

УДК 004.912

Научная статья

DOI: 10.35330/1991-6639-2025-27-1-111-119

EDN: UJJXRM

Исследование основных методов автоматической обработки, группировки и аннотирования информации

Д. В. Тихонов

Финансовый университет при Правительстве Российской Федерации (Ярославский филиал)
150003, Россия, г. Ярославль, ул. Кооперативная, 12а

Аннотация. В статье исследованы основные методы автоматической обработки, группировки и аннотирования информации. Показано, что методы автоматического анализа Data Mining базируются на использовании определенных статистических закономерностей (классификация, регрессия), поиске ключевых слов, однако не используют алгоритмы лингвистической обработки текстов. Таким образом, автоматический анализ текстовой информации, осуществляемый современными средствами аналитической обработки, не способен прорабатывать содержание текстов. Для сравнения двух простых предложений по содержанию был использован метод резолюций. Как показали исследования, при применении алгоритма унификации содержание предложений не учитывается. Поэтому как решение проблемы сравнительного анализа текстовой информации по содержанию были предложены новые алгоритмы работы с логико-лингвистическими моделями. Научная новизна полученных результатов состоит в методе быстрого извлечения набора локальных дескрипторов, описывающих все части изображения, что позволяет существенно ускорить процесс аннотирования и формировать более полный глобальный визуальный дескриптор изображения.

Ключевые слова: методы, автоматическая обработка, группировка, аннотирование, информация, Data Mining, метод резолюций

Поступила 25.12.2024, одобрена после рецензирования 28.01.2025, принята к публикации 03.02.2025

Для цитирования. Тихонов Д. В. Исследование основных методов автоматической обработки, группировки и аннотирования информации // Известия Кабардино-Балкарского научного центра РАН. 2025. Т. 27. № 1. С. 111–119. DOI: 10.35330/1991-6639-2025-27-1-111-119

MSC: 68U15

Original article

Research on main methods of automatic processing, grouping and annotation of information

D.V. Tikhonov

Financial University under the Government of the Russian Federation (Yaroslavl branch)
150003, Russia, Yaroslavl, 12a Cooperative street

Abstract. The article examines main methods of automatic processing, grouping and annotation of information. It is shown that methods of automatic analysis of Data Mining are based on the use of certain statistical patterns (classification, regression), search for keywords, but do not use algorithms for linguistic processing of texts. Thus, automatic analysis of text information carried out by modern means of analytical processing is not able to process texts content. To compare two simple sentences by content, the resolution method was used. Studies have shown the unification algorithm does not take into account

content of sentences. Because the solution to the problem of comparative analysis of text information by content, a new algorithm for working with logical and linguistic models is proposed. The scientific novelty of the results obtained lies in the method of quickly extracting a set of local descriptors describing all parts of the image, which allows us to significantly speed up the annotation process and form a more complete global visual descriptor of the image.

Keywords: methods, automatic processing, grouping, annotation, information, Data Mining, resolution method

Submitted 25.12.2024,

approved after reviewing 28.01.2025,

accepted for publication 03.02.2025

For citation. Tikhonov D.V. Research on main methods of automatic processing, grouping and annotation of information. *News of the Kabardino-Balkarian Scientific Center of RAS*. 2025. Vol. 27. No. 1. Pp. 111–119. DOI: 10.35330/1991-6639-2025-27-1-111-119

ВВЕДЕНИЕ

Автоматическая обработка информации относится к умственному когнитивному процессу с широким диапазоном характеристик: быстрый, параллельный, эффективный, требующий небольших когнитивных усилий и активного контроля или внимания со стороны субъекта. Этот тип обработки является результатом повторяющегося обучения одной и той же задаче. Однажды выученную автоматическую реакцию трудно подавить, изменить или игнорировать. Автоматическая обработка информации используется для решения квалифицированных задач и считается процессом, противоположным контролируемой обработке информации.

В 1950-х годах область когнитивной психологии сосредоточилась на ограничениях возможностей обработки информации человеком, например, на том, как мозг обрабатывает входящую информацию в виде стимулов. В 1958 году британский психолог Бродбент представил важную модель обработки информации с ограниченной пропускной способностью и был одним из первых, кто провел различие между автоматическими и контролируемыми процессами [1].

