитических сценариях оценки инвестиционной активности территорий.

Предложенный подход, основанный на многоэтапной обработке текстовых данных и применении нейросетевых архитектур, позволяет эффективно решать задачи структурирования и классификации информации из разнородных источников. Достигнутые результаты на реальных данных – точность классификации стадии проекта на уровне 96,0 % и средняя точность совместного предсказания стоимости в 85,3 % – свидетельствуют о практической применимости предложенного метода.

Анализ полученных результатов выявил закономерность в различиях точности при обработке разных типов информации. Категориальная информация извлекается с более высокой точностью (96,0 %) по сравнению с численной информацией (85,3 %), что объясняется фунда-

ментальными различиями в сложности этих задач. Численная информация требует не только правильной идентификации соответствующих фрагментов текста, но и их корректной нормализации с учетом различных форматов представления и единиц измерения.

Разработанная система демонстрирует потенциал для практического применения в области инвестиционного анализа, однако следует отметить ограничения, которые необходимо учитывать при практическом внедрении системы. Система требует наличия достаточного объема текстовой информации об анализируемых объектах, что может представлять проблему для проектов в ранних стадиях развития, а также для объектов, находящихся в регионах с низким уровнем представленности в открытых источниках.

Совершенствование системы требует систематического накопления и разметки примеров из реальных источников, что позволит повысить адаптивность модели к разнообразию текстовых представлений одного и того же типа информации. Применение методов активного обучения для оптимизации процесса разметки данных может существенно сократить трудозатраты при расширении обучающего набора.

Список литературы

1. **Абдурахманов К. Х.** Искусственный интеллект – основа устойчивого развития экономики: моногр. / Рос. экон. ун-т им. Г. В. Плеханова. М.: Изд-во РЭУ им. Г. В. Плеханова, 2023. 298 с.
2. **Аманкулова Н. А., Молмакова М. С., Каримова Г. Т.** Искусственный интеллект и геоинформационные системы // Бюллетень науки и практики. 2023. № 11. С. 278–287.
3. **Берлянт А. М.** Геоинформационное картографирование: моногр. М.: Астрейя, 1997. 64 с.
4. **Баранов Ю. Б., Берлянт А. М., Капралов Е. Г.** и др. Геоинформатика. Толковый словарь основных терминов. М.: ГИС-Ассоциация, 1999. 204 с.
5. **Миронова Ю. Н.** Геоинформационные системы // Актуальные проблемы гуманитарных и естественных наук. 2014. № 3 (62). С. 63–65.
6. **Хазратов Ф. Х.** Современные проблемы интеграции геоинформационных систем и интернет-технологий // Universum: технические науки: электрон. научн. журн. 2020. № 9 (78).
7. **Шайтупа С. В.** Интеллектуальный анализ геоданных // ПНиО. 2015. № 6 (18). С. 24–30.
8. **Devlin J., Chang M.-W., Lee K., Toutanova K.** BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding // Proceedings of NAACL-HLT 2019. P. 4171–4186.
9. **Burrough P. A., McDonnell R. A.** Principles of geographical information systems. Oxford: Oxford University Press, 1998, 333 p.
10. **Терентьева Н. А.** Прикладное значение ГИС-технологий // Научное сообщество студентов: Междисциплинарные исследования: сб. ст. по материалам LXXXIII Междунар. студ. науч.-практ. конф. № 24 (83). 2019.

References

1. **Abdurakhmanov K. H.** Iskusstvennyy intellekt – osnova ustoichivogo razvitiya ekonomiki: monografiya. Moscow: Publishing house of REU named after G. V. Plekhanov, 2023, 298 p. (in Russ.)
2. **Amankulova N. A., Molmakova M. S., Karimova G. T.** Iskusstvennyj intellekt i geoinformacionnye sistemy. *Byulleten' nauki i praktiki*. 2023, no. 11, pp. 278–287 (in Russ.)
3. **Berliant A. M.** Geoinformatsionnoe kartografirovaniye: monografiya. Moscow: Astreya, 1997, 64 p. (in Russ.)
4. **Baranov Yu. B., Berliant A. M., Kapralov E. G. et al.** Geoinformatika. Tolkoviy slovar osnovnikh terminov. Moscow: GIS-Association, 1999, 204 p. (in Russ.)

5. **Mironova Yu. N.** Geoinformatsionnyye sistemy. *Aktualnyye problemy gumanitarnykh i estestvennykh nauk*, 2014, no. 3 (62), pp. 63–65 (in Russ.)
6. **Khazratov F. Kh.** Sovremennyye problemy integratsii geoinformatsionnykh sistem i internet-tekhnologiy. *Universum: tekhnicheskiye nauki: elektronnyy nauchnyy zhurnal*, 2020, no. 9(78) (in Russ.)
7. **Shajtura S. V.** Intellektual'nyj analiz geodannyh. *PNiO*. 2015, no. 6 (18). P. 25.
8. **Devlin J., Chang M.-W., Lee K., Toutanova K.** BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. *Proceedings of NAACL-HLT 2019*, pp. 4171–4186.
9. **Burrough P. A., McDonnel R. A.** Principles of geographical information systems. Oxford: Oxford University Press, 1998, 333 p.
10. **Terenteva N. A.** Prikladnoye znachenije GIS-tekhnologiy. *Nauchnoye soobshchestvo studentov: mezhdistsiplinarnyye issledovaniya: sbornik statey*. 2019, no. 24(83) (in Russ.)

Информация об авторе

Браер Полина Сергеевна, студентка

Information about the Author

Polina S. Braer, Student

*Статья поступила в редакцию 25.11.2025;
одобрена после рецензирования 10.02.2026; принята к публикации 10.02.2026*
*The article was submitted 25.11.2025;
approved after reviewing 10.02.2026; accepted for publication 10.02.2026*

Научная статья

УДК 004.056

DOI 10.25205/1818-7900-2026-24-1-40-54

Виртуальная лаборатория по защите информации на объекте информатизации

**Марина Владимировна Тумбинская
Булат Ильнурович Гатауллин
Эндже Ильдаровна Хаерова**

Казанский национальный исследовательский технический университет
им. А. Н. Туполева – КАИ
Казань, Россия

tumbinskaya@inbox.ru, <https://orcid.org/0000-0003-3738-5242>

mr.bulgat12@mail.ru, <https://orcid.org/0009-0004-0202-0117>

engikhaer@gmail.com, <https://orcid.org/0009-0001-1872-2189>

Аннотация

Рост киберугроз обуславливает необходимость совершенствования методов подготовки специалистов в области информационной безопасности. В статье предложена виртуальная лаборатория, основанная на технологии виртуальной реальности (Unity3D), для моделирования практических кейсов по защите информации на объекте информатизации. Лаборатория включает сценарии по противодействию утечкам по техническим каналам, защите сетевого периметра и физической безопасности. Разработаны алгоритмы интерактивного взаимодействия, система автоматизированной проверки корректности решений и обратной связи. Экспериментальное исследование подтвердило эффективность разработки: среднее количество верно установленных средств защиты информации при использовании только виртуальной лаборатории составило на 6 % больше, чем при использовании комбинированного подхода (виртуальная лаборатория совместно с методическими указаниями).

Ключевые слова

виртуальная лаборатория, объект информатизации, технический канал утечки информации, программно-аппаратные средства защиты информации, физические средства защиты

Для цитирования

Тумбинская М. В., Гатауллин Б. И., Хаерова Э. И. Виртуальная лаборатория по защите информации на объекте информатизации // Вестник НГУ. Серия: Информационные технологии. 2026. Т. 24, № 1. С. 40–54. DOI 10.25205/1818-7900-2026-24-1-40-54

© Тумбинская М. В., Гатауллин Б. И., Хаерова Э. И., 2026

ISSN 1818-7900 (Print). ISSN 2410-0420 (Online)

Вестник НГУ. Серия: Информационные технологии. 2026. Том 24, № 1

Vestnik NSU. Series: Information Technologies, 2026, vol. 24, no. 1

Virtual Laboratory for Information Protection at the Informatization Facility

**Marina V. Tumbinskaya, Bulat I. Gataullin,
Endje I. Haerova**

Kazan national research technical university named after A. N. Tupolev – KAI
Kazan, Russian Federation

tumbinskaya@inbox.ru, <https://orcid.org/0000-0003-3738-5242>

mr.bulgat12@mail.ru, <https://orcid.org/0009-0004-0202-0117>

engikhaer@gmail.com, <https://orcid.org/0009-0001-1872-2189>

Abstract

Currently, due to the digitalization of processes and the growing number of cyber threats, many companies are in need of competent specialists with the necessary information security skills. The article offers a modern digital software product based on virtual reality technology. Virtual simulators can be actively used in the educational process. The use of such technologies makes it possible to carry out pilot industrial work in an environment as close to reality as possible. The virtual laboratory is a set of interactive models in a three-dimensional virtual space and has a wide range of practical tasks aimed at timely response to information security incidents and their prevention. The article presents an analysis of existing software tools and virtual simulators in the field of Information Security. Conclusions are drawn about the need to develop a virtual laboratory for information security at the informatization facility. The algorithm of the user's work with the virtual laboratory is presented. The interface forms and fragments of the virtual laboratory program code are described. The paper presents the stages of program development and the results of testing a virtual laboratory for information security at an informatization facility. The testing of the virtual laboratory confirmed the importance of using the simulator in the educational process as a tool for obtaining professional competencies. Based on the results of the statistical data of the experiment, it can be concluded that the use of an exclusively virtual laboratory has reduced the time spent on laboratory work by 30 minutes, while improving the quality of the absorbed material.

Keywords

virtual simulator, informatization facility, technical information leakage channel, software and hardware information protection tools, physical protection tools

For citation

Tumbinskaya M. V., Gataullin B. I., Haerova E. I. Virtual laboratory for information protection at the informatization facility. *Vestnik NSU. Series: Information Technologies*, 2026, vol. 24, no. 1, pp. 40–54 (in Russ.) DOI 10.25205/1818-7900-2026-24-1-40-54

Введение

На сегодняшний день получение знаний и навыков по различным областям нацелено на использование цифровых технологий. Все чаще в учебном процессе используют виртуальные лаборатории, AR-тренажеры, VR-симуляторы в таких областях, как энергетика, промышленность, авиация и др. [1]. На сегодняшний день популярность виртуальной и дополненной реальностей растет. Об активном развитии таких технологий свидетельствует исследование, проведенное компанией «ТМТ Консалтинг». По последним исследованиям, рынок виртуальной реальности в России в ближайшие годы будет расти в среднем на 37 % в год [2]. Такие прогнозные оценки позволяют сделать вывод, что виртуальная реальность займет важное место в тренде развития передовых технологий, в том числе в области образования. При подготовке специалистов в области защиты информации также могут применяться виртуальные тренажеры, позволяющие проводить опытно-промышленные работы в среде, максимально приближенной к реальности. Целью исследования является повышение качества подготовки специалистов в области защиты информации за счет разработки и практического применения виртуальной лаборатории, позволяющей имитировать решение профессиональных задач.

