



УНИВЕРСИТЕТ ИТМО

3



АЛЬМАНАХ

НАУЧНЫХ РАБОТ
МОЛОДЫХ
УЧЕНЫХ

2019

**МИНИСТЕРСТВО НАУКИ И ВЫСШЕГО ОБРАЗОВАНИЯ
РОССИЙСКОЙ ФЕДЕРАЦИИ**

УНИВЕРСИТЕТ ИТМО

**АЛЬМАНАХ
НАУЧНЫХ РАБОТ
МОЛОДЫХ УЧЕНЫХ
Университета ИТМО**

Том 3



УНИВЕРСИТЕТ ИТМО

Санкт-Петербург

2019

Альманах научных работ молодых ученых Университета ИТМО. Том 3. – СПб.: Университет ИТМО, 2019. – 210 с.

Издание содержит результаты научных работ молодых ученых, доложенные на XLVIII научной и учебно-методической конференции Университета ИТМО по тематикам: информационные технологии и программирование; инфокоммуникационные технологии; дизайн и урбанистика.

РЕДАКЦИОННАЯ КОЛЛЕГИЯ

Председатель редколлегии:

Бухановский Александр Валерьевич

доктор технических наук, директор мегафакультета трансляционных информационных технологий Университета ИТМО.

Члены редколлегии:

Зудилова Татьяна Викторовна

кандидат технических наук, доцент факультета инфокоммуникационных технологий

Хоружников Сергей Эдуардович

кандидат физико-математических наук, декан факультета инфокоммуникационных технологий

Парфенов Владимир Глебович

доктор технических наук, декан факультета информационных технологий и программирования

Матвеев Юрий Николаевич

доктор технических наук, профессор факультета информационных технологий и программирования

Духанов Алексей Валентинович

доктор технических наук, профессор института финансовых кибертехнологий

ISBN 978-5-7577-0606-1

ISBN 978-5-7577-0609-2 (Том 3)



УНИВЕРСИТЕТ ИТМО

Университет ИТМО – ведущий вуз России в области информационных и фотонных технологий, один из немногих российских вузов, получивших в 2009 году статус национального исследовательского университета. С 2013 года Университет ИТМО – участник программы повышения конкурентоспособности российских университетов среди ведущих мировых научно-образовательных центров, известной как проект «5 в 100». Цель Университета ИТМО – становление исследовательского университета мирового уровня, предпринимательского по типу, ориентированного на интернационализацию всех направлений деятельности.

© Университет ИТМО, 2019

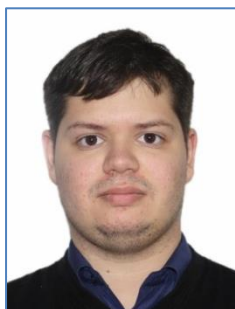
© Авторы, 2019

ВВЕДЕНИЕ

Издание содержит результаты научных работ молодых ученых, доложенные 29 января – 1 февраля 2019 года на XLVIII научной и учебно-методической конференции Университета ИТМО по тематике: информационные технологии и программирование; инфокоммуникационные технологии; дизайн и урбанистика.

Конференция проводится в целях усиления интегрирующей роли университета в области научных исследований по приоритетным направлениям развития науки, технологий и техники и ознакомления научной общественности с результатами исследований, выполненных в рамках государственного задания Министерства образования и науки РФ, программы развития Университета ИТМО на 2009–2018 годы, программы повышения конкурентоспособности Университета ИТМО среди ведущих мировых научно-образовательных центров на 2013–2020 гг., Федеральной целевой программы «Исследования и разработки по приоритетным направлениям развития научно-технологического комплекса России на 2014–2020 годы», грантов Президента РФ для поддержки молодых российских ученых и ведущих научных школ, грантов РФФИ, РГНФ, РНФ и Правительства РФ (по постановлению № 220 от 09.04.2010 г.) и по инициативным научно-исследовательским проектам, проводимым учеными, преподавателями, научными сотрудниками, аспирантами, магистрантами и студентами университета, в том числе в содружестве с предприятиями и организациями Санкт-Петербурга, а также с целью повышения эффективности научно-исследовательской деятельности и ее вклада в повышение качества подготовки специалистов.

**НАПРАВЛЕНИЕ
ИНФОРМАЦИОННЫЕ ТЕХНОЛОГИИ
И ПРОГРАММИРОВАНИЕ**



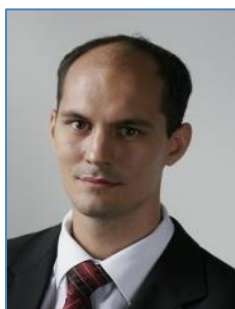
Большаков Никита Андреевич

Год рождения: 1996

Университет ИТМО, факультет информационных технологий и программирования, студент группы № M41212

Направление подготовки: 09.04.02 – Информационные системы и технологии

e-mail: snowmanikita@gmail.com



Шуранов Евгений Витальевич

Год рождения: 1980

Университет ИТМО, факультет информационных технологий и программирования, к.т.н.

e-mail: shuranov@speechpro.com

УДК 004.93

ПРИМЕНЕНИЕ ИСКУССТВЕННЫХ НЕЙРОННЫХ СЕТЕЙ ПРИ ДЕТЕКТИРОВАНИИ АКУСТИЧЕСКИХ СОБЫТИЙ

Большаков Н.А.

Научный руководитель – к.т.н. Шуранов Е.В.

Работа выполнена в рамках темы НИР № 618278 «Синтез эмоциональной речи на основе генеративных состязательных сетей».

В работе выполнен обзор методов и моделей применения нейронных сетей для детектирования акустических событий.

Ключевые слова: акустические события, детектирование, искусственные нейронные сети.

Акустическое событие – это звуковое явление определенной длительности, определенного времени возникновения, исчезновения и типа. Исследования в области детектирования таких событий находят применение при решении множества прикладных задач, связанных с определением вышеперечисленных параметров. В качестве примера можно привести задачу классификации акустической обстановки [1].

Искусственные нейронные сети (ИНС) являются математическим воплощением своих биологических аналогов. Идея воспользоваться живой нейронной системой как оригиналом для подражания при построении модели была высказана в 40-х годах прошлого века [2].

Долгое время использование ИНС было неэффективно, так как распознавание какого-либо образа достигалось только путем запоминания практически всех возможных его состояний. Поэтому широкое распространение технология получила относительно недавно, с появлением совокупности семейств различных методов машинного обучения, впоследствии названных глубоким обучением.

В этой работе предложено рассмотрение способов построения ИНС, реализованных командой Констанцкого университета в рамках DCASE Community Challenge 2018 [3]. Также в работе приведены особенности получения крупной нейронной сети с использованием нескольких предобученных высокоуровневых ИНС.

Визуальное представление структуры ИНС, используемых в первых двух анализируемых методах представлено на рис. 1.

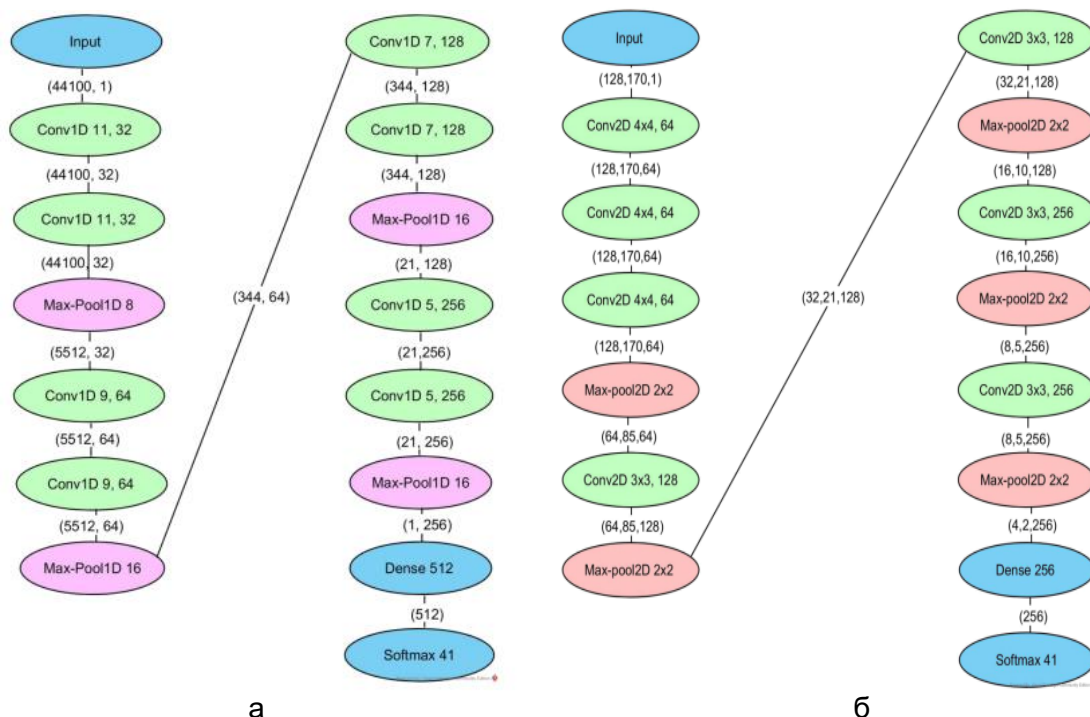


Рис. 1. Визуальное представление используемых ИНС: с обучением на необработанных данных (а); с обучением на спектрограммах (б)

Conv – это основные слои сверточной нейронной сети. Данные слои содержат фильтры для разных каналов, и обрабатывают предыдущие слои по частям с помощью так называемых ядер свертки. Как видно, обе используемые ИНС схожи, различие заключается в предварительной обработке данных акустического события. Для второй нейронной сети строятся спектрограммы, тогда возникает возможность обучения по двумерному изображению, где будут учитываться не только соседние значения в одномерном массиве, как при обычной интерпретации акустического события [4].