Задачу автоматизированной аналитической обработки текстовой информации пытаются решить многие иностранные и отечественные ученые. В частности, еще в 1979 году Н. Т. Кузин [2] описал методы частотной обработки текстовой информации, которые впоследствии были усовершенствованы в работах А. Бродера [3] и Д. В. Ланде [4]. Ученые в публикациях [5–7] обобщили данные по современным методам автоматического анализа Data Mining и Text Mining. Однако ни один из описанных методов не обеспечивает извлечение из текстовой информации знаний.

Цель статьи – проанализировать основные методы автоматической обработки, группировки и аннотирования информации, выявить их преимущества и недостатки для задач извлечения знаний из естественного языка.

В рамках достижения данной цели необходимо осуществить сравнительный анализ простых предложений природного языка с помощью метода резолюций Робинсона, обрисовать главные шаги метода сравнительного анализа текстовой информации, представленной в виде логико-лингвистических моделей.

МАТЕРИАЛЫ И МЕТОДЫ

Материалами исследования выступают научные публикации по тематике использования основных методов автоматической обработки, группировки и аннотирования информации. Используются методы анализа, обобщения и систематизации.

РЕЗУЛЬТАТЫ И ИХ ОБСУЖДЕНИЕ

Автоматизированная обработка информации включает в себя сбор компонентов информации, присутствующих в документе, с помощью программного обеспечения. Она использует такие технологии, как машинное обучение, компьютерное зрение, обработка естественного языка и распознавание текста. Автоматическая обработка документов в организации помогает сократить ручной труд, соблюдать требования, устранить проблемы и повысить скорость рабочих процессов.

Существует четыре распространенных способа обработки информации, проанализируем каждый из них [8].

1. Ручная обработка информации и документов. Ручная обработка документов подразумевает обработку соответствующей и важной информации из документов вручную и упорядочивание этих данных для принятия решений. Этот метод представляет собой трудоемкий процесс, обработка одного документа может занять до 20 минут (иногда и больше). Когда дело доходит до точности ручной обработки данных, она сравнительно низка по сравнению с другими доступными методами обработки данных и составляет всего 60–70 процентов точности. Этот метод также требует ручной рабочей силы для выполнения всей операции.

2. Компьютерное зрение. Компьютерное зрение относится к обучению компьютера работе с рядом форматов документов, чтобы обеспечить возможность идентификации символов и других элементов, управляемых данными, из документа. Это современная техника обработки данных, которая использует искусственный интеллект для получения значимых данных из изображений, видео, документов или чего-либо, что имеет цифровое или аналоговое существование. Это лучше объяснить искусственным интеллектом, который может заставить компьютер думать. Оно помогает видеть объекты, делать наблюдения, а затем понимать. Компьютерное зрение управляет другими методами, такими как оптическое распознавание меток и оптическое распознавание символов, и является расширенным набором этих методов обработки данных.

Этот процесс предполагает использование большого количества данных при повторном анализе до тех пор, пока не будут распознаны различия и данные из изображений или документов.

Компьютерное зрение использует две разные технологии – глубокое обучение и CNN – для достижения четкого распознавания. CNN относится к нейронным сетям свертки, которые помогают модели машинного обучения просматривать изображения, разбитые на пиксели с метками или тегами, для выполнения сверток и прогнозирования.

3. Оптическое распознавание символов. Оптическое распознавание символов, или OCR, идентифицирует данные из документов в форме символов и изображений и далее обрабатывает эти данные в поддающиеся учету форматы. Эти извлеченные данные затем преобразуются в машиночитаемую форму, которая в дальнейшем используется для обработки данных. OCR обрабатывает цифровые файлы, такие как квитанции о трудоустройстве, счета-фактуры, контракты, финансовые отчеты и т. д.

Оптическое распознавание символов помогает автоматизировать обработку документов и извлечение данных, что в конечном итоге позволяет организациям экономить драгоценные ресурсы и время. Эта технология анализирует текст, присутствующий на странице, идентифицирует символы и в дальнейшем превращает их в код, поддерживающий обработку информации в документе. Он имеет трехэтапную процедуру, которая включает предварительную обработку, распознавание символов и постобработку.

4. Интеллектуальная обработка документов IDP. IDP означает интеллектуальную обработку документов, которая преобразует полуструктурированную или неструктурированную информацию из документа в полезные данные. Примерно 80% данных всех организаций хранятся в полуструктурированной и неструктурированной форме, например, в счетах-фактурах, отчетах о прибылях и убытках и балансовых отчетах. Интеллектуальная обработка документов внесла революционные изменения в следующее поколение автоматизации обработки данных благодаря чрезвычайно быстрой обработке и таким возможностям, как извлечение и обработка документов различных форматов.