В статье описывается подход к разработке виртуальной лаборатории на базе платформы Unity3D для поддержки образовательного процесса в области технической защиты информа-

ции. Описан алгоритм взаимодействия пользователя с виртуальной лабораторией для трех сценариев практической подготовки (защита от утечки по техническим каналам, защита сетевого периметра, физическая защита объекта) в единой виртуальной среде. Проведено экспериментальное исследование, подтверждающее, что использование разработанной виртуальной лаборатории способствует повышению качества обучения за счет повышения точности выполнения кейсов и снижения частоты ошибок. Данный подход особенно актуален при подготовке специалистов по технической защите информации, поскольку традиционное практическое обучение связано с ограничениями: сложностью доступа к реальным объектам и недостаточной оснащенностью лабораторий современными средствами защиты информации. Реализация подхода позволит эффективно формировать необходимые профессиональные компетенции.

Анализ виртуальных тренажеров в области информационной безопасности

Результаты анализа литературных источников и патентного поиска показывают, что рынок технологий виртуальной и дополненной реальности в России растет с каждым годом. На сегодняшний день существуют различные виды виртуальных тренажеров для повышения компетенций в сфере информационной безопасности. К примеру, виртуальный тренажер «Технология VPN» [3], виртуальный учебник «Защита информации от искажения при передаче» [4], демонстрационная виртуальная среда «Монтаж средств ТЗИ» [5], виртуальный тренажер «Системы видеонаблюдения» [6], виртуальный тренажер «Монтаж телекоммуникационной стойки и сетевого оборудования» [7], виртуальный комплекс «Защита объекта от утечек информации по техническим каналам» [8], виртуальный комплекс «Обнаружение скрытых устройств и скрытых видеокамер» [9], виртуальный тренажер «Основы квантовой криптографии» [10].

В таблице сравнительного анализа виртуальных тренажеров в области информационной безопасности на основе патентного поиска представлены их ключевые особенности.

В процессе сравнительного анализа существующих виртуальных тренажеров, используемых для подготовки специалистов по информационной безопасности, были определены критерии сравнения: охват предметной области, количество изучаемых средств защиты информации (СрЗИ), наличие теоретического блока для изучения материала, тип проверки результатов и взаимодействия, а также доступность решения (платное или бесплатное). Сравнительный анализ существующих виртуальных тренажеров показал, что большинство проанализированных средств являются коммерческими решениями в области защиты информации, носят узкоспециализированный характер, обладают ограниченной иммерсивностью (реализованы преимущественно в формате 2D-симуляторов либо 3D-сред) без интеграции теоретического и практического блоков. Следовательно, можно сделать вывод о том, что целесообразность разработки специализированной виртуальной лаборатории, ориентированной на объект информатизации как средства повышения практических навыков и формирования профессиональных компетенций будущих специалистов в области информационной безопасности, является актуальной. Предмет исследования – виртуальный тренажер по защите информации на объекте информатизации. Методы исследования – методы системного анализа, моделирования. Практическая значимость работы заключается в применении виртуальной лаборатории в учебном процессе.

Разработка виртуальной лаборатории по защите информации на объекте информатизации

Для облегчения и ускорения работы пользователей необходимо, чтобы интерфейс виртуальной лаборатории был универсальным в использовании и предоставлял качественную ви-

Анализ виртуальных тренажеров в области информационной безопасности

Analysis of virtual simulators in the field of information security

Название	Охват предметной области	Количество изучаемых средств защиты информации	Наличие теоретического блока обучения	Тип проверки результатов	Тип взаимодействия с пользователем	Доступность решения
1	2	3	4	5	6	7
Виртуальный тренажер «Технология VPN»	Узкий (только VPN)	0 (объект изучения – программная технология и ее настройка)	Да	Автоматическая (без подсветки неправильно настроенного элемента)	2D-симулятор	Платно
Виртуальный учебник «Защита информации от искажения при передаче»	Узкий (изучения основ криптографии)	0 (модели алгоритмов)	Да	Нет	2D-симулятор	Платно
Виртуальный тренажер «Монтаж средств ТЗИ»	Узкий (навыков корректного монтажа средств защиты информации по (вибро-) акустическому каналу)	5	Нет	Автоматическая (уведомление о неверном размещении СрЗИ)	3D-среда	Платно
Виртуальный тренажер «Системы видеонаблюдения»	Узкий (монтаж, проектирование систем видеонаблюдения)	5	Нет	Автоматическая (уведомление о верном и неверном размещении СрЗИ)	3D-среда	Платно
Виртуальный тренажер «Монтаж телекоммуникационной стойки и сетевого оборудования»	Узкий (проектирование и монтаж стандартных телекоммуникационных стоек)	3–4	Да	Автоматическая (список допущенных обучаемым ошибок)	3D-среда	Платно
Виртуальный комплекс «Защита объекта от утечек информации по техническим каналам»	Узкий (исключительно технические каналы утечки информации)	10	Нет	Автоматическая (запись ошибок, допущенных в процессе)	3D-среда	Платно

Окончание табл.

1	2	3	4	5	6	7
Виртуальный тренажер «Обнаружение закладных устройств и систем видеонаблюдения»	Узкий (методы обнаружения закладных устройств и скрытых видеокamer)	8	Нет	Автоматическая (запись ошибок, допущенных в процессе)	3D-среда	Платно
Виртуальный тренажер «Основы квантовой криптографии»	Узкий (квантовая криптография)	0	Да	Автоматическая (запись ошибок, допущенных в процессе)	2D-симулятор	Платно

зуальную информацию [11]. В программном продукте реализовано два вида взаимодействия с пользователем. Первый вид взаимодействия является визуальным с использованием отображаемых кнопок на панели управления процессом обучения. Второй вид взаимодействия – это нативный, с использованием клавиш на клавиатуре, которые принято использовать для действий в игровом сообществе [12].

Процесс разработки виртуальной лаборатории был разделен на три основных этапа. На первом этапе была разработана система инвентаря для имитации подбора предметов и добавления их в виртуальный объект информатизации. На втором этапе была выполнена разработка и программная реализация модели, обеспечивающей размещение средств защиты информации в виртуальной среде и автоматизированную оценку корректности их установки. Третий этап включал в себя проектирование и разработку визуальных составляющих практического и теоретического блоков виртуальной лаборатории. Виртуальная лаборатория была разработана на платформе Unity3D с разделением событий на отдельные скрипты на языке программирования C#.

При разработке виртуальной лаборатории были проанализированы современные подходы к моделированию и анализу технических систем, в том числе методы, основанные на искусственных нейронных сетях [13–15], что позволило разработать многоуровневую систему автоматизированной проверки корректности размещения средств защиты информации, сочетающую анализ пространственного расположения объектов, учет их функционального назначения и соответствие заданным сценариям угроз. Компоненты виртуальной лаборатории позволяют моделировать процесс выбора и размещения средств защиты информации в виртуальной среде с последующей оценкой.

Разработанная виртуальная лаборатория поддерживает два режима работы: автономный – на базе персонального компьютера с традиционным интерфейсом (клавиатура, мышь), и иммерсивный – с применением системы виртуальной реальности Pico 4, обеспечивающей эффект присутствия и естественную механику взаимодействия с объектами виртуальной среды.

Организация работы пользователя в виртуальной лаборатории:

1. Выбор сценария. Студент запускает лабораторию и выбирает одно из трех виртуальных пространств для моделирования, т. е. виртуальная лаборатория состоит из отдельных модулей:

- защита информации от утечки по техническим каналам;
- программно-аппаратная защита сетевого периметра;
- физическая защита информации на объекте информатизации.

Обучающимся предлагается типовой набор виртуальных помещений (две серверные, два офисных помещения, две лаборатории), что позволяет студентам сконцентриро-

ваться на алгоритме действий, изучении средств защиты информации и их практическом применении.

2. Размещение средств защиты. Студент анализирует угрозы, каналы утечки информации и выбирает необходимые средства защиты информации из каталога. Требуется корректно разместить выбранные средства (например, камеры, датчики) на виртуальном объекте. При размещении средств защиты информации значения настраиваемых параметров их функционирования фиксируются.

3. Проверка решения. Система автоматически проверяет результат. Проверка решений в виртуальной лаборатории реализована путем сопоставления с эталонными решениями. Для каждой виртуальной комнаты экспертным путем (преподавателями-предметниками) определено эталонное решение по размещению средств защиты, основанное на действующих нормативных документах ФСТЭК России, ФСБ России, а также на положениях ГОСТ Р 50922 и руководящих документах (в том числе РД 78.36.003-2002). В разработанной виртуальной лаборатории заложена возможность множественности правильных решений. В базе данных системы хранятся несколько эталонных сценариев для каждого практического задания. Например, для контроля периметра помещения допустимы различные варианты размещения датчиков движения (в углах или на стенах) при условии обеспечения требуемой зоны обнаружения. Ошибочным считается отсутствие средства защиты в необходимой зоне либо выбор средства, не соответствующего типу угрозы (например, установка датчика движения вместо магнитоконтактного датчика для контроля открывания двери).

4. Обратная связь.

– «Корректно»: правильно установленные средства защиты информации подсвечиваются зеленым.

– «Некорректно»: ошибочные элементы выделяются красным с подсказками по исправлению.

5. Студент вносит коррективы и повторяет проверку до полного устранения ошибок.

Алгоритм взаимодействия обучаемого с программным обеспечением состоит из следующих шагов:

1. Погружение в виртуальную среду объекта информатизации.

2. Анализ уязвимостей – обследование инфраструктуры на наличие угроз информационной безопасности.

3. Выбор и изучение средств защиты информации.

4. Подготовка и установка – перенос в активную зону, выбор точек размещения и фиксация средств защиты информации.

5. Цикл – шаги 3–4 повторяются для всех необходимых средств.

6. Проверка и коррекция результата, исправление ошибок.

7. Завершение – задание считается выполненным после успешной проверки.

В виртуальной лаборатории представлено более 20 технических, программно-аппаратных и физических средств защиты информации, таких как виброакустический излучатель, коммутатор, датчик разбития стекла, камеры видеонаблюдения, магнитоконтактный датчик – геркон, датчик движения, межсетевой экран и другие. В работе представлены примеры внешнего вида технических средств (рис. 1–3).

В работе представлены интерфейсные формы пространства виртуальной лаборатории (рис. 4–6).