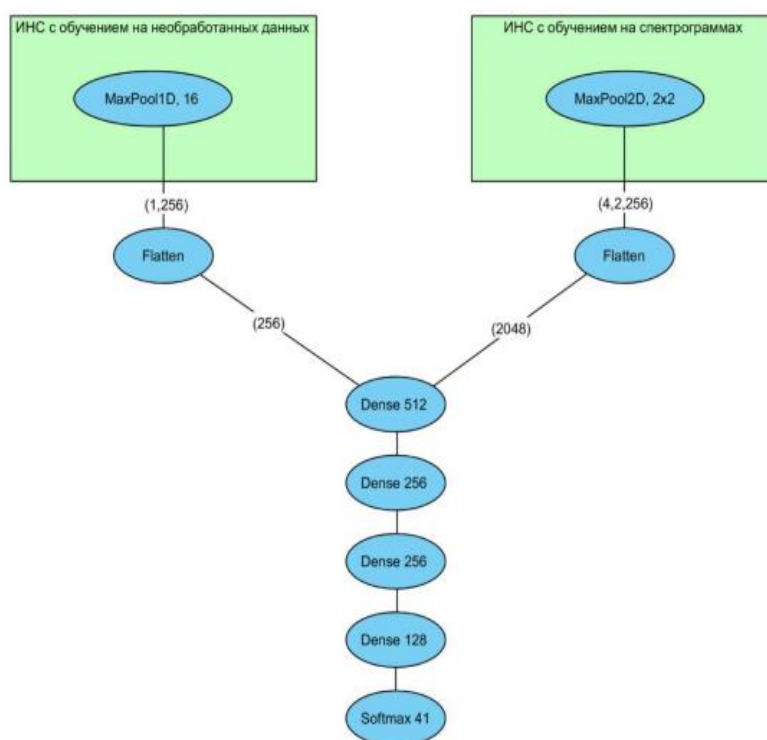


Рис. 2. Визуальное представление ИНС с обучением на комбинированных данных

Max-pooling – слой подвыборки, в котором происходит нелинейное уплотнение карты признаков, где несколько выбранных значений с помощью функции максимума уплотняются до одного [4].

Вышеописанные нейронные сети используют как и многослойные, но единственные для метода модели. В связи с этим интерес представляет возможность использования при обучении, кроме изначальных данных, весовые коэффициенты других ИНС. Визуальное представление реализации подобной высокоуровневой сети отражено на рис. 2.

Как видно из изображения (рис. 2), при построении были использованы уже знакомые нам сети с обучением на спектрограммах и необработанных данных, правда, уже без своих Dense-слоев и слоя Softmax, которые расширены и предоставлены последней моделью. В итоге подобный подход реализует собой очень простой и действенный способ повышения качества решения задачи путем минимальных усилий, в случае наличия нескольких простых ИНС.

Результаты проведенных тестирований данной нейронной сети представлены в таблице. Тестирование проводилось при обучении на 9000 объектов с 41 меткой. Поскольку задачей исследования являлось именно изучение предложенных методов, то результаты представлены на DCASE Community Challenge 2018, но они не только подтверждают эффективность изложенных методов, но и позволяют оценить повышение точности решений при использовании различных ИНС. Все методы были реализованы с использованием открытой нейросетевой библиотеки Keras, на языке программирования Python, с использованием вышеизложенных визуальных представлений в качестве основы.

Таблица. Результаты тестирования

Тестирование	Необработанные данные		Спектрограммы		Комбинированные данные	
Результат	Best Score	Worst Score	Best Score	Worst Score	Best Score	Worst Score
Значение	0,801	0,789	0,836	0,812	0,863	0,824

Литература

1. Учим компьютер различать звуки: знакомство с конкурсом DCASE и сборка своего аудио классификатора за 30 минут [Электронный ресурс]. – Режим доступа: <https://habr.com/ru/company/speechpro/blog/437818/> (дата обращения: 18.03.2019).
2. Мак-Каллок У.С., Питтс В. Логическое исчисление идей, относящихся к нервной активности. Архивная копия от 27 ноября 2007 на Wayback Machine // Автоматы. – 1956. – С. 363–384 (Перевод английской статьи 1943 г.).
3. Combining high-level features of raw audio and spectrograms for audio tagging [Электронный ресурс]. – Режим доступа: http://dcase.community/documents/challenge2018/technical_reports/DCASE2018_Wilhelm_109.pdf (дата обращения: 18.03.2019).
4. Keras: The Python Deep Learning library [Электронный ресурс]. – Режим доступа: <https://keras.io/> (дата обращения: 18.03.2019).

**Виноградова Таисия Борисовна**

Год рождения: 1995

Университет ИТМО, факультет информационных технологий
и программирования, студент группы № М41212Направление подготовки: 09.04.02 – Информационные системы
и технологии

e-mail: taisiyavinogr@gmail.com

**Рыбин Сергей Витальевич**

Год рождения: 1959

Университет ИТМО, факультет информационных технологий
и программирования, к.ф.-м.н., доцент

e-mail: rybin@speechpro.com

УДК 004.93

**СРАВНИТЕЛЬНЫЙ АНАЛИЗ МЕТОДОВ ОПРЕДЕЛЕНИЯ ЭМОЦИОНАЛЬНОЙ
ОКРАСКИ КОРОТКИХ ТЕКСТОВ****Виноградова Т.Б.****Научный руководитель – к.ф.-м.н., доцент Рыбин С.В.**

Работа выполнена в рамках темы НИР № 618278 «Синтез эмоциональной речи на основе генеративных состязательных сетей».

Проведен сравнительный обзор различных методов определения эмоциональной окраски текстов. Проанализируема их применимость к обработке коротких текстов на русском языке. Рассмотрены различные варианты постановки задачи и методы ее решение. Алгоритмы, основанные на использовании искусственных нейронных сетей, рассмотрены более подробно.

Ключевые слова: обработка естественного языка, анализ тональности, классификация текстов, нейронные сети, векторные представления текстов.

Автоматический анализ эмоциональной составляющей текстов играет важную роль при разработке систем синтеза эмоциональной речи. Эмотивная окраска речи влияет на восприятие человеком синтезированной речи. Искусственная речь, синтезированная с учетом особенностей эмоциональной окраски текста, на котором она основана, воспринимается более похожей на живую человеческую речь.

Рассмотрена задача автоматического определения эмоциональной окраски (тональности) коротких неформальных текстов на русском языке. Ограничение на длину рассматриваемых текстов в данном случае носит весьма условный характер. Так как наиболее популярными коллекциями данных, предназначенными для разработки моделей анализа неформальных текстов, являются публичные сообщения из социальных сетей (например, Twitter) и отзывы пользователей интернет-сайтов о различных товарах, услугах, рецензии фильмов, которые обычно содержат до 25 слов, удобно рассматривать именно эту длину сообщения в качестве ограничения.

Часто задача определения эмоциональной составляющей текста решается с помощью привлечения дополнительных данных, не являющихся текстовыми: аудиосигнала, соответствующего произнесенной фразе, видеозаписи говорящего, приложенного к тексту изображения, эмодиконов, знаков пунктуации и т.д. [1]. Тем не менее, наиболее интересной

кажется задача обработки исключительно текстовой информации, очищенной от знаков пунктуации, эмодиконов и других дополнительных данных, которые часто используются при решении подобных задач (например, для отзывов о кинофильмах, такими данными могут быть оценки фильмов по пятибалльной шкале).

Существует несколько различных вариантов постановки задачи анализа эмоциональной окраски. Среди постановок задачи, в которых пространство меток одномерно, можно выделить анализ полярности (классификация текстов на положительно и отрицательно эмоционально окрашенные) и анализ градации эмотивной составляющей в одномерном пространстве меток (определение не только полярности эмоциональной окраски, но и ее силы: от крайне негативной окраски до крайне позитивной). Более сложной задачей является определение эмоциональной окраски как принадлежности текста к одному из нескольких различных классов с метками «радость», «грусть», «отвращение», «страх» и другими типами меток-эмоций.

Существует два основных типа методов обработки текстов на естественном языке: методы, основанные на правилах (англ. rule-based), и методы, использующие машинное обучение.

Методы определения тональности текста, основанные на правилах, используют словари эмотивной лексики, которые включают в себя слова и фразы, наличие которых в тексте может указывать на определенную эмоциональную окраску. Такие алгоритмы не требуют большого количества данных, достаточно выделить присущий рассматриваемой предметной области эмотивный лексикон. Такие подходы полностью детерминированы, результаты работы этих алгоритмов интерпретируемы в терминах рассматриваемой предметной области. Недостатком подобных методов является невозможность масштабировать полученные модели: в разных контекстах и при рассмотрении различных предметных областей одни и те же слова могут нести различную эмоциональную окраску.

Алгоритмы машинного обучения основаны на использовании статистических данных. Для их использования требуется большое количество размеченных данных, которые часто бывает трудно получить. Тем не менее, преимущество этих алгоритмов заключается в возможности построения более общих моделей с помощью использования техник обучения с переносом (англ. transfer learning).

Среди алгоритмов машинного обучения особое место занимают искусственные нейронные сети. Они показывают передовые результаты в задачах анализа текстов большого объема и в некоторых случаях дают результаты распознавания, сравнимые по качеству с анализом, проведенным человеком [2]. В частности, высокие результаты при решении задач автоматического анализа текстов показывают сверточные (CNN) и рекуррентные (RNN) нейронные сети. Различные модификации рекуррентных нейронных сетей, такие как LSTM и GRU, помогают решить проблему затухания градиента при использовании алгоритма обратного распространения ошибки для обучения модели. Подобные методы позволяют контролировать количество информации, которая будет использоваться на каждом шаге обработки входной последовательности слов, регулируя с помощью весовых коэффициентов вклад разных частей входной последовательности в значение целевой функции.