Автоматизированная система управления документами использует технологии искусственного интеллекта, такие как обработка естественного языка, глубокое обучение, компьютерное зрение и машинное обучение, для классификации, категоризации и извлечения актуальной и важной информации, в конечном итоге проверяя данные. IDP – это следующий шаг в области оптического распознавания символов, поскольку преодолевает ограничения OCR при извлечении данных из всех нестандартных и сложных документов. Он имеет высокую точность, близкую к 100%, и обладает более быстрой функциональностью, чем другие методы извлечения данных, с возможностью обработки данных из сложных структур документов [9].

Для повышения точности анализа текстов разрабатываются методы предварительной лингвистической обработки, что требует, во-первых, значительных вычислительных затрат для лингвистического анализа индексированной коллекции текстов, во-вторых, разработки специализированной поисковой машины. Автоматизированное извлечение знаний из текста является одной из основных задач искусственного интеллекта и напрямую связано с пониманием текстов на естественном языке.

На сегодня существуют различные средства обработки текстовой информации. Для извлечения знаний из текстовой информации используются разные методы автоматического анализа Data Mining. Такие методы используют алгоритмы и средства искусственного интеллекта для исследования больших объемов информации и извлечения знаний, которые будут практически полезны и доступны для интерпретации человеком [5].

Основными методами Data Mining являются классификация, кластеризация, регрессия, поиск ассоциативных правил, аннотирование и автореферирование. Задача классификации сводится к определению класса объекта по его характеристикам, причем множество классов задается раньше времени. Классификация использует статистические корреляции для построения правил размещения документов в заданной категории; задача классификации – это задача распознавания, когда система относит новый объект к той или иной категории.

Классификация и регрессия предполагают осуществление двух обязательных этапов. Первый этап – выделение набора объектов, для которых известны значения зависимых и независимых переменных. На основе полученного набора строится модель определения значения зависимой переменной (функция классификации или регрессии). На втором этапе построенную модель применяют к анализируемым объектам. Недостатком классификации и регрессии является то, что разработчик системы должен фиксировать количество классов и характеристик, по которым будут проводиться исследования. Это означает, что если система не выявит признак или класс, к которому можно отнести, например, текстовый документ, он не будет корректно обработан.

Аннотирование – это процесс создания коротких сообщений об электронном тексте, позволяющих делать выводы о целесообразности его подробного изучения [10, 11].

Современные системы аналитической обработки текстовой информации обладают средствами автоматического составления аннотаций.

Метод аннотирования текста произвольной структуры предусматривает:

1. Формирование множества аннотированных фрагментов, являющихся целыми предложениями данного текста, содержат в своем составе глагол или краткое прилагательное и не являются вопросительным или восклицательным предложением.
2. Создание таблицы всех вероятных пар главных тематических узлов (здесь употребляется система продукций для установления черт структурных единиц текста, описанная ранее).
3. Отбор таких предложений, содержащих несколько различных тематических узлов, не встречавшихся ранее в тексте.

Осуществление автоматической аннотации является прикладной задачей, которая решается перед тем, как информация из заданного текста попадет в поисковый сервер. Автоматическое реферирование представляет собой создание кратких изложений материалов, аннотаций, дайджестов, то есть извлечение наиболее важных сведений из одного или нескольких документов и генерация на их основе лаконичных и информационно емких отчетов.

На сегодняшний день существует два основных направления автореферирования: квазиреферирование (основано на экстрагировании фрагментов документов, то есть выделении наиболее информативных фраз и формирование из них квазирефератов) и краткое изложение содержания первичных документов (дайджесты) [3].

Автоматическое реферирование и аннотирование используется в основном для экономии времени пользователям, создания каталогов информационных ресурсов, использования словарей-тезаурусов общего и специального назначения. Применяется автоматическое реферирование и аннотирование в корпоративных системах документооборота, поисковых машинах и каталогах ресурсов Интернет, автоматизированных информационно-библиотечных системах, каналах связи, службах рассылки новостей и т.д.

Поиск ассоциативных правил представляет собой способ поиска частичных зависимостей между объектами и субъектами. Найденные зависимости представляются в виде правил и используются для лучшего понимания природы анализируемых данных. То есть из большого количества наборов объектов определяются наиболее часто встречающиеся. При выявлении закономерностей можно с определенной вероятностью предсказать появление событий в будущем, что позволяет принимать решения. Такая задача является разновидностью задачи поиска ассоциативных правил и называется сиквенционным анализом.