Рис. 1. Внешний вид виброакустического излучателя с креплением на железобетонные перекрытия, бетонных перекрытий и иных несущих конструкций
Fig. 1. Appearance of a vibro-acoustic emitter mounted on reinforced concrete floors, concrete ceilings, and other load-bearing structures



Рис. 2. Внешний вид кодонаборной панели
Fig. 2. Appearance of the code dialing panel



Рис. 3. Внешний вид коммутатора
Fig. 3. Appearance of the switch



Рис. 4. Размещение средств защиты в виртуальной комнате «Защита информации от утечки по техническим каналам» с использованием гарнитуры виртуальной реальности Pico 4

Fig. 4. Placement of security tools in the virtual room “Protecting information from leakage through technical channels” using the Pico 4 virtual reality headset

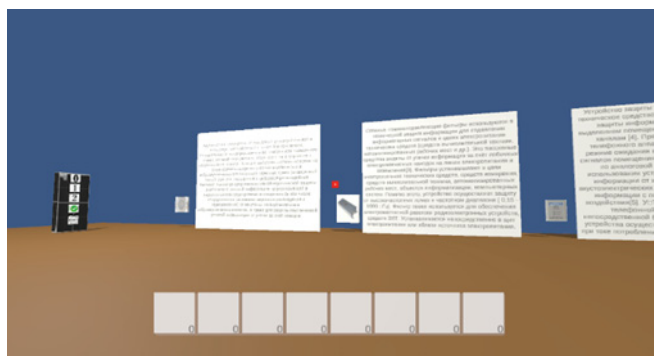


Рис. 5. Интерфейсная часть диалогового окна пространства теоретического блока обучения

Fig. 5. Interface part of the theoretical training block dialog box

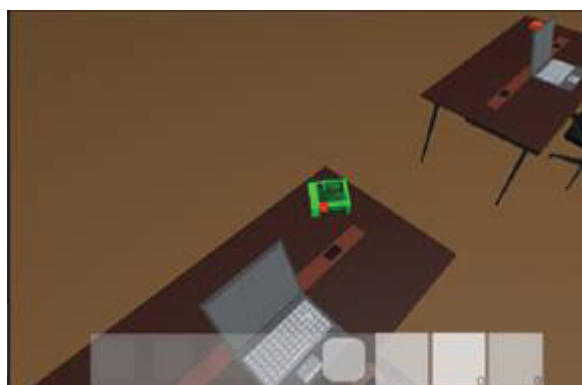


Рис. 6. Интерфейсная часть по идентификации верно и неверно установленных средств защиты информации

Fig. 6. The interface part for identifying correctly and incorrectly installed information security tools

На рис. 7 представлен фрагмент программного кода по разработке модуля системы идентификации объектов виртуальной лаборатории.

```

369 public void checkНастенныйВибро()
370 {
371     Debug.Log("Проверка настенный начало");
372     //CounterVibroSP157s = 0;
373     //Debug.Log("Я в проверкеVibroSP-157");
374     //if (КомплексныйПостановщикПомехs == null)
375     НастенныйВиброс = GameObject.FindGameObjectsWithTag("НастенныйВибро");
376     //if (Places == null)
377     НастенныйВиброPlaces = GameObject.FindGameObjectsWithTag("НастенныйВиброPlace");
378     //for (int i = 0; i <= КомплексныйПостановщикПомехs.Length; i++)
379     foreach (GameObject НастенныйВибро in НастенныйВиброс)
380     {
381         foreach (GameObject НастенныйВиброPlace in НастенныйВиброPlaces)
382         {
383             if (Vector3.Distance(НастенныйВибро.transform.position, НастенныйВиброPlace.transform.position) <= 5)
384             {
385                 CounterНастенныйВиброс++;
386                 Debug.Log("Я в проверке настенный вибро");
387                 Destroy(НастенныйВиброPlace);
388                 НастенныйВибро.GetComponent<Renderer>().material.color = Color.green;
389             }
390         }
391     }
392     foreach (GameObject НастенныйВиброPlace in НастенныйВиброPlaces)
393     {
394         НастенныйВиброPlace.GetComponent<MeshRenderer>().enabled = true;
395     }
396 }
397 }
398

```

Рис. 7. Фрагмент программного кода по разработке системы идентификации объектов

Fig. 7. Fragment of the software code for developing an object identification system

В работе предлагается модульная архитектура виртуальной лаборатории, благодаря которой добавление новых сценариев выполнения практических работ (например, защита речевой информации, монтаж систем контроля доступа) не требует переработки алгоритма взаимодействия пользователя с системой, инвентаря и автоматизированной проверки. Это позволяет рассматривать разработку как масштабируемую платформу для создания виртуального лабораторного практикума по технической защите информации с возможностью расширения функционала в дальнейшем без изменения базового программного ядра. Кроме того, в отличие от узкоспециализированных тренажеров, ориентированных на один тип средств или канал утечки, предложенная лаборатория объединяет технические, программно-аппаратные и физические средства защиты, что соответствует комплексному характеру профессиональной деятельности специалиста.

Экспериментальное исследование

В процессе исследования была проведена апробация виртуальной лаборатории по защите информации на объекте информатизации. Апробация проводилась в течение двух месяцев. Для апробации были привлечены группы студентов КНИТУ-КАИ им. А. Н. Туполева направления подготовки 10.03.01 «Информационная безопасность» – 42 человека. Студенты были разделены на две группы по 21 человеку. Первая использовала виртуальную лабораторию в сочетании с традиционными методическими указаниями, вторая – только виртуальный тренажер. В рамках апробации фиксировались следующие два параметра:

1. Время, затраченное на одну попытку выполнения практического задания по каждой виртуальной комнате (мин.).
2. Количество верно установленных средств защиты информации при каждой попытке выполнения практического задания по каждой виртуальной комнате (шт.).

Результат апробации показал, что среднее время размещения средств защиты информации в виртуальной комнате «Защита информации от утечки по техническим каналам» (ТКУИ) с использованием виртуальной лаборатории совместно с методическими указаниями составило 127 минут. При использовании только виртуальной лаборатории среднее время составило 105 минут. Среднее время размещения средств защиты информации сократилось на 17 % у студентов, выполнявших работу только с использованием виртуальной лаборатории «Защита информации от утечки по техническим каналам». Аналогичная тенденция наблюдалась при выполнении лабораторных работ в виртуальных комнатах «Программно-аппаратная защита сетевого периметра» и «Физическая защита информации на объекте информатизации». Среднее время сократилось на 34 и 35 % соответственно. В среднем время выполнения лабораторных работ по каждой виртуальной комнате с использованием виртуального тренажера сократилось на 30 минут. На рис. 8 представлена гистограмма распределения среднего времени размещения средств защиты информации по каждой виртуальной комнате.

Фактором снижения среднего времени выполнения лабораторной работы по каждой виртуальной комнате является снижение среднего времени размещения каждого средства защиты информации. В результате статистического анализа можно сделать вывод о том, что среднее время размещения одного средства защиты информации по виртуальным комнатам «Защита информации от утечки по техническим каналам», «Программно-аппаратная защита сетевого периметра», «Физическая защита информации на объекте информатизации» снизилось на 17; 40 и 30 % соответственно, что в среднем по трем комнатам составило примерно 32 %. На рис. 9 представлена гистограмма распределения среднего времени размещения одного средства защиты информации.

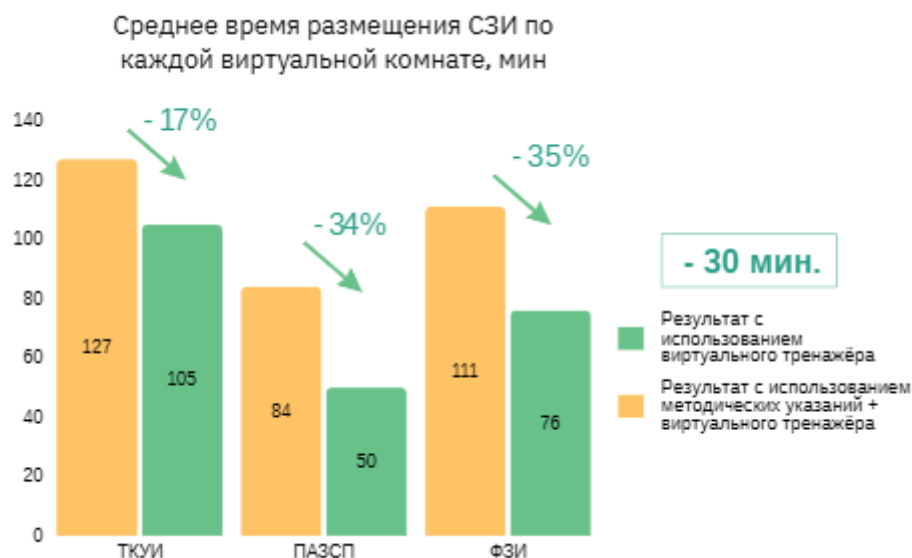


Рис. 8. Гистограмма распределения среднего времени размещения средств защиты информации по каждой виртуальной комнате

Fig. 8. Histogram of the distribution of the average placement time of information protection tools for each virtual room

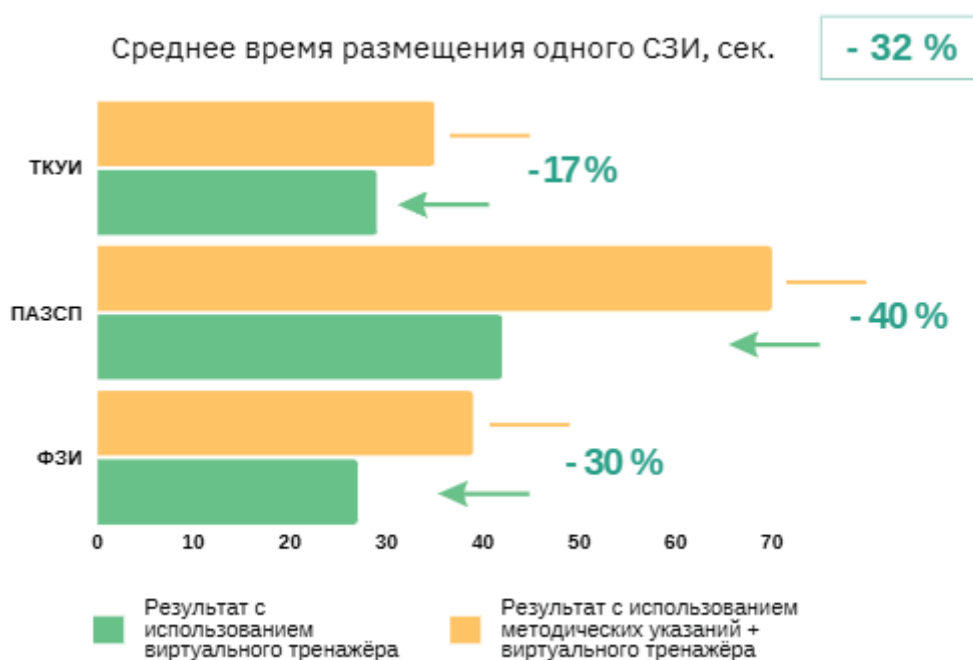


Рис. 9. Гистограмма распределения среднего времени размещения одного средства защиты информации

Fig. 9. Histogram of the distribution of the average placement time of one information protection tool

В рамках апробации были получены статистические данные среднего количества корректно установленных средств защиты информации при каждой попытке отдельно в каждой виртуальной комнате. По полученным данным можно сделать вывод о том, что виртуальный тре-

нажер является эффективным инструментом для обучения студентов. Гистограмма показывает улучшение результатов обучающихся с каждой попыткой выполнения практического задания. На рис. 10 представлена гистограмма распределения среднего количества корректно установленных средств защиты информации при каждой попытке.