Механизм внимания, используемый при анализе текстов, помогает выделить фрагменты входного текста, которые должны оказать наибольшее влияние на значение оптимизируемой функции. Хотя нейросетевые архитектуры с использованием этого механизма дают значительный прирост качества распознавания при анализе длинных текстов, для текстов небольшой длины разница между моделями, использующими и не использующими механизм внимания, может не быть существенной [3].

В неформальных текстах могут встречаться опечатки, ошибки и неточности. Существуют алгоритмы получения векторного представления текстовых данных, которые могут использоваться для обработки зашумленных текстов [4]. Подобные алгоритмы основаны на нейросетевых подходах и позволяют учитывать не только само слово, но и

контекст, в котором оно употребляется. Еще одним преимуществом такого подхода является то, что при его использовании размерность полученных векторных представлений имеет порядок мощности алфавита. Представление данных в пространстве невысокой размерности позволяет ускорить обработку больших объемов данных.

Таким образом, для решения задачи автоматического определения эмоциональной окраски текста, целесообразно использовать подходы на основе искусственных нейронных сетей. При работе с короткими текстами использование механизма внимания, вероятно, не улучшит результаты распознавания. Если в качестве набора данных для обучения модели выбраны неформальные тексты, можно решить проблему зашумленности данных уже на этапе получения векторных представлений.

Литература

1. Liu B., Zhang L. A survey of opinion mining and sentiment analysis // Mining Text. Springer U.S. – 2012. – P. 415–463.
2. Arkhipenko K. et al. Comparison of Neural Network Architectures for Sentiment Analysis of Russian Tweets // Proceedings of the International Conference «Dialogue 2016». – 2016. – V. 15. – P. 50–58.
3. Cliche M. BB_twtr at SemEval-2017 Task 4: Twitter Sentiment Analysis with CNNs and LSTMs // Proceedings of the 11th International Workshop on Semantic Evaluation (SemEval-2017). – 2017. – P. 573–580.
4. Malykh V. Robust word vectors for Russian language [Электронный ресурс]. – Режим доступа: <http://val.maly.hk/pdf/ainl2016.pdf> (дата обращения: 06.01.2019).



Зено Бассель

Год рождения: 1986

Университет ИТМО, факультет информационных технологий
и программирования, аспирант

Направление подготовки: 09.06.01 – Информатика и вычислительная
техника

e-mail: basilzeno@gmail.ru



Матвеев Юрий Николаевич

Год рождения: 1955

Университет ИТМО, факультет информационных технологий
и программирования, д.т.н.

e-mail: matveev@speechpro.com

УДК 004.89

ОБУЧЕНИЕ РАСПУТАННОГО ПРЕДСТАВЛЕНИЯ С ИСПОЛЬЗОВАНИЕМ ГЕНЕРАТИВНЫХ СОСТЯЗАТЕЛЬНЫХ СЕТЕЙ

Зено Б.

Научный руководитель – д.т.н. Матвеев Ю.Н.

В работе предложена комплексная среда для изучения парных кодов распутанной идентичности и представления для набора данных. Разработанная структура состоит из одного генератора, двух кодеров и одного дискриминатора.

Ключевые слова: генеративные состязательные сети, сверточные слои, распутанное представление, кодер, декодер.

За последние несколько лет производительность систем распознавания лиц значительно улучшилась вследствие использования сверточных нейронных сетей, обученных с крупномасштабными базами данных лиц [1, 2]. Тем не менее, вариации позы по-прежнему являются областью исследования для многих реальных сценариев распознавания лиц.

Существующие методы, которые рассматривают варианты поз, можно разделить на две категории. Одна категория пытается подготовить ручную поза-инвариантные признаки [1, 2], в то время как другая применяет метод синтеза для генерации изображений лица конкретного человека. Например, TP-GAN [3] и FF-GAN [4] пытаются восстановить изображение фронтального вида из любого изображения лица с большой позой. DR-GAN [5] может изменить позу входного изображения лица. Однако эти методы могут манипулировать только ограниченными позами изображения лица. Кроме того, они требуют полной аннотации атрибутов для обучения моделей. Некоторые из них, основанные на GAN-методах [4–6], обычно имеют схему с одним путем (сеть кодера-декодера сопровождается сетью дискриминатора). В то время как другие [3] имеют двухсторонний дизайн. Например, кодер (E) отображает входные изображения в скрытое пространство (Z), а затем подается в декодер (G) с вектором позы для генерации новых видов.

На практике, чтобы обучать значительный набор обучающих данных без дополнительных ограничений, были предложены различные структуры GAN. Они изучают интерпретируемые и значимые скрытые представления в неконтролируемой среде, такой как InfoGAN, BiGAN, или в контролируемой среде. Несмотря на все усилия в этой области, эти

подходы не учитывают один из основных принципов генерации изображения лица, а именно распутывание идентичности и позы лица.

Обучение распутанных кодов идентификации и позы является сложной задачей, поскольку:

1. компьютеры должны «представлять», как будет выглядеть объект после применения трехмерного вращения;
2. генерирование нескольких видов должно сохранять ту же «идентичность»;
3. разработанные алгоритмы решения этого класса задач обучают GAN в контролируемых условиях.

Несмотря на все усилия, предпринимаемые в этой области, по-прежнему отсутствуют согласованные рамки для обучения неконтролируемого распутанного представления позы и личности. Исходя из этого, в работе автор предложил комплексную структуру для обучения парных кодов распутанной идентичности и представления для набора данных.

Наша структура состоит из одного генератора, в отличие от [3]. Генератор имеет несколько сверточных слоев и набор остаточных блоков. Вход генератора представляет собой конкатенацию идентификационного кода и кода позы, в то время как требуемый код позы подается в многослойный перцептрон (MLP) для генерации параметров слоев AdaIn (Adaptive Instance Normalization), которые они вводили в каждом остаточном блоке генератора. Архитектура предложенного нами GAN-метода показана на рисунке.

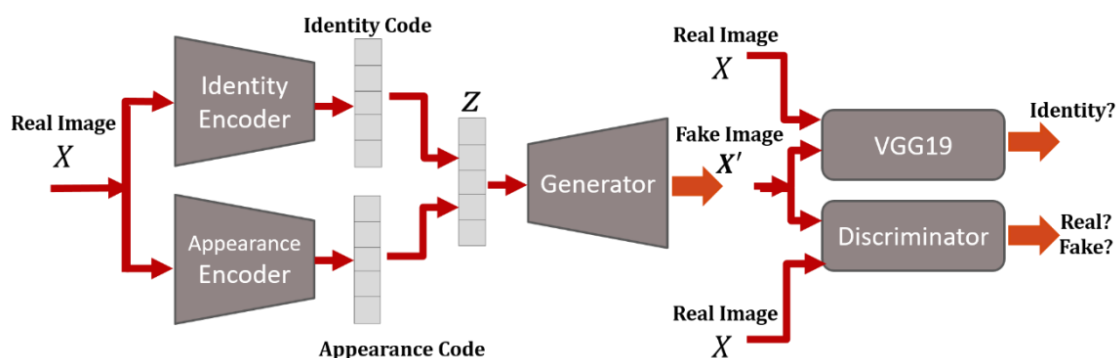


Рисунок. Архитектура модели GAN

Для проведения экспериментов используется набор данных Multi-PIE. Multi-PIE – это крупнейшая база данных с маркировкой для оценки распознавания лиц в условиях позы, освещения и выражения в контролируемых условиях. База содержит 337 тождеств и предоставляет аннотации вертикальных углов поворота позы головы, таких как 90°, 45°, 0°.

По итогам проведенных экспериментов предложенная модель генерирует новые изображения лиц с новыми позами, но изображения размыты немного, поскольку идентичность не сохраняет 100%. Для улучшения качества изображения необходимо обучить модель длительное время с большим количеством итераций.

Литература

1. Chen D., Cao X., Wen F. and Sun J. Blessing of dimensionality: High-dimensional feature and its efficient compression for face verification // Proceedings of the 2013 IEEE Conference on Computer Vision and Pattern Recognition. – 2013. – P. 3025–3032.
2. Schroff F., Kalenichenko D. and Philbin J. FaceNet: A unified embedding for face recognition and clustering [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1503.03832.pdf> (дата обращения: 06.01.2019).
3. Huang R., Zhang S., Li T. and He R. Beyond face rotation: Global and local perception GAN for photorealistic and identity preserving frontal view synthesis [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1704.04086.pdf> (дата обращения: 06.01.2019).

4. Yin X., Yu X., Sohn K., Liu X. and Chandraker M. Towards large-pose face frontalization in the wild [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1704.06244.pdf> (дата обращения: 06.01.2019).
5. Tran L., Yin X. and Liu X. Disentangled representation learning GAN for pose-invariant face recognition [Электронный ресурс]. – Режим доступа: https://www.researchgate.net/publication/316017161_Disentangled_Representation_Learning_GAN_for_Pose-Invariant_Face_Recognition (дата обращения: 06.01.2019).
6. Bao J., Chen D., Wen F., Li H., Hua G. Towards Open-Set Identity Preserving Face Synthesis [Электронный ресурс]. – Режим доступа: http://openaccess.thecvf.com/content_cvpr_2018/papers/Bao_Towards_Open-Set_Identity_CVPR_2018_paper.pdf (дата обращения: 06.01.2019).