Кластеризация – это разбиение множества документов на кластеры (группы документов с общими признаками), которые представляют собой подмножества, смысловые параметры которых заранее неизвестны. Численные методы кластеризации базируются на определении кластера как множества документов. Для задачи кластеризации характерен поиск групп наиболее схожих объектов.

Результат кластеризации зависит от природы данных и от представления кластеров.

Все описанные выше методы автоматического анализа Data Mining обеспечивают определенную структуризацию текстовой информации, ее обобщение или аннотирование. Однако для извлечения знаний из электронных текстов, в частности сравнения и выявления в них совпадений, необходимы средства автоматического лингвистического анализа. Основным методом, используемым сегодня для логического сравнения текстовой информации, является метод резолюций Робинсона.

Например, пусть есть два простых предложения, для каждого из которых построена логическая модель.

Первое предложение: "Эксперт отвечает на вопросы слушателей":

$$P(x_1, x_2[x_3]). \quad (1)$$

Соответствует (эксперт, вопросы [слушателей]).

Второе предложение: "Эксперт анализирует вопросы специалистов":

$$P'(x_1, x_2[x_3']). \quad (2)$$

Анализирует (эксперт, вопросы [специалистов]).

По алгоритму унификации ищем подстановку $Q = \{P/P', x_3/x_3'\}$. Если осуществить подстановку в выражение (2), будем иметь множества, содержащие литералы с одинаковыми предикатами.

После этого, применяя метод резолюций к выражениям (1) и (2), получим резольвенту, что не равно пустому множеству, это свидетельствует о том, что предложения одинаковы.

Анализируя содержание заданных предложений, можно сделать вывод, что в данном случае в алгоритме унификации нельзя было применять подстановку P/P' , так как предикаты разные по содержанию и не являются синонимами.

Метод резолюций не позволяет определить это в процессе замены, потому что не анализирует содержание слов, входящих в предложение естественного языка [7]. Это означает, что для корректного сравнения текстовых документов по содержанию необходимы новые алгоритмы лингвистического анализа, которые обеспечат содержательную обработку текстовой информации.

Одним из таких алгоритмов может быть алгоритм сравнения логико-лингвистических моделей предложений естественного языка, включающий следующие этапы.

1. Построение логико-лингвистических моделей [8]. На этом этапе каждому предложению природного языка ставится в соответствие логическая формула, представляющая собой одномерный массив слов, из которых состоят предложения, упорядоченные в соответствии с тем, какую синтаксическую роль они выполняют.

2. Идентификация. Происходит очередной просмотр элементов всех логико-лингвистических моделей: предикатов, предикатных переменных (субъектов), предикатных переменных (аргументов), предикатных констант. Среди составляющих логико-лингвистических моделей ищутся однокоренные слова, синонимы, активные и пассивные формы однокоренных глаголов.

3. Замена тождественных переменных. Если на этапе идентификации найдены тождественные переменные, во всех логико-лингвистических моделях происходит их переобозначение, благодаря чему одни и те же слова (даже если они имеют разные грамматические рамки) будут обозначаться одинаково и соответственно иметь идентичное содержание.

4. Логический вывод. После идентификации и замены тождественных переменных применяется система продукции, содержащая правила сравнения логико-лингвистических моделей. Такие правила позволяют через установленные связи между словами переходить к представлению значений слов посредством комбинаций элементарных компонентов содержания.

В качестве эффективного метода автоматической обработки, группировки и аннотирования информации был предложен метод автоматического аннотирования изображений

на основе обучающего набора изображений, разделенного на однородные текстово-визуальные группы, а также предложен алгоритм для реализации данного метода, который отличается тем, что аннотирование нового изображения осуществляется с помощью обучающих изображений небольшого количества визуально схожих групп. В рамках алгоритма для всех учебных, а также аннотированных изображений образуется глобальный визуальный дескриптор. Для этого с изображения извлекается набор локальных дескрипторов, который кодируется с помощью словаря визуальных слов. Поскольку этот этап автоматической обработки, группировки и аннотирования информации является наиболее вычислительно затратным, предложен быстрый метод извлечения набора локальных дескрипторов, что позволяет избежать повторных вычислений при наложении областей расчета дескрипторов и существенно ускоряет процесс аннотирования. Также предложен процесс вычисления цветовых локальных дескрипторов, использование которых позволяет повысить точность аннотирования, алгоритмы формирования словаря визуальных слов и кодирование набора локальных дескрипторов в глобальный визуальный дескриптор.