Рис. 10. Гистограмма распределения среднего количества верно установленных средств защиты информации при каждой попытке
 Fig. 10. Histogram of the distribution of the average number of correctly installed information protection tools for each attempt

Результат эксперимента показал, что выполнение практических работ с использованием только виртуальной лаборатории обеспечивает высокий процент количества верно размещенных средств защиты информации на объекте информатизации за все использованные попытки. Представленные результаты демонстрируют динамику среднего количества верно установленных средств защиты информации (СЗИ) по каждой попытке выполнения лабораторных работ в трех виртуальных комнатах: «Защита информации от утечки по техническим каналам» (ТКУИ), «Программно-аппаратная защита сетевого периметра» (ПАЭСП) и «Физическая защита информации» (ФЗИ).

В группе, использовавшей только виртуальный тренажер (без методических указаний), наблюдается более выраженная положительная динамика. В комнате ТКУИ количество верно установленных СЗИ увеличилось на 450 % (с 12 до 66), в ПАЭСП – на 206 % (с 18 до 55), в ФЗИ – на 175 % (с 20 до 55). В группе, использовавшей традиционные методические указания совместно с тренажером, прирост составил: в ТКУИ – 32 % (с 40 до 53), в ПАЭСП – 124 % (с 21 до 47), в ФЗИ – 122 % (с 23 до 51).

Таким образом, использование виртуальной лаборатории способствует более интенсивному усвоению материала и позволяет обучающимся быстрее достигать высоких результатов за счет активного взаимодействия с иммерсивной средой. Полученные данные подтверждают эффективность разработанной виртуальной лаборатории как инструмента формирования профессиональных компетенций в области защиты информации.

На рис. 11 представлена диаграмма распределения среднего количества верно установленных средств защиты информации по каждой виртуальной комнате.



Рис. 11. Диаграмма распределения среднего количества верно установленных средств защиты информации по каждой виртуальной комнате

Fig. 11. Distribution chart of the average number of correctly installed information security tools for each virtual room

Экспериментальное исследование подтвердило эффективность разработки: среднее количество верно установленных средств защиты информации при использовании только виртуальной лаборатории составило на 6 % больше, чем при использовании комбинированного подхода (виртуальная лаборатория совместно с методическими указаниями). Таким образом, экспериментальное исследование подтверждает эффективность виртуальной лаборатории в учебном процессе. Иммерсивная среда обеспечивает безопасное, наглядное и интерактивное моделирование реальных условий защиты информации, что особенно актуально в условиях ограниченной материально-технической базы учебных заведений и высокой стоимости реальных технических средств защиты. Полученные данные имеют практическую ценность разработанного программного продукта как инструмента подготовки квалифицированных специалистов в области информационной безопасности.

Заключение

Результаты проведенного исследования свидетельствуют о том, что разработанная виртуальная лаборатория по защите информации на объекте информатизации, включающая технические, программно-аппаратные и физические средства защиты, обеспечивает более глубокое усвоение учебного материала по сравнению с традиционными методами обучения. Апробация данной лаборатории в образовательном процессе показала, что обучающиеся успешно трансформируют теоретические знания в практические навыки и компетенции в условиях, максимально приближенных к реальным профессиональным задачам в области информационной безопасности. В работе основное внимание уделено корректности выбора типа и места установки средства защиты, что является базовым навыком для начинающих специалистов. В дальнейшем планируется дополнить лабораторию модулем конфигурирования, в котором обучающемуся потребуется не только установить, но и настроить устройство (например, задать политику межсетевого экрана или определить зону обнаружения датчика).

Список литературы

1. Дудырев Ф. Ф., Максименкова О. В. Симуляторы и тренажеры в профессиональном образовании: педагогические и технологические аспекты // Вопросы образования. 2020. № 3. С. 255–276. DOI 10.17323/1814-9545-2020-3-255-276
2. Вольнов М. М., Китов А. А., Горячкин Б. С. Виртуальная реальность: виды, структура, особенности, перспективы развития // Информационные технологии. 2023.

3. Виртуальный тренажер «Технология VPN» [Электронный ресурс] // Портал учебного оборудования. URL: <https://labstand.ru/catalog/virtualnye-trenazhery-i-emulyatory-zashhita-informaczii/virtualnyj-trenazher-tehnologiya-vpn> (дата обращения: 13.11.2025).
4. Виртуальный учебник «Кибербезопасность: Управление и реагирование на инциденты безопасности» // Портал учебного оборудования. URL: <https://labstand.ru/catalog/virtualnye-trenazhery-i-emulyatory-zashhita-informaczii/virtualnyj-uchebnik-kiberbezopasnost-upravlenie-i-reagirovanie-na-inczidenty-bezopasnosti> (дата обращения: 13.11.2025).
5. Демонстрационная виртуальная среда «Монтаж средств ТЗИ» // DOKISAN. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/demonstracionnaya-virtualnaya-sreda-montazh-sredstv-tzi> (дата обращения: 13.11.2025).
6. Виртуальный тренажер «Системы видеонаблюдения» // DOKISAN. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/virtualnyj-trenazhior-sistemy-videonabljudeniya> (дата обращения: 13.11.2025).
7. Виртуальный тренажер «Монтаж и сервисное обслуживание систем охраны» (МССО) // DOKISAN. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/mssso> (дата обращения: 13.11.2025).
8. Виртуальный тренажер «Защита объекта от утечек информации по техническим каналам» // DOKISAN. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/defend-exam> (дата обращения: 13.11.2025).
9. Виртуальный тренажер «Обнаружение скрытых устройств и скрытых видеокамер» // DOKISAN. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/ozu> (дата обращения: 13.11.2025).
10. Виртуальный тренажер «Основы квантовой криптографии» // DOKISAN. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/virtualnyj-trenazhior-osnovy-kvantovoj-kriptografii> (дата обращения: 13.11.2025).
11. Шарипов Р. Р., Макаров С. П., Кассирова А. А. Разработка программного комплекса потокового шифра RC4 для обучающихся по дисциплине «криптография». Международный научно-исследовательский журнал. 2024. № 9(147). С. 1–8. DOI 10.60797/IRJ.2024.147.15
12. Хаерова Э. И., Гарифуллин Р. Ф., Тумбинская М. В. VR-тренажер по обработке конфиденциальных данных на физическом носителе // Информатизация образования и науки. 2024. № 2 (62). С. 62–76.
13. Гибадуллин Р. Ф., Лекомцев Д. В., Перухин М. Ю. Анализ параметров промышленных сетей с применением нейросетевой обработки // Искусственный интеллект и принятие решений. 2020. № 1. С. 80–87.
14. Гизатуллин З. М., Гизатуллин Р. М., Мубараков Р. Р. Моделирование помех в электронном устройстве при воздействии импульсного магнитного поля с использованием искусственной нейронной сети // Журнал радиоэлектроники. 2024. № 5. С. 1–15.
15. Гизатуллин З. М., Мубараков Р. Р. Прогнозирование помех электростатического разряда в электронном устройстве с использованием искусственной нейронной сети // Журнал радиоэлектроники. 2024. № 8. С. 1–18.

References

1. Dudyrev F. F., Maksimenkova O. V. Simulators and training devices in vocational education: pedagogical and technological aspects. *Voprosy obrazovaniya*, 2020, no. 3, pp. 255–276. DOI 10.17323/1814-9545-2020-3-255-276 (in Russ.)
2. Volnov M. M., Kitov A. A., Goryachkin B. S. Virtual reality: types, structure, features, development prospects. *Information technologies*, 2023, pp. 1–18. (in Russ.)

3. Virtual simulator “VPN Technology”. *Portal of educational equipment*. URL: <https://labstand.ru/catalog/virtualnye-trenazhery-i-emulyatory-zashhita-informaczii/virtualnyj-trenazher-tehnologiya-vpn> (date of access: 11/13/2025). (in Russ.)
4. Virtual textbook “Cybersecurity: Managing and responding to security incidents”. *Portal of educational equipment*. URL: <https://labstand.ru/catalog/virtualnye-trenazhery-i-emulyatory-zashhita-informaczii/virtualnyj-uchebnik-kiberbezopasnost-upravlenie-i-reagirovanie-na-inczidenty-bezopasnosti> (date of access: 11/13/2025). (in Russ.)
5. Demonstration virtual environment “Installation of TZI equipment”. *DOKISAN*. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/demonstracionnaya-virtualnaya-sreda-montazh-sredstv-tzi> (date of access: 11/13/2025). (in Russ.)
6. Virtual simulator “Video surveillance systems”. *DOKISAN*. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/virtualnyj-trenazh-jor-sistemy-videonabljudeniya> (date of access: 11/13/2025). (in Russ.)
7. Virtual simulator “Installation and maintenance of security systems” (MSSO). *DOKISAN*. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/mssso> (date of access: 11/13/2025). (in Russ.)
8. Virtual simulator “Protecting an object from information leaks through technical channels”. *DOKISAN*. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/defend-exam> (date of access: 11/13/2025). (in Russ.)
9. Virtual simulator “Detection of bugs and hidden video cameras”. *DOKISAN*. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/ozu> (date of access: 11/13/2025). (in Russ.)
10. Virtual simulator “Fundamentals of quantum cryptography”. *DOKISAN*. URL: <https://dokisan.com/catalog/virtualnye-trenazhery/virtualnyj-trenazh-jor-osnovy-kvantovoj-kriptografii> (date of access: 13.11.2025). (in Russ.)
11. **Sharipov R. R., Makarov S. P., Kassirova A. A.** Development of a software package for the RC4 stream cipher for students majoring in cryptography. *International Research Journal*, 2024, no. 9 (147), pp. 1–8. DOI 10.60797/IRJ.2024.147.15. (in Russ.)
12. **Khaerova E. I., Garifullin R. F., Tumbinskaya M. V.** VR simulator for processing confidential data on a physical medium. *Informatization of education and science*, 2024, no. 2 (62), pp. 62–76. (in Russ.)
13. **Gibadullin R. F., Lekomtsev D. V., Perukhin M. Yu.** Analysis of industrial network parameters using neural network processing. *Artificial Intelligence and Decision Making*, 2020, no. 1, pp. 80–87. (in Russ.)
14. **Gizatullin Z. M., Gizatullin R. M., Mubarakov R. R.** Modeling of interference in an electronic device exposed to a pulsed magnetic field using an artificial neural network. *Journal of Radio Electronics*, 2024, no. 5, pp. 1–15. (in Russ.)
15. **Gizatullin Z. M., Mubarakov R. R.** Prediction of electrostatic discharge interference in an electronic device using an artificial neural network. *Journal of Radio Electronics*, 2024, no. 8, pp. 1–18. (in Russ.)