**Казиева Назым**

Год рождения: 1976

Университет ИТМО, факультет информационных технологий
и программирования, аспирантНаправление подготовки: 09.06.01 – Информатика и вычислительная
техника

e-mail: kaznaz@list.ru

**Кухарев Георгий Александрович**

Год рождения: 1941

Университет ИТМО, факультет информационных технологий
и программирования, д.т.н., профессор

e-mail: gakukharev@etu.ru

УДК 612.087.1**ФУНКЦИОНАЛЬНЫЕ СХЕМЫ АЛГОРИТМОВ ШТРИХОВОГО КОДИРОВАНИЯ
ДЛЯ ЗАДАЧ ЛИЦЕВОЙ БИОМЕТРИИ****Казиева Н.****Научный руководитель – д.т.н., профессор Кухарев Г.А.**

Работа выполнена в рамках темы НИР № 617040 «Синтез эмоциональной речи на основе глубокого машинного обучения».

В работе рассмотрено использование функциональных схем алгоритмов штрихового кодирования для задач лицевой биометрии. Представлены инструменты для моделирования задач использования штриховых кодов и моделирования процессов построения штриховых кодов.

Ключевые слова: штриховые коды, генерация QR-кодов, BIO QR-коды, лицевая биометрия, генерация цветных штриховых кодов.

При выполнении прикладных задач, связанных с лицевой биометрии, применяются инструменты для моделирования задач использования штриховых кодов и моделирования процессов построения штриховых кодов. В выполненном исследовании в роли такого инструмента предлагаются различные программные модули, которые реализованы в пакете MATLAB. Алгоритмы, реализованные в них, представлены в виде функциональных схем. Такое представление является наиболее наглядным для представления моделируемых процессов.

В работе [1] выполнен обзор существующих решений в биометрии, также были предложены новые решения применения штрихового кодирования (ШК) в лицевой биометрии. Исходя из изученного, новые решения ШК для лицевой биометрии и задач, связанных с ней, могут быть выполнены путем применения цветных ШК. В процессе исследования разработаны базовые алгоритмы применения ШК в приложениях лицевой биометрии. К базовым алгоритмам относятся: генерация цветных линейных и двумерных ШК; генерация QR-кодов «ИНФО» и «АНТРО»; способы их компоновки в структуре ШК с изображением лица (ИЛ) или слоями ИЛ; алгоритмы онлайн и офлайн генерации ШК и варианты их реализации; структура сообщений для различных вариантов реализации ШК; алгоритмы кодирования и декодирования сообщений ШК.

Исходными данными в разрабатываемых ШК являются: изображение лица, которое может содержать следующую информацию: ИЛ цветное, ИЛ черно-белое, ИЛ анфас, ИЛ

профиль; размер ИЛ; фенотип лица; морфотип лица; координаты антропометрических точек (АПТ), также могут содержать «карту глубин» и краткую документальную информацию (о человеке или базе, которому принадлежит лицо).

Необходимо отметить, что наиболее важным и информативным способом представления ИЛ достигается в рамках цифровой лицевой антропометрии – через представление его набором ключевых точек, а также расстояний между ними и соотношениями расстояний [2]. Цифровая лицевая антропометрия включает: автоматическое определение координат ключевых точек (антропометрических точек) на ИЛ; оценку всех базовых (габаритных) размеров лица и его частей по расположению координат АПТ; оценку координат АПТ и границ примитивов лица; вычисление соотношений между выбранными координатами АПТ; составление сводных таблиц по этим соотношениям; проведение специальных статистических (сравнительных) исследований по ним и поиск закономерностей их изменений (например, для возрастных, половых и расовых групп людей и их популяций). При этом именно координаты АПТ на лице каждого человека отображают его индивидуальность, представленную на «обобщенной (согласованной) координатной сетке АПТ», в рамках которой производятся все морфометрические измерения для дальнейшего их использования [3].

Исходя из изложенного, далее в работе представлены соответствующие функциональные схемы алгоритмов построения, соответствующих ШК. На рис. 1 показана общая функциональная схема генерации QR-кода с АПТ для ИЛ (АНТРО). Кроме того, возможна генерация документальной информации в отдельный QR-код (ИНФО). Источником исходных сцен могут быть офлайн или онлайн сцены, в работе разработаны несколько сценариев обработки ИЛ и получение АПТ. Следует отметить, что АНТРО может содержать информацию о фенотипе лица, морфотипе лица, размере ИЛ, более детально описанный в работе [4].

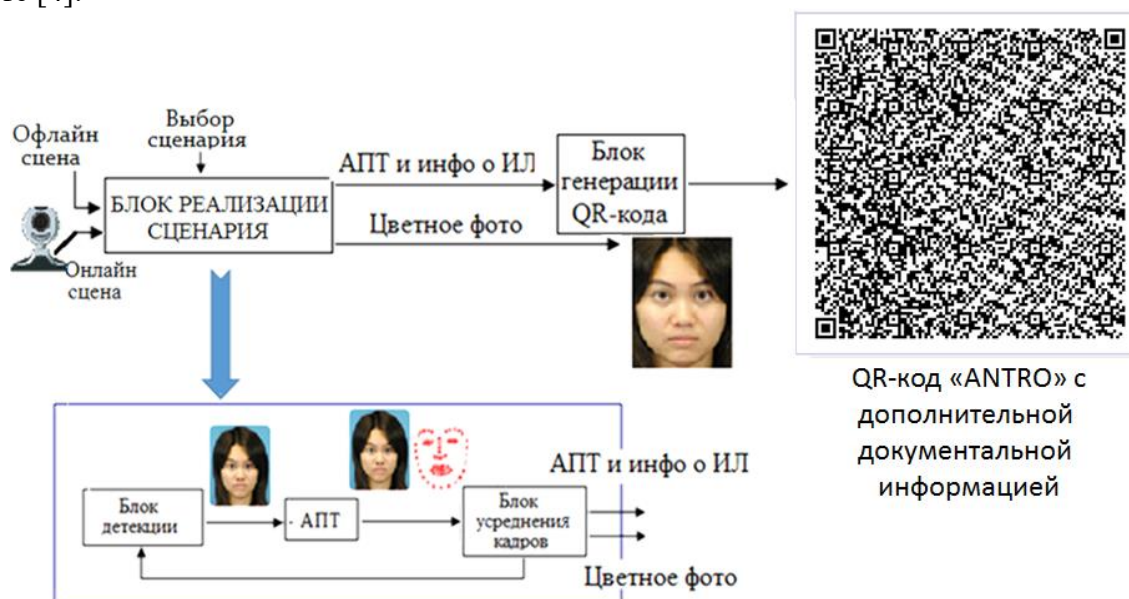


Рис. 1 Функциональная схема построения QR-кода с АПТ для ИЛ

Выходные данные, полученные в результате построения QR-кода (АНТРО, ИНФО), обработанное ИЛ, собираются в контейнер, который называется BIO QR-код. В рамках исследований авторами предложены новые решения по построению BIO ШК (линейные, двумерные) для передачи информации.

Таким образом, BIO ШК как контейнеры, должны содержать всю перечисленную выше информацию. Примеры построения BIO QR-кодов показаны на рис. 2. На рис. 2, а, показан BIO QR-код с ИЛ в формате GRAY. На рис. 2, б, представлен BIO QR-код с цветными ИЛ анфас и профиль в виде цветного изображения в формате «*.png».

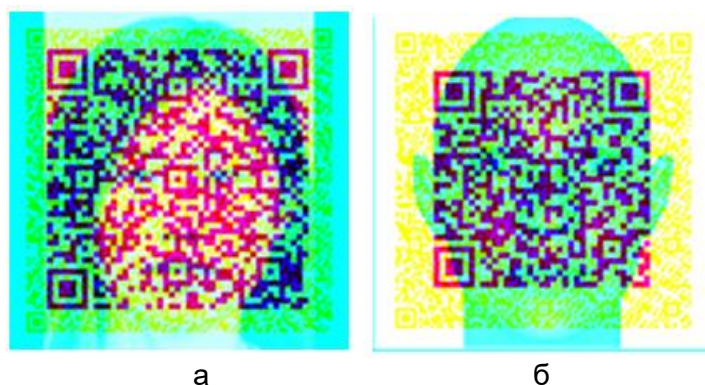


Рис. 2. Новые решения по построению BIO QR-кодов

Функциональная схема декомпозиции сгенерированных BIO QR-кодов показана на рис. 3. Аналогичная функциональная схема применяется для BIO ШК (линейных).



Рис. 3. Функциональная схема контроля записи и расшифровки QR-кодов

В процессе исследования получены свидетельство на «Программу для ЭВМ» алгоритмов, созданных в рамках вышеизложенных функциональных схем алгоритмов штрихового кодирования для задач лицевой биометрии. Зарегистрированными алгоритмами являются:

1. программа онлайн-генерации QR-кода с координатами АПТ ИЛ. В программе реализуются: онлайн-детекция ИЛ в последовательности кадров; определение в каждом кадре координат АПТ лица и их усреднение по нескольким кадрам; запись усредненных значений координат в QR-коде; запись QR-кода в памяти компьютера; контрольное чтение QR-кода из памяти и декодирование АПТ из него. В QR-код включается INFO QR-код;
2. программа детекции лиц фанатов из видеосцен и создания базы изображений лиц фанатов: выполняет анализ потока видеосцен с мест скопления футбольных фанатов (на входах/выходах и трибунах стадионов, в фан-зонах, на улицах). Цель анализа – детекция лиц, окрашенных в цвета национальных флагов или укрытых цветными полосами и масками. Исходные сцены, найденные лица и их фенотипы, сохраняются в базе изображений лиц фанатов (БИЛФ);
3. программа распознавания лиц фанатов в видеосценах. В программе выполняется: детекция лиц, окрашенных в цвета национальных флагов или укрытых цветными полосами и масками; найденные лица проверяются на соответствие пороговому значению по размеру и фенотипу; прошедшие эту проверку лица сравниваются с лицами из БИЛФ. Поиск и подбор лиц из БИЛФ осуществляется по критерию минимального расстояния, в метрике L2 и с рангом 1–10.