Выводы

Методы автоматического анализа Data Mining базируются на использовании определенных статистических закономерностей (классификация, регрессия), поиске ключевых слов, однако не используют алгоритмы лингвистической обработки текстов. Таким образом, автоматический анализ текстовой информации, осуществляемый современными средствами аналитической обработки, не способен прорабатывать содержание текстов. Для сравнения двух простых предложений по содержанию был использован метод резолюций. Как показали исследования, при применении алгоритма унификации содержание предложений не учитывается. Поэтому как решение проблемы сравнительного анализа текстовой информации по содержанию предложен новый алгоритм работы с логико-лингвистическими моделями. Предложен метод быстрого извлечения набора локальных дескрипторов, описывающих все части изображения, что позволяет существенно ускорить процесс аннотирования и сформировать более полный глобальный визуальный дескриптор изображения.

СПИСОК ЛИТЕРАТУРЫ

1. *Hammar A.* Automatic Information Processing. In: Seel, N.M. (eds). Encyclopedia of the Sciences of Learning. Springer. Boston, MA. 2012. DOI: 10.1007/978-1-4419-1428-6_494
2. *Khazaei E., Alimohammadi A.* An Automatic User Grouping Model for a Group Recommender System in Location-Based Social Networks. ISPRS Int. J. Geo-Inf. 2018. 7(2):67. DOI: 10.3390/ijgi7020067
3. *Ячная В. О., Луцив В. Р., Малашин Р. О.* Современные технологии автоматического распознавания средств общения на основе визуальных данных // КО. 2023. Т. 47. № 2. С. 287–305. URL: <https://cyberleninka.ru/article/n/sovremennye-tehnologii-avtomaticheskogo-raspoznavaniya-sredstv-obscheniya-na-osnove-vizualnyh-dannyh> (дата обращения: 18.02.2024)
4. *Назаров Т. Р., Мамедова Н. А.* Автоматизированное решение задачи детектирования промышленных объектов на ортофотоплане с помощью нейронной сети // Программные продукты и системы. 2023. Т. 36. № 1. С. 144–158. URL: <https://cyberleninka.ru/article/n/avtomatizirovannoe-reshenie-zadachi-detektirovaniya-promyshlennyh-obektov-na-ortofotoplane-s-pomoschyu-neyronnoy-seti> (дата обращения: 18.02.2024)

5. Власов С. О., Гладышев А. И., Богуславский А. А., Соколов С. М. Решение задачи обнаружения объекта с помощью нейросетевых технологий // Препринты ИПМ им. М. В. Келдыша. 2023. № 16. 27 с. DOI: <https://doi.org/10.20948/prepr-2023-16>
6. Гайсин А. Э. Анализ существующих методов автоматического текстового анализа // Вестник науки. 2023. № 6(63). Т. 4. С. 254–258. URL: <https://cyberleninka.ru/article/n/analiz-suschestvuyuschih-metodov-avtomaticheskogo-tekstovogo-analiza> (дата обращения: 18.02.2024)
7. Пригодич Н. Д., Коробко С. С. Применение программных методов для автоматизированной обработки источников личного происхождения // Историческая информатика. 2023. № 1(43). С. 1–9. URL: <https://cyberleninka.ru/article/n/primenenie-programmnyh-metodov-dlya-avtomatizirovannoy-obrabotki-istochnikov-lichnogo-proishozhdeniya> (дата обращения: 18.02.2024)
8. Li H., Yuan D., Ma X., Cui D., Cao L. Genetic algorithm for the optimization of features and neural networks in ECG signals classification // Scientific Reports. 2017. Vol. 7. No. 1. Pp. 1–12. DOI: 10.1038/srep41011
9. He X., Zhao K., Chu X. AutoML: A Survey of the State-of-the-Art // Knowledge-Based Systems. 2021. Vol. 212. P. 106622. DOI: 10.1016/j.knosys.2020.106622
10. Jin H., Song Q., Hu X. Auto-keras: An efficient neural architecture search system // Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining. 2019. Pp. 1946–1956. DOI: 10.1145/3292500.3330648
11. Real E., Aggarwal A., Huang Y., Le Q. V. Regularized evolution for image classifier architecture search // Proceedings of the AAAI Conference on Artificial Intelligence. 2018. Vol. 33. Pp. 4780–4789. DOI: 10.1609/aaai.v33i01.33014780