Информация об авторах

Тумбинская Марина Владимировна, кандидат технических наук, доцент

Гатауллин Булат Ильнурович, студент

Хаерова Эндже Ильдаровна, магистрант

Information about the Authors

Marina V. Tumbinskaya, Candidate of Technical Sciences

Bulat I. Gataullin, Student

Endje I. Haerova, Master's Student

Статья поступила в редакцию 16.01.2026;

одобрена после рецензирования 23.03.2026; принята к публикации 30.03.2026

The article was submitted 16.01.2026;

approved after reviewing 23.03.2026; accepted for publication 30.03.2026

Научная статья

УДК 004.83

DOI 10.25205/1818-7900-2026-24-1-55-66

Методы генерации критических вопросов для анализа аргументации на русском языке

Диаб Хайдар

Новосибирский государственный университет
Новосибирск, Россия
k.diab1@g.nsu.ru

Аннотация

Исследование посвящено разработке и реализации модуля генерации критических вопросов к текстам, содержащим аргументацию, с использованием схем аргументации Д. Уолтона и больших языковых моделей. В рамках исследования для проведения экспериментов было разработано два корпуса текстов: англоязычный корпус, состоящий из 200 аргументов, и русскоязычный корпус из 92 текстов, из которых 54 текста являются новостными статьями, а 38 текстов – эссе или аналитическими статьями. Для тестирования подхода к генерации критических вопросов было использовано несколько больших языковых моделей: mistralai/Mistral-7B-Instruct-v0.3, Qwen/Qwen2.5-7B-Instruct и google/flan-t5-large. Результаты исследования показали, что модель google/flan-t5-large имеет лучшие результаты на англоязычном корпусе (совокупный балл – 0,361), а модель mistralai/Mistral-7B-Instruct-v0.3 – на русскоязычном корпусе (совокупный балл – 0,285; BERTScore – 0,722). Реализованный модуль генерации критических вопросов может быть использован для построения компоненты интерпретации результатов анализа достоверности информации.

Ключевые слова

генерация критических вопросов, анализ аргументации, большие языковые модели, схемы аргументации Д. Уолтона, проверка достоверности информации, обработка естественного языка, интерпретируемые модели ИИ

Для цитирования

Хайдар Д. Методы генерации критических вопросов для анализа аргументации на русском языке // Вестник НГУ. Серия: Информационные технологии. 2026. Т. 24, № 1. С. 55–66. DOI 10.25205/1818-7900-2026-24-1-55-66

Methods for Generating Critical Questions for Analyzing Russian-Language Arguments

Diab Haidar

Novosibirsk State University
Novosibirsk, Russian Federation
k.diab1@g.nsu.ru

Abstract

In this paper, we aim to develop an efficient workflow for critical questions generation from argumentative texts based on argumentation schemes and modern large language models. We have prepared two corpora for this task: an English corpus with 200 argumentative texts and a Russian corpus with 92 texts (54 news texts, 38 essays/analitics). We have also carried out an experiment with three different models for critical questions generation: mistralai/Mistral-7B-Instruct-v0.3, Qwen/Qwen2.5-7B-Instruct, google/flan-t5-large. The results have shown that the google/flan-t5-large

© Хайдар Д., 2026

model is the best for critical questions generation for the English corpus (composite score 0.361), while for the Russian corpus, it was the best performance of the mistralai/Mistral-7B-Instruct-v0.3 model (composite score 0.285, BERTScore 0.722). The critical questions generation module can be used as a component in information reliability analysis systems.

Keywords

critical question generation, argumentation, large language models, Walton's argumentation schemes, argumentation analysis, information reliability verification, natural language processing, interpretable AI models

For citation

Haidar D. Methods for generating critical questions for analyzing Russian-language arguments. *Vestnik NSU. Series: Information Technologies*, 2026, vol. 24, no. 1, pp. 55–66 (in Russ.) DOI 10.25205/1818-7900-2026-24-1-55-66

Введение

Оценка достоверности информации является одной из основных проблем в современных условиях повсеместного распространения информации через социальные сети и новостные агрегаторы. Эта проблема часто решается с помощью методов классификации текстов. Однако эти методы не раскрывают мотивы, лежащие в основе принятого решения, а для практических приложений они очень важны. В данной статье предлагается метод генерации критических вопросов (Critical Question Generation, CQG), которые используются для объяснения процесса проверки аргументов, в частности, для указания найденных достоверных источников информации, базовых предположений и возможных интерпретаций.

В контексте практического применения метод CQG может быть использован совместно с модулем анализа текста и модулем проверки фактов. Критические вопросы, сгенерированные этим методом, являются отправной точкой для поиска доказательств достоверности фактов, а полученные при этом результаты служат основной для составления отчета о выполненных проверках и найденных источниках. Этот подход особенно актуален для русскоязычного сегмента сети Интернет, для которого характерны высокий уровень вариативности в стилистике и многообразие информационных источников.

1. Обзор подходов

В отличие от репродуктивных методов построения вопросов, ориентированных на пересказ содержания, метод генерации критических вопросов нацелен на обнаружение возможных слабых мест аргументации (например, несостоятельность доказательств, недостоверность источников, логические неточности в причинно-следственных связях, неадекватность примеров, игнорирование исключений и т. п.). Теоретическую основу этого подхода составляют схемы аргументации Д. Уолтона и связанные с ними критические вопросы [1; 2]. С инженерной точки зрения генерация критических вопросов хорошо встраивается в технологию генерации ответов с дополненной выборкой (Retrieval-Augmented Generation, RAG): сначала формулируются проверочные подзадачи, затем для каждой подзадачи извлекаются документы, а ответ формируется на основе найденных в документах свидетельств [3]. В задачах фактчекинга (fact-checking) такой сценарий естественно сочетается с крупными датасетами, используемыми для проверки утверждений и извлечения доказательств (например, FEVER) [4], а также с нейросетевыми подходами к извлечению и ранжированию доказательств [5].

Критические вопросы связаны с аргументативными схемами (в частности, схемами Д. Уолтона), которые описывают типовые паттерны рассуждений и наборы вопросов для проверки надежности реализуемого в них вывода. Современные исследования в области генерации критических вопросов рассматривают эту задачу как автоматическое создание последовательности критических вопросов, направленных на анализ входного аргумента с использованием таких методов работы с языковыми моделями, как преобразование одних текстовых последователь-

ностей в другие (seq2seq) или выполнение инструкций пользователей (instruction-following). Для получения воспроизводимых результатов в данной работе используются открытые языковые модели и единый протокол автоматической оценки полученных результатов [6].

2. Подготовка данных и корпуса

В экспериментах использовались два корпуса: англоязычный корпус (Е-корпус) и русскоязычный корпус (R-корпус). Е-корпус содержит 200 примеров, данные в нем представлены в формате [Аргумент (тезис с кратким обоснованием), Набор критических вопросов]. R-корпус сгенерирован на основе данных из RSS-ленты и содержит 92 примера (новости и аналитические материалы). Средняя длина материала в R-корпусе составляет 444,5 символа, медиана – 351,5, диапазон длин – от 200 до 2 171. Данные для каждого документа включают следующие элементы: идентификатор, источник, дату публикации, заголовок и текстовые данные. Этапы предварительной обработки данных включают удаление HTML-элементов, нормализацию пробелов, исключение коротких материалов и удаление дубликатов. Удаление дубликатов основано на хеш-функции SHA-1 для текстовых данных. Фильтр минимальной длины для снижения доли коротких заметок без выраженной аргументации составлял 200 символов. Материалы для ручной аннотации включают руководство по схемам аргументации и шаблон с таблицей для выделения фрагментов, содержащих аргументацию, а также схемы Уолтона и критические вопросы [7].

Промпт для моделей построен в стиле instruction-following и фиксирует роль системы (ассистент по критическому мышлению), требуемое число вопросов и строгий формат ответа. Важной функцией протокола является контроль формата результатов работы модели. Модель должна вернуть JSON-массив строк, где каждая строка заканчивается знаком вопроса. Англоязычный корпус использовался как контрольный материал для проверки корректности реализации протокола.

После того как ИИ-модель сгенерировала текст критических вопросов, результат подвергается нормализации. На этом этапе удаляются служебные токены и поясняющий текст модели, а длинная строка разбивается на отдельные вопросы. Пустые строки также удаляются. В результате формируется список чистых критических вопросов, пригодных для использования в автоматизированных конвейерах, где каждый вопрос служит отдельным поисковым запросом для извлечения доказательств. Например, модель может сгенерировать строку:

«Вопрос 1: Каковы основания утверждения?»\nВопрос 2: Какие доказательства поддерживают вывод?»

После нормализации она преобразуется в список:

[«Вопрос 1: Каковы основания утверждения?», «Вопрос 2: Какие доказательства поддерживают вывод?»].

Модели запускаются с использованием одинаковых условий генерации для возможности сравнения качества формулирования вопросов для каждой модели, независимо от различий между ними. Например, были установлены ограничения на длину входного контекста для моделей и длину сгенерированного текста. При обучении моделей использовался тот же запрос и то же количество примеров вопросов. Были сгенерированы сводные результаты для всего корпуса, где для каждого показателя было рассчитано среднее значение по всем моделям, а затем были составлены сравнительные таблицы и графики.

3. Методология и архитектура конвейера

Конвейер CQG реализован в виде следующих этапов: (1) подготовка аргументативного фрагмента, (2) формирование подсказки на основе инструкций модели, аргументативного

контекста и требований к формату вывода, (3) формирование нескольких вопросов для одного входного примера, (4) постобработка и нормализация результатов, (5) автоматическая оценка на основе заранее определенного набора метрик. Для устойчивой обработки в экспериментах используется строгий формат вывода: JSON-массив строк. Это позволяет автоматически извлекать вопросы, измерять долю корректных ответов и строить сводные отчеты [1; 8; 9].

4. Постановка эксперимента

В ходе эксперимента были задействованы три модели, иллюстрирующие различные архитектурные стратегии: instruct-LLM (реализующая архитектуру исключительно с декодером) и модели класса seq2seq. Для всех моделей были установлены единые значения параметров генерации: температура – 0,7 (регулирует степень случайности при формировании выходной последовательности), top_p – 0,9 (ядро выборки, обеспечивающее связность порождаемых токенов). На каждый входной образец формировалось по три вопроса, при этом применялись ограничения на длину как входных, так и выходных токенов. Оценка качества выполнялась с привлечением обширного набора метрик, включавшего лексические и семантические показатели (BLEU, ROUGE-L, METEOR, BERTScore, SemSim, SemanticAcc), метрику формальной корректности вывода (Parseable) и метрику лексического разнообразия (Distinct-2).

Далее приводится детальная характеристика каждой из использованных метрик.

BLEU и METEOR – лексические меры совпадения с эталонными вопросами. Принимают значения в интервале от 0 до 1: чем ближе результат к единице, тем выше степень дословного соответствия.

ROUGE-L – метрика, опирающаяся на поиск наибольшей общей подпоследовательности в паре «сгенерированный вопрос – эталонный вопрос». Диапазон значений – от 0 до 1; более высокие значения свидетельствуют о более выраженном структурном подобии формулировок.

BERTScore – семантическая метрика, использующая контекстные векторные представления; позволяет сопоставлять близкие по смыслу высказывания даже при значительном лексическом расхождении. Принимает значения в пределах от 0 до 1; рост показателя указывает на повышение качества семантического сопоставления.

SemSim (semantic similarity) – косинусное сходство между эмбедингами сгенерированного и эталонного вопросов, вычисленное на уровне целых предложений. Значения лежат в интервале [0; 1].

SemanticAcc – мера семантической точности, отражающая долю случаев, в которых сгенерированный вопрос оказался семантически близким хотя бы к одному из эталонных критических вопросов.

Distinct-2 – показатель лексического разнообразия: доля уникальных биграмм относительно общего числа биграмм, содержащихся во всех сгенерированных вопросах. Приближение к единице соответствует более вариативному использованию лексики.