Следует отметить, что применения функциональных схем алгоритмов штрихового кодирования для задач лицевой биометрии более эффективно показывает процесс создания, компоновки новых решений BIO ШК (линейных и двумерных).

Литература

1. Kaziyeva N., Kukharev G.A., Matveev Y.N. Barcoding in biometrics and its development // Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics). – 2018. – V. 11114. – P. 464–471.
2. Jayaratne Yasas S.N., Zwahlen Roger A. Application of Digital Anthropometry for Craniofacial Assessment // Craniomaxillofacial Trauma and Reconstruction. – 2014. – V. 7. – № 2/2014. – P. 101–107.
3. Deutsch C.K. et al. The Farkas System of Craniofacial Anthropometry: Methodology and Normative Databases // Handbook of Anthropometry. – 2012. – P. 561–573.
4. Кухарев Г.А., Казиева Н., Цымбал Д.А. Технологии штрихового кодирования для задач лицевой биометрии: современное состояние и новые решения // Научно-технический вестник информационных технологий, механики и оптики. – 2018. – Т. 18. – № 1. – С. 72–86.

**Лизунова Инна Александровна**

Год рождения: 1996

Университет ИТМО, факультет информационных технологий
и программирования, студент группы № М41212Направление подготовки: 09.04.02 – Информационные системы
и технологии

e-mail: inna.lizunova@gmail.com

**Махныткина Олеся Владимировна**

Год рождения: 1982

Университет ИТМО, факультет информационных технологий
и программирования, к.т.н., доцент

e-mail: makhnytkina@corp.ifmo.ru

УДК 004.9**АНАЛИЗ ТОНАЛЬНОСТИ ТЕКСТОВ С ИСПОЛЬЗОВАНИЕМ НЕЙРОСЕТЕВЫХ
МОДЕЛЕЙ****Лизунова И.А.****Научный руководитель – к.т.н., доцент Махныткина О.В.**

Работа выполнена в рамках темы НИР № 618278 «Синтез эмоциональной речи на основе генеративных состязательных сетей».

Целью работы являлось исследование и проведение сравнительного анализа различных методов анализа тональности текстов. В работе рассмотрено использование нейросетевых моделей при построении векторного представления слов и последующей классификации текстов на основе их эмоциональной окраски.

Ключевые слова: анализ тональности, нейросеть, машинное обучение, наивный байесовский классификатор, логистическая регрессия, метод опорных векторов, сверточная нейронная сеть.

Анализ тональности текстов – активно развиваемое направление в сфере автоматической обработки текстов, которое занимается изучением эмоциональной окраски текстовых документов и изложенных в них мнений. Актуальность данной темы обусловлена развитием и ростом популярности социальных сетей, онлайн-овых рекомендательных сервисов, новостных порталов и блогов, которые содержат большое количество мнений людей по различным вопросам [1]. Перспективным направлением в этой области является использование нейросетевых моделей.

В работе осуществлено сравнение различных методов классификации текстов по их эмоциональной окраске, реализованных на языке программирования Python. Для обработки был выбран корпус коротких текстов Юлии Рубцовой [2], сформированный на основе русскоязычных сообщений из социальной сети Twitter. Данный корпус содержит 114 991 позитивных сообщений и 111 923 негативных.

Работа с русскоязычными сообщениями из сети имеет ряд особенностей. В таких сообщениях содержится множество упоминаний имен других пользователей вида «@имя пользователя» и url-ссылок. Поэтому в ходе предварительной обработки текста были проведены следующие операции: замена всех url-ссылок на токен «URL», замена всех упоминаний других пользователей на токен «USER», замена символа «ё» на символ «е»,

Далее было получено векторное представление слов с помощью технологии word2vec [3]. Данный инструмент основан на обучении нейронной сети. С его помощью каждому слову из входного корпуса был сопоставлен вектор на основе контекстной близости слов, т.е. слова, употребляемые в схожем контексте, а значит, близкие по смыслу, имеют близкие координаты. На рис. 1 представлен фрагмент визуализации полученного векторного представления, отображенного в двумерное пространство, на примере 20 наиболее близких слов для «iphone», «футболка», «аэропорт» и «фильм».

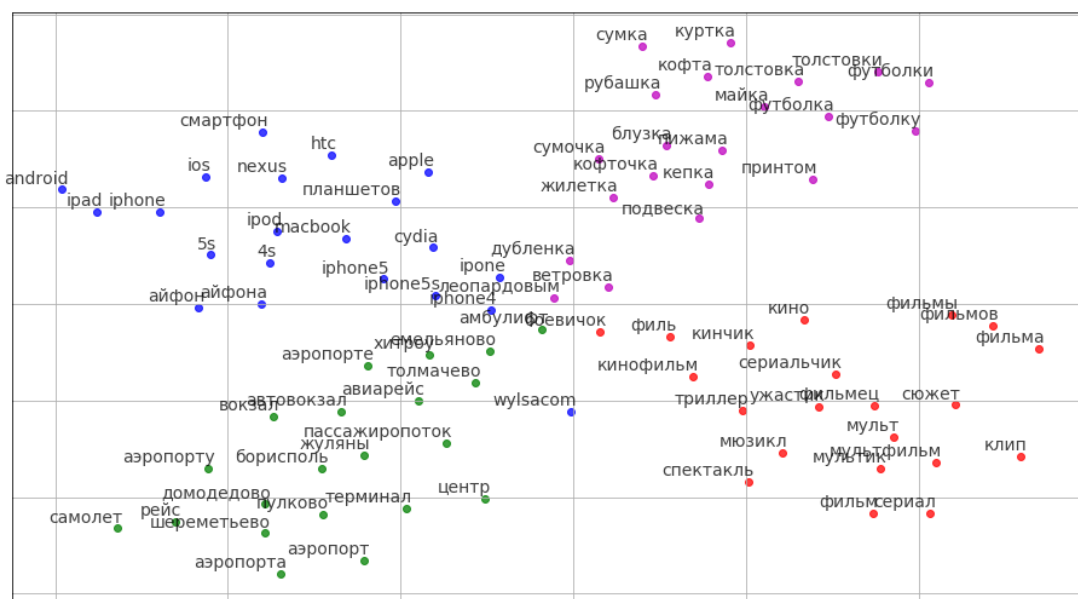


Рис. 1. Визуализация кластеров похожих слов

Для решения задачи классификации были рассмотрены следующие методы: наивный байесовский классификатор (NB), логистическая регрессия (LR), метод опорных векторов (SVM) [4] и сверточная нейронная сеть (CNN) [5, 6].

В качестве оценок эффективности работы алгоритмов используются следующие параметры: точность (precision), полнота (recall), F-мера (f1-score). Оценивание производится на основе следующих показателей: TP – истинно-положительное решение, TN – истинно-отрицательное решение, FP – ложноположительное решение, FN – ложноотрицательное решение, и осуществляется по формулам:

$$Precision = \frac{TP}{TP + FP},$$

$$Recall = \frac{TP}{TP + FN}.$$

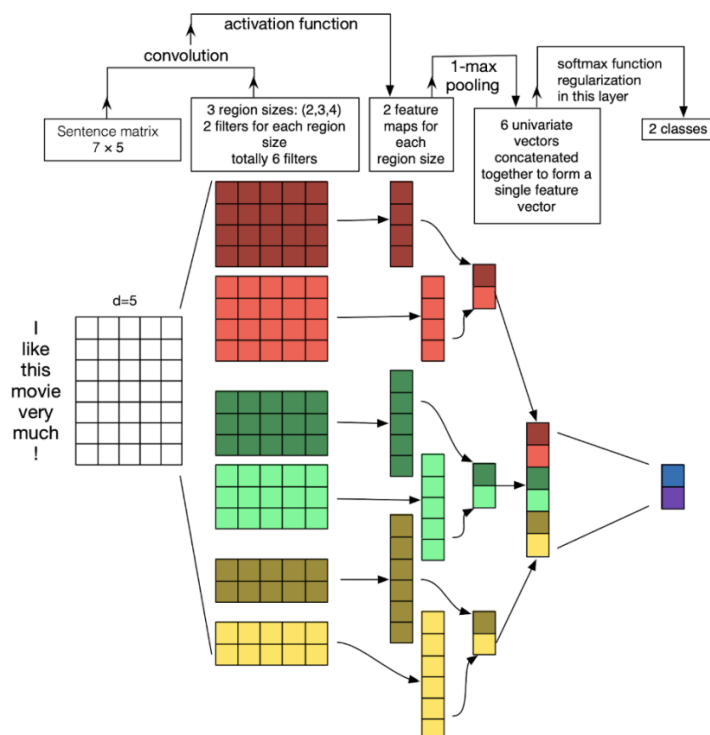


Рис. 2. Архитектура сверточной нейронной сети

F-мера является оценкой, основанной на результатах оценки точности и полноты, и вычисляется по следующей формуле

$$F_1 = \frac{2 \cdot \text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}.$$

Оценки, полученные в результате проведения классификации с помощью различных методов, представлены в таблице.

Таблица. Оценки работы различных классификаторов

Method	Precision	Recall	F1-score
NB	0,6245	0,6242	0,6238
LR	0,7278	0,7278	0,7278
SVM	0,7573	0,7572	0,7572
CNN	0,7798	0,7796	0,7795

Таким образом, наилучшую точность классификации показала сверточная нейронная сеть. Полученная в результате максимальная точность классификации составила 77,95%. В перспективе планируется увеличить точность классификации путем расширения диапазона рассматриваемых методов анализа тональности сообщений, включая использование других нейросетевых моделей, а также с помощью различных комбинаций методов предварительной обработки данных.