REFERENCES

1. Hammar Å. Automatic Information Processing. In: Seel, N.M. (eds). Encyclopedia of the Sciences of Learning. Springer. Boston, MA. 2012. DOI: 10.1007/978-1-4419-1428-6_494
2. Khazaei E., Alimohammadi A. An Automatic User Grouping Model for a Group Recommender System in Location-Based Social Networks. *ISPRS Int. J. Geo-Inf.* 2018. 7(2):67. DOI: 10.3390/ijgi7020067
3. Yachnaya V.O., Lutsiv V.R., Malashin R.O. *Sovremennyye tekhnologii avtomaticheskogo raspoznavaniya sredstv obshcheniya na osnove vizual'nykh dannykh* [Modern technologies for automatic recognition of means of communication based on visual data]. KO. 2023. Vol. 47. No. 2. Pp. 287–305. URL: <https://cyberleninka.ru/article/n/sovremennyye-tehnologii-avtomaticheskogo-raspoznavaniya-sredstv-obshcheniya-na-osnove-vizualnykh-dannykh> (access date: 02.18.2024). (In Russian)
4. Nazarov T.R., Mamedova N.A. Automated solution to the problem of detecting industrial objects on an orthomosaic using a neural network. *Programmnyye produkty i sistemy* [Software products and systems]. 2023. Vol. 36. No. 1. Pp. 144–158. URL: <https://cyberleninka.ru/article/n/avtomatizirovannoe-reshenie-zadachi-detektirovaniya-promyshlennykh-obektov-na-ortofotoplane-s-pomoschyu-neuronnoy-seti> (date of access: 02.18.2024). (In Russian)
5. Vlasov S.O., Gladyshev A.I., Boguslavsky A.A., Sokolov S.M. Solving the problem of object detection using neural network technologies. Preprints of the Keldysh IPM. 2023. No. 16. 27 p. DOI: <https://doi.org/10.20948/prepr-2023-16>. (In Russian)
6. Gaisin A.E. Analysis of existing methods of automatic text analysis. *Vestnik nauki* [Bulletin of Science]. 2023. No. 6(63). Vol. 4. Pp. 254–258. URL: <https://cyberleninka.ru/article/n/analiz-suschestvuyuschih-metodov-avtomaticheskogo-tekstovogo-analiza> (date of access: 02.18.2024). (In Russian)

7. Prigodich N.D., Korobko S.S. Application of software methods for automated processing of sources of personal origin. *Istoricheskaya informatika* [Historical informatics]. 2023. No. 1(43). Pp. 1–9. URL: <https://cyberleninka.ru/article/n/primenenie-programmnyh-metodov-dlya-avtomatizirovannoy-obrabotki-istochnikov-lichnogo-proishozhdeniya> (date of access: 02.18.2024). (In Russian)
8. Li H., Yuan D., Ma X., Cui D., Cao L. Genetic algorithm for the optimization of features and neural networks in ECG signals classification. *Scientific Reports*. 2017. Vol. 7. No. 1. Pp. 1–12. DOI: 10.1038/srep41011
9. He X., Zhao K., Chu X. AutoML: A Survey of the State-of-the-Art. *Knowledge-Based Systems*. 2021. Vol. 212. P. 106622. DOI: 10.1016/j.knsys.2020.106622
10. Jin H., Song Q., Hu X. Auto-keras: An efficient neural architecture search system. *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. 2019. Pp. 1946–1956. DOI: 10.1145/3292500.3330648
11. Real E., Aggarwal A., Huang Y., Le Q. V. Regularized evolution for image classifier architecture search. *Proceedings of the AAAI Conference on Artificial Intelligence*. 2018. Vol. 33. Pp. 4780–4789. DOI: 10.1609/aaai.v33i01.33014780

Финансирование. Исследование проведено без спонсорской поддержки.

Funding. The study was performed without external funding.

Информация об авторе

Тихонов Дмитрий Владимирович, канд. техн. наук, доцент кафедры «Экономика и финансы», Финансовый университет при Правительстве Российской Федерации (Ярославский филиал); 150003, Россия, г. Ярославль, ул. Кооперативная, 12а; Dtihonov1987@yandex.ru, ORCID: <https://orcid.org/0009-0001-2293-6390>, SPIN-код: 4195-0317

Information about the author

Dmitry V. Tikhonov, Candidate of Engineering Sciences, Associate Professor of the Department of Economics and Finance, Financial University under the Government of the Russian Federation (Yaroslavl branch); 150003, Russia, Yaroslavl, 12a Cooperative street; Dtihonov1987@yandex.ru, ORCID: <https://orcid.org/0009-0001-2293-6390>, SPIN-code: 4195-0317