Parseable – доля примеров, для которых ответ модели был успешно интерпретирован как корректный JSON-массив вопросов. Выражается в процентах; чем выше процентное значение, тем лучше соблюдение формата вывода.

Score – составной показатель качества (принимая значения от 0 до 1), интегрирующий ключевые лексические и семантические метрики в единую величину, предназначенную для сравнительного ранжирования моделей. Вычисление производилось по формуле:

$$\text{Score} = 0,15 \cdot \text{BLEU} + 0,15 \cdot \text{ROUGE-L} + 0,10 \cdot \text{METEOR} + 0,20 \cdot \text{BERTScoreF1} + 0,25 \cdot \text{SemSim} + 0,15 \cdot \text{SemanticAcc}.$$

Метрики Distinct-2 и Parseable в составной индекс не включались и рассматривались изолированно, поскольку первая характеризует вариативность лексического оформления, а вто-

рая – формальную корректность структурирования ответа, но ни одна из них напрямую не отражает содержательное качество порождаемых критических вопросов.

5. Результаты и анализ

В табл. 1–2 приведены агрегированные метрики для англоязычного и русскоязычного корпусов. Для наглядности также построены графики сравнения моделей и распределения длин документов (рис. 1 и рис. 2).

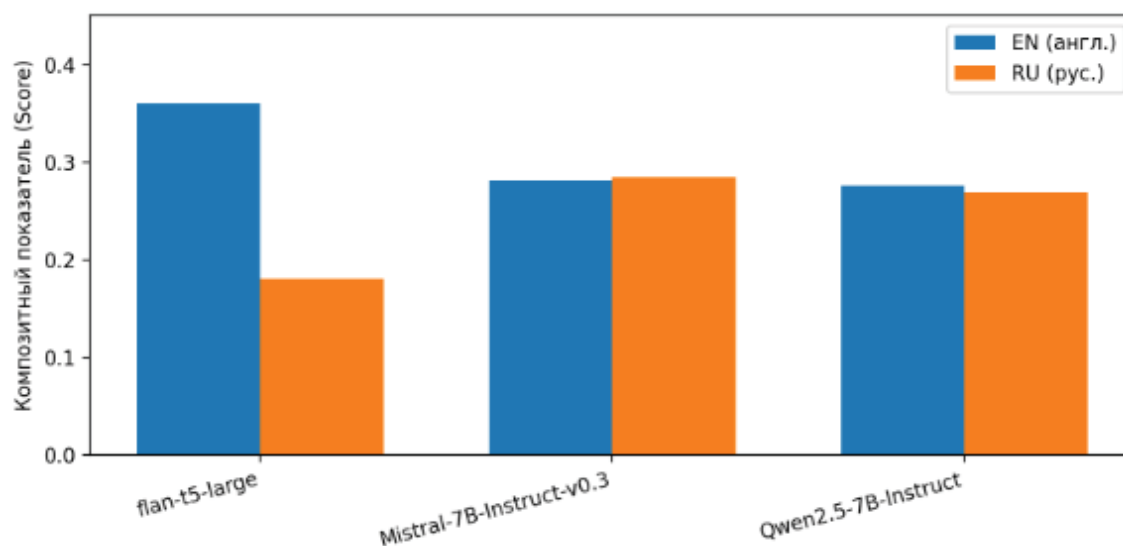


Рис. 1. Композитный показатель качества

Fig. 1. Composite quality diagram

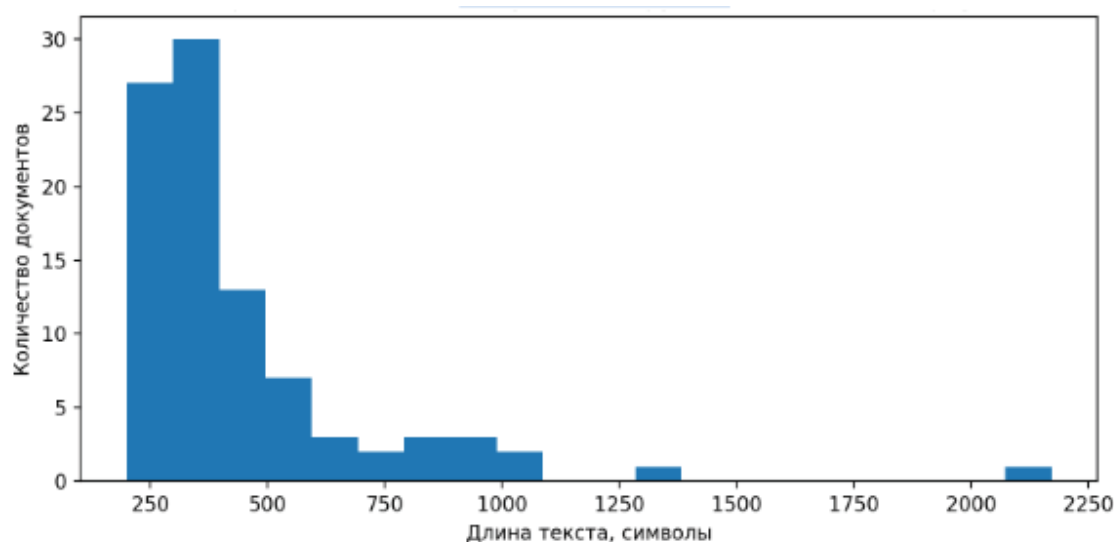


Рис. 2. Распределение длины документов в R-корпусе

Fig. 2. Distribution of document lengths in the R-corpus

Таблица 1

Результаты для англоязычного корпуса (агрегированные метрики)

Table 1

Results for the English-language corpus (aggregated metrics)

Модель	BLEU (лекс., 0-1, ↑)	METEOR (лекс., 0-1, ↑)	BERTScore (се- мант., 0-1, ↑)	SemSim (коси- нус, 0-1, ↑)	Distinct-2 (разн., 0-1, ↑)	Parseable (JSON, %, ↑)	Score (компози- 0-1, ↑)
flan-t5-large	0,096	0,202	0,348	0,646	0,987	100	0,361
Mistral-7B- Instruct-v0.3	0,042	0,161	0,249	0,619	0,906	100	0,281
Qwen2.5-7B- Instruct	0,043	0,153	0,252	0,608	0,950	100	0,276

Таблица 2

Результаты по русскоязычному корпусу (агрегированные показатели)

Table 2

Results on the Russian-language corpus (aggregated metrics)

Модель	BLEU (лекс., 0-1, ↑)	METEOR (лекс., 0-1, ↑)	BERTScore (семант., 0-1, ↑)	SemSim (косинус, 0-1, ↑)	Distinct-2 (разн., 0-1, ↑)	Parseable (JSON, %, ↑)	Score (компози- 0-1, ↑)
Mistral-7B- Instruct-v0.3	0,066	0,198	0,722	0,439	0,880	100	0,285
Qwen2.5-7B- Instruct	0,060	0,140	0,707	0,408	0,959	100	0,270
flan-t5-large	0,033	0,013	0,531	0,273	0,916	97	0,181

Примечание к табл. 1–2. Метрики BLEU и METEOR оценивают лексическое соответствие с эталонными вопросами; метрики BERTScore и SemSim оценивают семантическую близость формулировок; Distinct-2 измеряет разнообразие вопросов; Parseable указывает долю случаев, в которых ответ может быть правильно извлечен в виде массива JSON; Score представляет собой агрегированный показатель качества.

6. Качественный анализ: схемы Уолтона и примеры генерации

6.1. Схемы аргументации и типовые критические вопросы

Схема	Краткое описание	Типовые критические вопросы (примеры)
expert_opinion	Аргумент от мнения эксперта/авторитета	Насколько источник компетентен и надежен? Есть ли независимые подтверждения/контрмнения?
cause_effect	Причинно-следственная связь	Действительно ли А вызывает В? Есть ли альтернативные причины и контрпримеры?
consequences	Аргумент от последствий	Насколько вероятны заявленные последствия? Насколько последствия значимы и есть ли альтернативы?
practical_reasoning	Практическое рассуждение (цель – действие)	Приводит ли действие к цели и каковы риски? Есть ли более безопасные/эффективные альтернативы?
example	Аргумент от примера	Достаточно ли примеров и репрезентативны ли они? Существуют ли контрпримеры/исключения?

6.2. Примеры генерации критических вопросов по схемам аргументации

6.2.1. R-корпус. Примеры

Пример 1 (схема expert_opinion; корпус RSS, модель Mistral-7B-Instruct-v0.3).

Аргумент: Марвин Мински: гений, когнитивный ученый, крестный отец ИИ. Когда мы говорим об искусственном интеллекте, мы неизбежно говорим о наследии Марвина Мински. Этот ученый не просто стоял у истоков области искусственного интеллекта; он в значительной степени определил, чем станет эта область.

Ответ модели:

[
«Какие доказательства подтверждают утверждение, что Марвин Мински определил искусственный интеллект и его область?»,

«Существуют ли альтернативные объяснения того, как были определены искусственный интеллект и его область, помимо вклада Марвина Мински?»,

«Какие предположения лежат в основе аргумента о вкладе Марвина Мински в определение искусственного интеллекта и его области?»

]

Комментарий: Вопросы предназначены для проверки основания утверждения (доказательства), наличия других объяснений (альтернатив) и лежащих в основе утверждения предположений. Это набор обычных вопросов для аргументации с помощью авторитета.

Пример 2 (схема expert_opinion; корпус RSS, модель Mistral-7B-Instruct-v0.3).

Аргумент: Соединенные Штаты взяли на себя ответственность за авиационный инцидент в Вашингтоне 29 января. Более конкретно, власти США признали свою ответственность за столкновение пассажирского самолета American Airlines и вертолета Black Hawk в воздухе над Вашингтонским национальным аэропортом имени Рональда Рейгана 29 января. Об этом признании сообщила газета The Washington Post со ссылкой на соответствующее судебное заявление Министерства юстиции.

Ответ модели:

[

«Какой документ лежит в основе судебного заявления?»»,

«Существуют ли альтернативные объяснения столкновения самолета и вертолета, помимо ответственности США?»»,

«Какие доказательства подтверждают, что Соединенные Штаты признали свою ответственность за столкновение?»»

]

Комментарий: Поставленные выше вопросы позволяют проверить аргумент путем идентификации документа, рассмотрения альтернативных точек зрения и предоставления подтверждающих доказательств, которые имеют решающее значение для проверки аргумента.

Пример 3 (схема «причина – следствие»; корпус RSS, модель Mistral-7B-Instruct-v0.3).

Аргумент: Группа компаний Litres изъяла из продажи около 4500 книг из-за опасений, связанных с законодательством.

Ответ модели:

[

«Какие конкретно законодательные акты вызывают опасения у издательства?»»,

«Существуют ли альтернативные способы распространения или публикации книг, которые соответствовали бы действующему законодательству?»»,

«На каком основании было решено, что изъятые книги не соответствуют действующему законодательству?»»

]

Комментарий: Вопросы исследуют взаимосвязь между причиной и следствием, альтернативные объяснения и ставят под сомнение основания для вывода, снижая вероятность правильного понимания «из-за» как подразумевающей причинно-следственную связь.

Пример 4 (схема expert_opinion; RSS-корпус, модель Mistral-7B-Instruct-v0.3).