Литература

1. Мальцева А.В., Махныткина О.В., Шилкина Н.Е. Изучение поведенческих паттернов пользователей социальных сетей: возможности Big Data // Жизнь исследования после исследования: как сделать результаты понятными и полезными. – 2016. – С. 988–991.
2. Рубцова Ю.В. Построение корпуса текстов для настройки тонового классификатора // Программные продукты и системы. – 2015. – № 1(109). – С. 72–78.
3. Mikolov T., Chen K., Corrado G., Dean J. Efficient Estimation of Word Representations in Vector Space [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1301.3781.pdf> (дата обращения: 06.01.2019).

4. Воронцов К.В. Лекции по линейным алгоритмам классификации [Электронный ресурс]. – Режим доступа: <http://www.machinelearning.ru/wiki/images/6/68/voron-ML-Lin.pdf> (дата обращения: 06.01.2019).
5. Cliche M. BB_twtr at SemEval-2017 Task 4: Twitter Sentiment Analysis with CNNs and LSTMs // Proceedings of the 11th International Workshop on Semantic Evaluation (SemEval-2017). – 2017. – P. 573–580.
6. Zhang Y., Wallace B. A Sensitivity Analysis of (and Practitioners' Guide to) Convolutional Neural Networks for Sentence Classification [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1510.03820.pdf> (дата обращения: 06.01.2019).

**Макаров Ростислав Николаевич**

Год рождения: 1995

Университет ИТМО, факультет информационных технологий и программирования, студент группы № М4221

Направление подготовки: 09.04.02 – Информационные системы и технологии

e-mail: makarov-@outlook.com

**Рыбин Сергей Витальевич**

Год рождения: 1959

Университет ИТМО, факультет информационных технологий и программирования, к.ф.-м.н., доцент

e-mail: rybin@speechpro.com

УДК 004.934.5

**ПРЕДСКАЗАНИЕ ГРАММАТИЧЕСКИХ И ПРОСОДИЧЕСКИХ ХАРАКТЕРИСТИК
ТЕКСТА ДЛЯ СИСТЕМ TTS****Макаров Р.Н.****Научный руководитель – к.ф.-м.н., доцент Рыбин С.В.**

Работа выполнена в рамках темы НИР № 618278 «Синтез эмоциональной речи на основе генеративных состязательных сетей».

В работе представлены результаты дополнения признакового пространства грамматическими признаками и просодическими признаками для систем TTS. Были собраны базы и обучены модели на основе рекуррентных нейронных сетей.

Ключевые слова: машинное обучение, нейронные сети, синтез речи, тональный контур, обработка естественного языка.

Цель работы – построение модели тонального контура с помощью глубокого обучения. На данном этапе работы стояла задача дополнить пространство признаков для предсказания акустических характеристик грамматическими признаками текста и просодическими характеристиками звука.

Сначала был проведен поиск баз данных, которые можно использовать для этой задачи. Был найден национальный корпус русского языка [1]. Интерес представляет офлайн версия основного корпуса русского языка со снятой морфологической омонимией. Объем словоупотреблений около 1 млн. База данных представлена множеством документов формата «html» с различными текстами. Каждый текст разделен по предложениям, внутри их по словам, и каждому слову дана полная грамматическая характеристика (часть речи, род, число, форма и т.д.). В качестве целевых признаков были извлечены такие как: числительное, глагол, существительное, вводное слово, частица, союз, местоимение-прилагательное, местоимение-существительное, прилагательное, местоимение-наречие, предикат, предлог, междометие, прилагательное-числительное, инициалы, иностранные слова, местоимение-предикат, пунктуация.

После извлечения необходимой информации из базы была построена нейронная сеть для классификации слов по частям речи. Архитектура сети представлена на рисунке. На вход сети подается предложение русского языка, которое затем переводится в векторное пространство,

размерностью 200 значений, embedding слоем. После чего пропускается через два рекуррентных слоя LSTM [2] (выходной размер по 256 значений). И затем сжимается полносвязным слоем, формируя выход сети.

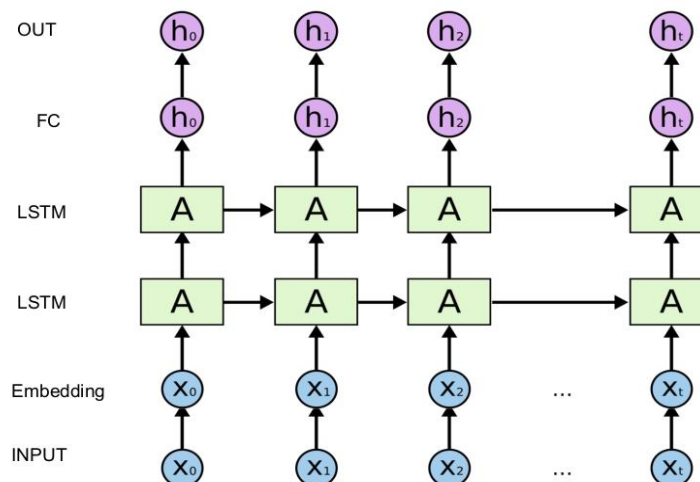


Рисунок. Архитектура нейронной сети

Также было взято 5 текстов книг, общим объемом словоупотребления 870 тысяч, из библиотеки Максима Машкова [3]. Из которых были извлечены признаки для частеречной (морфологической) разметки и разметки интонационных контуров с помощью программы VVTools, предоставленной базовым предприятием Центром речевых технологий (ЦРТ) [4].

После некоторых объединений целевые признаки в задаче частеречной разметки получились следующими: числительное, существительное, прилагательное, глагол, причастие, «странное» слово, наречие, местоимение, пунктуация.

Интонационные контура разделены на следующие категории:

1. завершенность: нисходящая интонация (ИК-11);
2. специальный вопрос: фразовое ударение на вопросительном слове (ИК-30);
3. восклицание без вопросительного слова (восходящая интонация) (ИК-40);
4. восклицание с вопросительным словом (нисходящая интонация) (ИК-50);
5. простой общий вопрос: (восходящая интонация) (ИК-70);
6. незавершенность, вопрос: восходящая интонация (ИК-110);
7. отсутствие тонального контура.

Для последних двух задач была использована та же архитектура нейронной сети.

Анализ полученных результатов. Все результаты были оценены общей точностью по всей базе (что дает довольно общую и неполную оценку, в связи с разным количеством слов каждого класса в базе). Более точно описывающая метрика, использованная автором – F1-мера для каждого класса.

Общая точность предсказания частей речи на базе национального корпуса русского языка составила 96%. F1-мера 13/19 классов составила более 90%, наихудшая F1-мера – у иностранных слов, 65%.

Общая точность предсказания частей речи на базе VVTools составила 96%, F1-мера 7/10 классов составила более 90%, наихудшая F1-мера – у причастия – 35%.

Общая точность предсказания просодических признаков составила 98,5%, F1-мера 4/7 превысила 90%, наихудшая F1-мера получилась у ИК-30 и составила 82%.

Выводы

1. Извлечены признаки из национального корпуса русского языка и построена модель частеречной разметки.

2. Извлечены признаки с помощью системы VVTools и построены модели предсказания грамматических и просодических характеристик.
3. Обучены модели и оценены результаты.

Литература

1. Национальный корпус русского языка [Электронный ресурс]. – Режим доступа: <http://www.ruscorpora.ru> (дата обращения: 16.03.2019).
2. Hochreiter S., Schmidhuber J. Long short-term memory // Neural computation. – 1997. – V. 9. – № 8. – P. 1735–1780.
3. Библиотека Максима Мошкова [Электронный ресурс]. – Режим доступа: <http://www.lib.ru> (дата обращения: 16.03.2019).
4. Чистиков П.Г., Хомицевич О.Г., Рыбин С.В. Статистические методы автоматического определения мест и длительности пауз в системах синтеза речи // Изв. вузов. Приборостроение. – 2014. – Т. 57. – № 2. – С. 28–32.



Мамаев Никита Константинович

Год рождения: 1996

Университет ИТМО, факультет информационных технологий
и программирования, студент группы № М41211

Направление подготовки: 09.04.02 – Информационные системы
и технологии

e-mail: nikita.mamaev1@gmail.com



Черных Ирина Александровна

Год рождения: 1981

Университет ИТМО, факультет информационных технологий
и программирования, инженер

e-mail: chernykh-i@speechpro.com

УДК 004.89

АНАЛИЗ ФУНКЦИОНАЛЬНЫХ ВОЗМОЖНОСТЕЙ ДИАЛОГОВОЙ ПРОГРАММНОЙ БИБЛИОТЕКИ DEERPAVLOV

Мамаев Н.К.

Научный руководитель – Черных И.А.

Работа выполнена в рамках темы НИР № 618278 «Синтез эмоциональной речи на основе генеративных состязательных сетей».

Целью работы являлся анализ качества работы компонентов автоматической диалоговой системы, представленных в программной библиотеке DeepPavlov: двух компонентов, выполняющих автоматическое исправление ошибок, а также компонента, выполняющего ранжирование документов.

Ключевые слова: диалоговая система, вопросно-ответная система, чатбот, проверка правописания, исправление ошибок, ранжирование документов.