Аргумент: Мосгорсуд признал законным приговор по делу о выводе из России 1 трлн рублей. Мосгорсуд оставил в силе приговор 11 бывшим сотрудникам банка «Новое время», признанным виновными в участии в организованном преступном сообществе и незаконной банковской деятельности. Они участвовали в схеме по выводу из России средств на сумму около 1 трлн руб.

Ответ модели:

[

«Какие доказательства подтверждают, что схема по выводу из России средств суммой около 1 трлн рублей незаконна?»»,

«Есть ли альтернативные объяснения для участия в организованном преступном сообществе и незаконной банковской деятельности, кроме участия в схеме по выводу из России средств?»»,

«На каких предположениях основан аргумент о незаконности схемы по выводу из России средств?»»

]

Комментарий: вопросы уточняют источник/основания решения и качество доказательств, а также требуют рассмотреть альтернативные объяснения – это усиливает проверяемость утверждения.

Пример 5 (схема consequences; корпус RSS, модель Mistral-7B-Instruct-v0.3).

Аргумент: Президент Венесуэлы взаимодействует с Генеральным секретарем Организации Объединенных Наций в отношении предполагаемых угроз со стороны Соединенных Штатов, согласно пресс-релизу Организации Объединенных Наций. Президент Венесуэлы г-н Мадуро информирует Генерального секретаря Организации Объединенных Наций о «напряженности в регионе», вызванной Соединенными Штатами, которая может иметь опасные последствия.

Ответ модели:

[
«На каком основании считается, что существуют угрозы со стороны Соединенных Штатов?»,
«Возможно ли, что другие региональные игроки также представляют угрозу для Венесуэлы?»,
«Какие факты или данные подтверждают вывод о том, что эти угрозы могут привести к опасным последствиям?»
]

Комментарий: Заданные вопросы основаны на доказательствах существования угроз, вероятности последствий и рассмотрении альтернативных источников угроз, что является основой для аргументации.

6.2.2. Е-корпус. Примеры

Пример 1 (схема expert_opinion; Е-корпус, модель Mistral-7B-Instruct-v0.3).

Аргумент: Marvin Minsky: a personality associated with the concepts of genius, cognitive science, and the founding influence in artificial intelligence. The mention of artificial intelligence is always linked to the legacy of Marvin Minsky. This scientist not only found himself at the core of the founding of this science, but also significantly guided its development.

Ответ модели:

[
«What substantiates the claim that Marvin Minsky founded artificial intelligence and defined the field thereof? “
“Is there any other possible explanation for the founding and definition of the field of artificial intelligence, apart from Marvin Minsky’s contributions?”
“What underlying assumptions are there for the claim that Marvin Minsky had a defining impact in the field of artificial intelligence?”
]

Комментарий: Эти вопросы предназначены для проверки обоснованности утверждения (доказательства), наличия альтернативных объяснений и лежащих в основе предположений, таким образом представляя собой набор проверок для аргумента, основанного на авторитете.

Пример 2 (схема expert_opinion; Е-корпус, модель Mistral-7B-Instruct-v0.3).

Аргумент: The United States admitted responsibility for the plane crash in Washington on January 29. The US government admitted responsibility for the collision of an American Airlines passenger jet and a Black Hawk helicopter in the skies over Ronald Reagan Washington International

Airport on January 29. This was reported by The Washington Post, citing a court filing by the Department of Justice.

Ответ модели:

[
«On what document is the court filing based?»,
«Are there any explanations for the plane-helicopter collision other than US responsibility?»,
«What evidence supports the US admitting responsibility for the collision?»
]

Комментарий: Вопросы позволяют определить источник утверждения, альтернативные версии события и данные, подтверждающие утверждение, тем самым рационализируя проверку аргумента на основе первичных и независимых источников.

Пример 3 (схема *cause_effect*; E-корпус, модель Mistral-7B-Instruct-v0.3).

Аргумент: Компания Litres изъяла из продажи 4500 книг из-за юридических проблем. Группа компаний Litres изъяла из продажи около 4500 книг.

Ответ модели:

[
«What specific laws are Litres Publishing House concerned about?»,
«Are there alternative ways to sell or publish such books that would comply with existing laws?»,
«On what basis were the conclusions reached regarding the non-compliance of the books that were removed with the law?»
]

Комментарий: Вопросы исследуют причинно-следственную связь (причину действия), альтернативные причины действия и основополагающие доводы для вывода, тем самым снижая вероятность неправильной интерпретации причинно-следственной связи, выраженной словами «из-за».

7. Обсуждение

Можно утверждать, что набор критических вопросов может рассматриваться как промежуточный уровень между текстом и окончательным выводом. В задачах анализа достоверности такая структура может считаться полезной, поскольку она помогает в проверке аргументации. Ее также можно считать подходящей для интеграции с поисковым модулем, поскольку каждый вопрос требует поиска фактов, подтверждающих его, и интеграции этих фактов в окончательную оценку аргументации [9; 10].

Отдельное практическое преимущество CQG – снижение зависимости от одного запроса. Даже если исходное утверждение сформулировано неоднозначно, его декомпозиция в виде 3–6 вопросов повышает полноту поиска и помогает собрать доказательства разных типов: прямые факты, контрпримеры, статистику, заявления компетентных источников. Это создает основу для последующих исследований в области RAG-фактчекинга (проверки фактов с использованием Retrieval-Augmented Generation) и объяснимого искусственного интеллекта (Explainable AI), что важно для повышения прозрачности и доверия к системам автоматической проверки достоверности информации. [11; 12].

Если рассматривать CQG как модуль в полной (end-to-end) системе проверки утверждений, то сгенерированные вопросы можно использовать как набор поисковых запросов. По каждому вопросу запускается поиск информации из внешних источников (например, из Интернета и медиаархивов). Далее, окончательный вывод составляется только на основе собранных доказательств, при этом на каждом этапе предоставляется текстовое обоснование по принципу «какой вопрос проверяется?» и «как он подтверждается источником?» [13].

Следует подчеркнуть, что использованный в настоящей работе русскоязычный корпус обладает сравнительно скромным объемом – 92 текста, из которых 54 приходятся на новостные сообщения, а 38 представляют собой эссе либо аналитические публикации. По этой причине полученные результаты надлежит трактовать исключительно как пилотные: они дают возможность зафиксировать наиболее заметные тенденции в качестве автоматического порождения критических вопросов и провести сопоставление моделей в рамках поставленной экспериментальной задачи, однако не позволяют распространять сделанные наблюдения на весь спектр русскоязычных текстов аргументативного характера. В перспективе представляется оправданным расширение корпуса путем привлечения более широкого круга источников, жанров и тематик, а также углубление процедуры ручной аннотации аргументативных структур и эталонных критических вопросов.

Заключение

В работе построен воспроизводимый конвейер генерации критических вопросов для англоязычных и русскоязычных текстов и выполнено сравнение нескольких открытых языковых моделей. Получены агрегированные метрики качества, иллюстрирующие устойчивую генерацию вопросов с контролируемым форматом вывода и разнообразием формулировок. Результаты могут быть использованы как основа для дальнейшей разработки модулей с объяснениями их результатов в системах проверки достоверности информации, что обеспечивает прозрачность и понятность выводов, а также возможность объяснения процесса принятия решения.

В качестве известного ограничения настоящего исследования следует указать относительно незначительный объем русскоязычного корпуса, насчитывающего 92 текста, вследствие чего полученные данные правомерно квалифицировать как предварительные. В планах дальнейшей работы предусматривается существенное расширение корпуса, привлечение более гетерогенных по своей природе образцов аргументативных текстов, а также совершенствование процедуры ручной аннотации аргументативных схем и эталонных критических вопросов.

Следующим этапом развития может стать включение модуля CQG в прототип системы проверки утверждений с извлечением информации из внешних источников и последующей агрегацией доказательств. В таком сценарии вопросы выполняют роль механизма декомпозиции: система проверяет отдельные аспекты утверждения, а затем формирует итоговое решение и краткое объяснение на основе найденных подтверждающих фрагментов. Это направление естественным образом продолжает тему автоматического обнаружения недостоверной информации и может служить основой для дальнейших исследований, связанных с развитием систем для проверки достоверности информации, способных объяснять результаты своей работы.

Список литературы / References

1. **Walton D., Reed C., Macagno F.** Argumentation Schemes. Cambridge: Cambridge University Press, 2008.
2. **Song Y., Heilman M., Beigman Klebanov B., Deane P.** Applying argumentation schemes for essay scoring // Proceedings of the First Workshop on Argumentation Mining. Baltimore, Maryland: Association for Computational Linguistics, 2014. P. 69–78.
3. **Gao Y., Xiong Y., Gao X., Jia K., Pan J., Bi Y., Dai Y., Sun J., Wang H.** Retrieval-Augmented Generation for Large Language Models: A Survey // arXiv preprint. 2023. arXiv:2312.10997
4. **Thorne J., Vlachos A., Christodoulopoulos C., Mittal A.** FEVER: a large-scale dataset for fact extraction and verification // Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies,

- Volume 1 (Long Papers). New Orleans, Louisiana: Association for Computational Linguistics, 2018. P. 809–819.
5. **Soleimani A., Monz C., Worring M.** BERT for evidence retrieval and claim verification // Advances in Information Retrieval. Cham: Springer International Publishing, 2020. P. 359–366.
 6. **Ruiz-Dolz R., Lawrence J.** Detecting argumentative fallacies in the wild: problems and limitations of large language models // Proceedings of the 10th Workshop on Argument Mining. Singapore: Association for Computational Linguistics, 2023. P. 1–10.
 7. **Ruiz-Dolz R., Taverner J., Lawrence J., Reed C.** NLAS-multi: a multilingual corpus of automatically generated natural language argumentation schemes // arXiv preprint. 2024. arXiv:2402.144.
 8. **Chung Y.-L., Kuzmenko E., Tekiroglu S.S., Guerini M.** CONAN – COunter NArratives through nichesourcing: a multilingual dataset of responses to fight online hate speech // Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics. Florence, Italy: Association for Computational Linguistics, 2019. P. 2819–2829.
 9. **Visser J., Lawrence J., Reed C., Wagemans J., Walton D.** Annotating argument schemes // Argumentation. 2021. Vol. 35, No. 1. P. 101
 10. **Lawrence J., Visser J., Reed C.** BBC Moral Maze: test your argument // In: COMMA. 2018.
 11. **Maertens R., Roozenbeek J., Basol M., van der Linden S.** Long-term effectiveness of inoculation against misinformation: three longitudinal experiments // Journal of Experimental Psychology: Applied. 2021. Vol. 27, No. 1. P. 1–16.
 12. **Xu F., Lin Q., Han J., Zhao T., Liu J., Cambria E.** Are large language models really good logical reasoners? A comprehensive evaluation and beyond // arXiv preprint. 2023. arXiv:2306.09841 cs.
 13. **Musi E., Carmi E., Reed C., Yates S., O'Halloran K.** Developing misinformation immunity: How to reason-check fallacious news in a human–computer interaction environment // Social Media + Society. 2023. Vol. 9, No. 1. Article 20563051221150407.

Сведения об авторе

Диаб Хайдар, магистрант

Information about the Author

Diab Haidar, Master's Student

Статья поступила в редакцию 08.03.2026;

одобрена после рецензирования 15.04.2026; принята к публикации 15.04.2026

The article was submitted 08.03.2026;

approved after reviewing 15.04.2026; accepted for publication 15.04.2026

Правила оформления текста рукописи

Авторы представляют статьи на русском или английском языке объемом от 0,5 авторского листа (20 тыс. знаков) до 1 авторского листа (40 тыс. знаков), включая иллюстрации (1 иллюстрация форматом 190 × 270 мм = 1/6 авторского листа, или 6,7 тыс. знаков). Публикации, превышающие указанный объем, допускаются к рассмотрению только после индивидуального согласования с редакцией журнала.

Текст рукописи должен быть представлен в редколлегию в виде файла MS Word (.doc, .docx). Гарнитура Times New Roman, размер шрифта 11, межстрочный интервал 1, размеры полей – стандартные значения текстового редактора. Форматирование – выравнивание по ширине страницы, переносы слов включены, каждый новый абзац начинается с красной строки. Не допускается ручное форматирование абзацев (пробелами, лишними переводами строк, разрывами страниц).

Структура статьи

- Индекс УДК (универсальной десятичной классификации). Выравнивание по левому краю
- Название статьи. Выравнивание по центру, полужирный шрифт
- ФИО авторов (полностью). Выравнивание по центру, полужирный шрифт
- Места работы всех авторов. Выравнивание по центру, курсив
- Адреса электронной почты, ORCID авторов
- Аннотация статьи
- Ключевые слова, не более 10
- Благодарности, сведения о финансовой поддержке
- Название статьи **на английском языке**. Выравнивание по центру, полужирный шрифт
- ФИО авторов **на английском языке** (полностью). Выравнивание по центру, полужирный шрифт
- Места работы авторов **на английском языке**. Выравнивание по центру, курсив
- Аннотация статьи **на английском языке (Abstract)**, 200–250 слов
- Ключевые слова **на английском языке (Keywords)**, не более 10
- Благодарности, сведения о финансовой поддержке **на английском языке**, если есть соответствующий раздел на русском языке (**Acknowledgements**)
- Основной текст
- Список литературы / **References**
- Сведения об авторах

Требования к оформлению основного текста и иллюстративных материалов

Основной текст должен быть представлен в структурированном виде, рекомендуется использовать подзаголовки – например: Введение, Методика..., Выводы, Результаты, Заключение.

Подзаголовки отделяются и набираются полужирным шрифтом. В целях выделения частей текста и отдельных слов и словосочетаний допускается использование курсива или полужирного шрифта. Подчеркивание, разрядка, изменение основного кегля и выделение цветом не используются.

Иллюстрации к рукописи статьи должны быть приложены в виде отдельных файлов. При этом в тексте должно содержаться включенное изображение с указанием имени файла. Все иллюстрации, содержащие схемы, графики, алгоритмы и т. п., должны быть представлены в векторном виде (.ai, .eps, .cdr). Скриншоты и другие растровые изображения должны быть представлены в максимально высоком качестве, без каких-либо потерь и искажений (.jpg, .tif). Все иллюстрации должны иметь подрисуночную подпись – свое название. Надписи к таблицам и подписи к иллюстрациям приводятся **на двух языках (русском и английском)**.

Примеры:

Рис. 1. Диаграмма производительности...

Fig. 1. Performance diagram...

Таблица 1

Сравнение алгоритмов...

Table 1

Comparison of algorithms...

Нумерация последовательная и неразрывная от начала статьи. Не допускается использование других наименований, кроме «Рис.» / «Fig.», «Таблица» / «Table», и усложнение нумерации (например, «Рис. 3.2.»). Ссылка на иллюстрацию в тексте должна быть приведена в круглых скобках, например: (рис. 1), (табл. 1).

Формулы должны быть набраны с использованием редактора MathType либо встроенного редактора формул MS Word. Кегль основных символов – 11, греческие символы набираются прямым шрифтом, латинские – курсивом. Нумеруются только те формулы, на которые автор ссылается в тексте.

Abstract

Аннотация статьи на английском языке (Abstract) не должна быть дословным переводом русскоязычной аннотации. Раздел Abstract, как и основной текст, должен быть структурирован, в нем должно содержаться описание цели работы, методов исследования, научной значимости, выводов / результатов. Требуется качественный перевод на английский язык (при необходимости просим авторов обращаться к профессиональным переводчикам). **Объем Abstract 200–250 слов.**

Список литературы / References

Список литературы и список литературы на английском языке (References) размещаются в общем разделе. Рекомендуемое количество цитируемых в статье источников – не менее 10, в список желательно включать ссылки на актуальные работы по теме исследования, особенно в иностранных периодических изданиях.

В тексте статьи ссылки на литературу указываются цифрами в квадратных скобках, при необходимости указываются номера страниц, например: [2; 3. С. 15].

Список литературы нумеруется в порядке цитирования и оформляется в соответствии с ГОСТ Р 7.0.5-2008 на библиографическое описание (знаки тире в описании опускаются). Ссылки на неопубликованные работы, а также на интернет-ресурсы (кроме электронных изданий, поддающихся библиографическому описанию) оформляются в виде сноски.

В Список литературы ссылки на источники следует включать на оригинальном языке опубликования. Каждый источник должен быть также оформлен на английском языке (References) по международному стандарту для публикаций в области информатики IEEE Style со следующими отличиями:

- инициалы авторов указываются после фамилии;
- название статьи не берется в кавычки, отделяется точкой;

- отсутствует союз «and» перед фамилией последнего автора;
- в диапазоне страниц – удвоенная «р» (например, «pp. 2–9»);
- год издания указывается после места издания (для книг) и сразу после названия журнала (для периодики).
- Перевод источника на английский язык:
- если источник имеет выходные данные на английском языке, то для формирования References **следует использовать именно эти данные**;
- если оригинальная публикация не содержит выходных данных на английском языке, то допускается транслитерация названия материала на латинский алфавит в сочетании с переводом на английский язык в квадратных скобках. В конце описания указывается, на каком языке написана эта работа, например, (in Russ.). При транслитерации можно воспользоваться интернет-ресурсом <http://ru.translit.ru/>, рекомендуется выбрать стандарт BSI. Место издания не транслитерируется, указывается полностью на английском языке, например: Moscow. Название издательства / издателя, как правило, транслитерируется. Для журналов, у которых есть официальное название на английском языке, – использовать его (проверить на сайте журнала, или, например, в библиотеке WorldCat), если названия на английском языке нет, использовать транслитерацию по системе BSI. Не следует самостоятельно переводить названия журналов.

Если у цитируемого источника есть **цифровой идентификатор DOI** (<https://search.crossref.org/>), его требуется обязательно указывать в конце библиографической ссылки.

Примеры оформления ссылок. Каждый источник в том же пункте дублируется на английском языке (References).

Источник на русском языке, перевод на английский доступен в метаданных статьи

1. Журавлев С. С., Рудометов С. В., Окольников В. В., Шакиров С. Р. Применение модельно-ориентированного проектирования к созданию АСУ ТП опасных промышленных объектов // Вестник НГУ. Серия: Информационные технологии. 2018. Т. 16, № 4. С. 56–67. DOI 10.25205/1818-7900-2018-16-4-56-67

Zhuravlev S. S., Rudometov S. V., Okolnishnikov V. V., Shakirov S. R. Model-Based Design Approach for Development Process Control Systems of Hazardous Industrial Facilities. *Vestnik NSU. Series: Information Technologies*, 2018, vol. 16, no. 4, pp. 56–67. (in Russ.) DOI 10.25205/ 1818-7900-2018-16-4-56-67

Источник на английском языке. Оформляем согласно требованиям для References. Приводим только 1 раз.

2. Telnov V. I. Optimization of the Beam Crossing Angle at the ILC for E + e- and yy Collisions. *Journal of Instrumentation*, 2018, vol. 13, no. 03, pp. P03020–P03020. DOI 10.1088/1748-0221/13/03/p03020

Метаданные источника доступны только на русском языке

3. Жижимов О. Л., Федотов А. М., Шокин Ю. И. Технологическая платформа массовой интеграции гетерогенных данных // Вестник НГУ. Серия: Информационные технологии. 2013. Т. 11, вып. 1. С. 24–41.

Zhizhimov O. L., Fedotov A. M., Shokin Yu. I. Tekhnologicheskaya platforma massovoi integratsii geterogennykh dannykh [Technology Platform for the Mass Integration of Heterogeneous Data]. *Vestnik NSU. Series: Information Technologies*, 2013, vol. 11, no. 1, pp. 24–41. (in Russ.)

Сведения об авторах

Последний раздел статьи – информация об авторе / авторах **на русском и английском языках**:

- ФИО полностью, ученая степень, ученое звание;
- идентификаторы автора, такие как ResearcherID (всем авторам рекомендуется использовать данные сервисы для ведения актуального списка своих публикаций);
- контактный телефон (не публикуется).

Если статья представляется на английском языке, необходимо приложить перевод на русский язык названия, аннотации, ключевых слов, сведений об авторе.

Доставка материалов

Материалы предоставляются в редакцию по электронной почте inftech@vestnik.nsu.ru.

Порядок рецензирования

Все статьи сначала проходят проверку на заимствование и только после этого отправляются на рецензирование. Редакционный совет не допускает к публикации материал, если имеется достаточно оснований полагать, что он является плагиатом.

Тип рецензирования статей – двухуровневое, одностороннее анонимное («слепое»).

Для каждой статьи редколлегией выбираются рецензенты, научная деятельность которых связана с темой представленного материала. Ответственный секретарь журнала обращается к ним с просьбой дать экспертную оценку статье либо помочь организовать рецензирование.

Рецензии для журнала «Вестник НГУ. Серия: Информационные технологии» составляются по единой схеме и подразумевают оценку по следующим критериям: соответствие тематике журнала, оригинальность и значимость результатов, качество изложения материала.

Заполненный бланк рецензии высылается на электронный адрес редакции. В зависимости от экспертных заключений статья может быть принята редакционным советом к опубликованию, рекомендована автору к доработке (с последующим повторным рецензированием либо без него) или отклонена (с предоставлением автору мотивированного отказа). Автору на электронный адрес высылается текст рецензии без указания ФИО рецензента и его контактных данных.

Все рецензии хранятся в редакции журнала не менее 5 лет. Редколлегия журнала обязуется при поступлении соответствующего запроса направлять копии рецензий в Министерство науки и высшего образования Российской Федерации.

Вестник Новосибирского государственного университета
Серия: Информационные технологии
Научный журнал
2026. Т. 24, № 1

Учредитель
Новосибирский государственный университет

Адрес учредителя, издателя
ул. Пирогова, 2, Новосибирск, 630090, Россия

ВрИО главного редактора
д-р физ.-мат. наук, проф. М. М. Лаврентьев

Адрес редакции
ул. Пирогова, 1, Новосибирск, 630090, Россия

Дата выхода в свет 22.06.2026
Тираж 27 экз.
Цена 1329 руб.

Адрес типографии
ул. Пирогова, 2, Новосибирск, 630090, Россия