Диалоговые системы имеют широкий спектр применения, начинающийся со служб поддержки и заканчивающийся средствами обучения языку и личными помощниками. Разработка и ввод в эксплуатацию подобных систем позволяет автоматизировать многие задачи, подразумевающие работу с естественным языком. Долгое время основными технологиями, применяемыми в обработке естественного языка (Natural Language Processing, NLP), было машинное обучение, а именно линейные модели (метод опорных векторов, логистическая регрессия), обучаемые на векторах очень высоких размерностей, но очень разреженных, так как для векторного представления слов использовались подходы one-hot и bag-of-words. На сегодняшний день сфере NLP характерен переход как к новым моделям векторизации слов и текстов – GloVe, word2vec, fastText, – так и к новым вычислительным архитектурам, среди которых распространенными являются сверточные, рекуррентные и генеративно-состязательные сети, а также их модификации.

При построении диалоговой системы необходимо учитывать, что в качестве входных данных зачастую будут подаваться зашумленные текстовые фрагменты, возможно, содержащие орфографические, грамматические, синтаксические и лексико-семантические ошибки. Это чревато возникновением ошибок в работе диалоговой системы: в частности, возникновением ложных связей при обучении.

В библиотеке DeepPavlov представлены два компонента, условно названные levenshtein и brillmoore, разработанные для автоматического исправления орфографических ошибок. Оба компонента имеют архитектуру pipeline, и в обоих используется n -граммная статистическая модель языка, которая при анализе текста позволяет оценить вероятность появления той или иной последовательности слов. Алгоритм, использующий статистическую модель, является вторым и последним звеном в архитектуре обоих компонентов. Для ускорения работы со статистической моделью используется модуль KenLM, аппроксимирующий данную модель вариацией сглаживающего метода Kneser-Ney [1].

В частности, компонент levenshtein определяет кандидатов на исправление на основе расстояния редактирования Damerau-Levenshtein (D.-L.). Расстояние D.-L. между двумя словами определяется как наименьшее количество операций вставки, удаления, замены и перемещения соседних символов, необходимое для приведения одного слова к форме другого. Принцип поиска кандидатов на исправление можно обобщить следующим образом: если слово, подлежащее анализу, отличается от «ближайшего» по расстоянию D.-L. словарного слова не менее чем на 1 правку, но не более чем на N правок, это слово становится кандидатом на замену.

Компонент brillmoore определяет кандидатов на исправление на основе noisy channel-подхода, а точнее его версии, предложенной Е. Brill и R.C. Moore [2]. Реализованный в компоненте алгоритм предусматривает обучение на паре слов, одно из которых содержит ошибку, а другое является орфографически корректным. В отличие от предыдущего компонента, эта модель позволяет учитывать ошибки произвольного вида вместе с ошибкой, сохраняя информацию о ее позиции в слове.

Наконец, статистическая модель языка вычисляет показатель уверенности для исправленного варианта последовательности слов, после чего он перемножается с показателем от предыдущего звена в архитектуре, и на основе наивысшего произведения показателей выбирается окончательный кандидат на замену.

Авторы произвели сравнительное тестирование работы компонентов levenshtein и brillmoore относительно работы программного обеспечения (ПО) LanguageTool 3.3 (www.languagetool.org) на данных, предоставленных авторами ООО «ЦРТ». Данные имели вид набора диалогов между клиентом и сотрудником службы поддержки. Данные, прошедшие обработку с помощью компонентов из библиотеки DeepPavlov, были сравнены с данными, ошибки в которых были исправлены вручную. Полученные показатели:

1. компонент brillmoore: точность: 0,2318, полнота: 0,4904, F-мера: 0,3148;
2. компонент levenshtein: точность: 0,2318, полнота: 0,4904, F-мера: 0,3148;
3. ПО LanguageTool 3.3: точность: 0,2111, полнота: 0,5481, F-мера: 0,3048.

Результаты работы компонентов достаточно близки. Тем не менее, компонент brillmoore предполагает обучение, что гипотетически позволяет получать более высокие результаты при обучении и тестировании на данных со специфической лексикой и ошибками.

В качестве одного из подходов к созданию целеориентированной диалоговой системы, или, по крайней мере, как один из ее элементов, может рассматриваться вопросно-ответная система, вариант реализации которой представлен в библиотеке DeepPavlov в виде компонента под названием «Question Answering Model for SQuAD Dataset» [3]. Вторая часть названия компонента отсылает к известной в сфере data mining задаче построения вопросно-ответной системы на основе набора данных SQuAD (Стэнфордского вопросно-ответного корпуса). Формулировка данной задачи звучит так: «по данному текстовому отрывку и вопросу алгоритм должен определить позицию начала и конца ответа в отрывке, заведомо содержащем ответ. В принципе работы компонента лежит нейронная сеть с архитектурой R-Net [4].

Естественной надстройкой над моделью вопросно-ответной системы можно назвать модель вопросно-ответной системы на открытой базе знаний. Условие новой задачи подразумевает, что поиск ответа будет производиться по некоторому условно открытому множеству документов, возможно содержащих информацию, не принадлежащую лишь к

одной сфере знаний. Компонент (или Навык/Skill, согласно терминологии разработчиков), реализующий описанную надстройку, входит в библиотеку DeepPavlov и носит название «Open Domain Question Answering Skill on Wikipedia» [5]. Алгоритм, позволяющий использовать предыдущую архитектуру для решения более широкой задачи – это алгоритм ранжирования набора документов по текстовому запросу на основе показателя TF-IDF. Иными словами, этот компонент по некоторому входному тексту упорядочивает документы из множества согласно их близости к тексту по показателю TF-IDF.

Навык «ODQA on Wikipedia» имеет модульную структуру, и компонент, производящий ранжирование, может функционировать отдельно от компонента «QA Model for SQuAD Dataset». Поскольку у нас не было возможности произвести тестирование Навыка, мы предприняли тестирование компонента, производящего ранжирование, в качестве подготовки к предыдущему.

Итак, мы произвели тестирование TF-IDF ранжирующего компонента на данных, предоставленных нам ООО «ЦРТ». Данные включили набор документов типа «вопрос-ответ» на формальном языке (часто задаваемые вопросы, или FAQ, с сайта одной из компаний-поставщика услуг населению), а также набор переформулировок на естественном языке, составленных лингвистами вручную для некоторых вопросов из FAQ. Тестирование проводилось три раза: на переформулировках, на исходных вопросах и на исходных документах. Для оценки качества работы компонента вычислялся показатель Exact Match (EM): количество верных результатов (первый сопоставленный документ является истинным связанным) делилось на количество вопросов, переданных в программу.

Полученные показатели:

- тестирование на переформулировках вопросов – EM: 0,2887, всего запросов: 665;
- тестирование на исходных вопросах – EM: 0,431, всего запросов: 85348;
- тестирование на исходных документах – EM: 0,606, всего запросов: 85348.

Полученные уровни точности представляются ожидаемыми для используемого набора данных. Заметим, что для этого набора характерно следующее: многие вопросы сформулированы очень кратко, формулировка ответов зачастую сильно отличается от формулировки соответствующих вопросов, в большинстве документов присутствует одинаковая лексика. Возможным путем повышения точности решения данной задачи представляется использование ранжировщика на основе моделей word2vec или doc2vec, позволяющих учитывать семантическую близость текстов.

Литература

1. DeepPavlov's documentation: Automatic spelling correction pipelines [Электронный ресурс]. – Режим доступа: http://docs.deeppavlov.ai/en/latest/components/spelling_correction.html (дата обращения: 19.01.2019).
2. Brill E., Moore R.C. An Improved Error Model for Noisy Channel Spelling Correction [Электронный ресурс]. – Режим доступа: <http://www.aclweb.org/anthology/P00-1037> (дата обращения: 19.01.2019).
3. DeepPavlov's documentation: Question Answering Model for SQuAD Dataset [Электронный ресурс]. – Режим доступа: <http://docs.deeppavlov.ai/en/latest/components/squad.html> (дата обращения: 19.01.2019).
4. Natural Language Computing Group, Microsoft Research Asia. R-Net: Machine Reading Comprehension with Self-Matching Networks [Электронный ресурс]. – Режим доступа: <https://www.microsoft.com/en-us/research/wp-content/uploads/2017/05/r-net.pdf> (дата обращения: 20.01.2019).
5. DeepPavlov's documentation: Open Domain Question Answering Skill on Wikipedia [Электронный ресурс]. – Режим доступа: <http://docs.deeppavlov.ai/en/latest/skills/odqa.html> (дата обращения: 20.01.2019).

**Мирзаянова Светлана Владимировна**

Год рождения: 1984

Университет ИТМО, факультет информационных технологий
и программирования, студент группы № М4221Направление подготовки: 09.04.02 – Информационные системы
и технологии

e-mail: ptitzalone@gmail.com

**Рыбин Сергей Витальевич**

Год рождения: 1959

Университет ИТМО, факультет информационных технологий
и программирования, к.ф.-м.н., доцент

e-mail: rybin@speechpro.com

УДК 004.934**ИСПОЛЬЗОВАНИЕ НЕЙРОННЫХ СЕТЕЙ ПРИ ОТБОРЕ ПРИЗНАКОВ В ЗАДАЧАХ
МАШИННОГО ОБУЧЕНИЯ****Мирзаянова С.В.****Научный руководитель – к.ф.-м.н., доцент Рыбин С.В.**

Работа выполнена в рамках темы НИР № 618278 «Синтез эмоциональной речи на основе генеративных состязательных сетей».

Целью работы являлся анализ применимости нейросетевых архитектур при отборе признаков в задачах определения тональности текста. Были рассмотрены два инструмента из набора word2vec и особенности их применения. Итогом проведенной работы служила обученная модель word2vec для дальнейшего получения признаков слов, для использования в классификаторе тональности текста.

Ключевые слова: анализ тональности, векторизация, рекуррентные нейронные сети, word2vec, word embedding, механизм внимания.

В современных задачах машинного обучения часто возникает необходимость обработки текста. С текстом связаны такие проблемы как фильтрация спама, автоматический перевод, голосовые помощники, выявление семантической близости слов, анализ тональности и многие другие. Задача анализа тональности текста может быть интересна для применения в различных системах, таких как синтез эмоциональной речи или рекомендательные системы, а также при анализе удовлетворенности клиентов и пользователей какого-либо бизнес-проекта.

Для использования методов машинного обучения в задачах анализа текста необходимо предоставлять слова или предложения в виде набора признаков, являющегося вектором. Векторизация слов применяется в большинстве задач NLP [1]. Представление текста в виде векторов позволяет использовать методы машинного обучения и нейронные сети для решения данных задач. Однако недостаточно представить текст в виде векторов слов. Простые вектора вида bag-of-words и one-hot vectors не содержат семантической информации [2]. Для задачи определения тональности семантическая информация несет ключевой смысл и нельзя ее упускать.

Получить семантическую информацию позволяют инструменты word2vec. Они позволяют объединить слова в кластеры в зависимости от контекста, в котором данные слова встречаются. Инструменты word2vec строятся на основе простой нейронной сети.

Семантическая близость слов оценивается как косинусное расстояние между векторами, представляющими данные слова [3]:

$$\cos(\varphi) = \frac{A \cdot B}{\|A\| \|B\|} = \frac{\sum_{i=1}^n A_i \times B_i}{\sqrt{\sum_{i=1}^n (A_i)^2} \times \sqrt{\sum_{i=1}^n (B_i)^2}},$$

где A и B – вектора, представляющие слова.

Различают два инструмента в составе word2vec. Первый называется Continue Bag of Words (CBOW) или последовательный мешок слов. При использовании данного инструмента в модель последовательно подаются наборы слов из текста, не учитывая порядок слов в тексте. Данный метод позволяет сформировать вектор слова в зависимости от контекста. Другой инструмент называется Skip-gram или словосочетание с пропуском. С его помощью по слову угадывается его контекст.

Увидеть особенности строения моделей CBOW и Skip-gram можно на рисунке ниже.

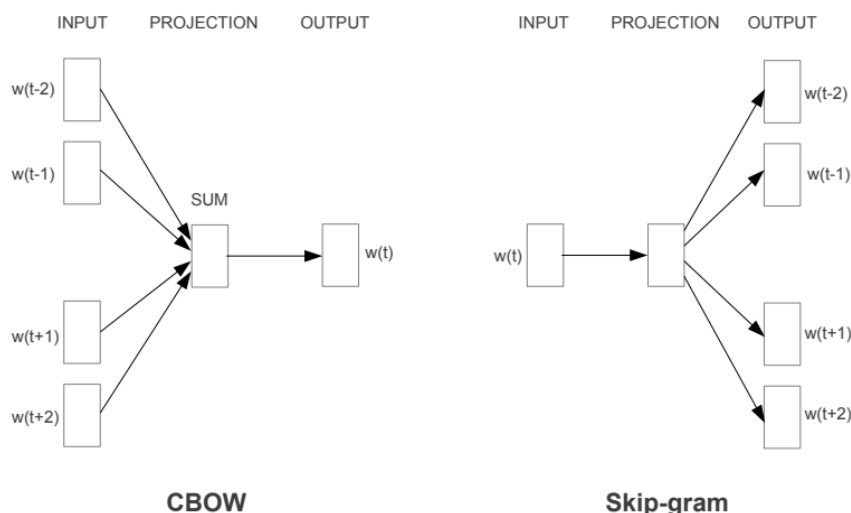


Рисунок. Схематическое изображение архитектур инструментов word2vec [3]

В рамках проводимого исследования были реализованы обе модели word2vec. Обучение моделей производилось на основе собственной базы данных, собранной из отзывов к кинофильмам с сайта «Кинопоиск». В результате проверки полученных зависимостей выяснилось, что модель CBOW лучше находит зависимости между словами по данной базе и разделяет слова разной тональности в пространстве. В таблице ниже представлены результаты сравнения векторов, полученных с помощью моделей CBOW и Skip-gram, обученных на собственной базе данных.

Таблица. Сравнение результатов обучения моделей CBOW и Skip-gram

Инструмент word2vec	CBOW	Skip-gram
Сходство между словами:		
«идеальный» и «отвратительный»	0,366	0,197
«замечательный» и «безупречный»	0,689	0,399
«плохо» и «ужасно»	0,790	0,431
Список ближайших слов к слову:		
«отлично»	великолепно, идеально, удачно, превосходно, достойно, безупречно, убедительно, неплохо	неразумный, шикарно, безукоризненно, драматически, превосходно, сколько-нибудь, эбигейл, меринг

Инструмент word2vec	CBOW	Skip-gram
«отвратительно»	ужасно, мерзко, безликий, убогий, шаблонный, плохо, удручающий, несмешной	ужасно, промотать, протухший, бесталанный, выговаривать, дилетант, промолчать, панфил

Целью использования инструментов word2vec при отборе признаков в задаче анализа тональности текста является представление слов в векторной форме таким образом, чтобы слова одинаковой ориентации тональности имели близкие вектора, а слова разной ориентации тональности имели сильно различающиеся вектора.

Как можно видеть из полученных результатов, близость полученных векторов слов зависит от ориентации их тональности. Сходство между словами одной тональности, как положительной, так и отрицательной, больше, чем сходство между словами разной тональности. Однако можно заметить, что в результатах модели CBOW слова, описывающие тональность, имеют более близкие вектора, чем эти же слова в модели Skip-gram. Но соотношение расстояний между словами при этом примерно одинаково в разных моделях.

Также интересно рассмотреть список ближайших слов в полученном векторном пространстве к интересующим нас тональным словам. Модель CBOW помещает рядом слова схожей тональности, которые могут служить, в некотором роде, синонимами рассматриваемому слову. Модель Skip-gram дает несколько другие результаты. В приведенном выше примере можно видеть, что среди ближайших векторов можно обнаружить как слова, близкие по смыслу, так и совершенно далекие слова. Такой результат связан с логикой работы модели Skip-gram. Здесь близкие вектора получают слова, встречающиеся в одном контексте. Таким образом, ближайшими словами оказываются слова, чаще всего присутствующие рядом с тональным словом, но не обязательно также описывающие тональность.

В результате проведенного исследования можно сделать вывод, что в задаче определения тональности текста для отбора признаков по имеющейся базе данных больше подходит инструмент CBOW, представляющий вектора одной тональности близкими векторами. В дальнейшем этот результат будет использоваться для получения векторов признаков для классификатора тональности.

Литература

1. Maas A.L., Daly R.E., Pham P.T., Huang D., Ng A.Y., Potts C. Learning Word Vectors for Sentiment Analysis // Association for Computational Linguistics. – 2011. – P. 142–150.
2. Kaliyev A., Rybin S.V., Matveev Y. The Pausing Method Based on Brown Clustering and Word Embedding // Lecture Notes in Computer Science. – 2017. – V. 10458. – P. 741–747.
3. Mikolov T., Chen K., Corrado G., Dean J. Efficient Estimation of Word Representations in Vector Space [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1301.3781.pdf> (дата обращения: 06.01.2019).



Мироненко Александр Алексеевич

Год рождения: 1995

Университет ИТМО, факультет информационных технологий
и программирования, студент группы № М4221

Направление подготовки: 09.04.02 – Информационные системы
и технологии

e-mail: hexkritor@yandex.com



Матвеев Юрий Николаевич

Год рождения: 1955

Университет ИТМО, факультет информационных технологий
и программирования, д.т.н.

e-mail: matveev@speechpro.com

УДК 004.93

АНАЛИЗ АЛГОРИТМОВ ТРЕКИНГА ЛИЦ С КОРРЕКТИРОВКОЙ ОШИБОК ТРЕКИНГА

Мироненко А.А.

Научный руководитель – д.т.н. Матвеев Ю.Н.

В работе были рассмотрены различные алгоритмы, используемые для трекинга лиц в видеопотоке. Технологии трекинга лиц могут быть использованы во многих областях современной человеческой деятельности. Несмотря на большое количество готовых решений в этой области, их исследование по-прежнему представляет интерес, так как есть простор для их усовершенствования или появления новых.

Ключевые слова: трекинг лиц, метод Виолы-Джонса, OpenCV, трекинг.

Цель работы – провести анализ работы алгоритмов трекинга, имеющих реализацию в библиотеке OpenCV, на наличие ошибок трекинга, в частности ошибки назначения трека при пересечении лиц или долгого отсутствия лица в видеопотоке.

На сегодняшний день системы трекинга находят широкое практическое применение. В частности, они используются в видеонаблюдении, дополненной реальности, мобильных приложениях, робототехнике, мониторинге дорожного движения. Задача трекинга лиц является одной из них. Существует множество методов и алгоритмов трекинга лиц [1], но, несмотря на их высокую точность работы, в некоторых ситуациях возникают ошибки [2]. Одной из них является ошибка назначения трека при пересечении лиц или долгого отсутствия лица в видеопотоке.

Для оценки работы алгоритмов использована видеозапись, являющаяся частью базы Elliptical Head Tracking [3], разрешением 128×96 пикселей и длительностью 101 кадр. На данной видеозаписи присутствуют фрагменты, при которых возникает отсутствие лица в кадре, путем поворота головы человека, а также пересечением лица другим человеком. Камера, использованная для видеозаписи, находится в движении, что усложняет условия работы трекеров. Для подсчета точности работы использованы метрики MOTA [4], оценивающая частоту появления ошибок при трекинге, и MOTP [4], оценивающая отклонение позиций трекера с истинным значением. Для точной оценки эффективности трекинга для данной задачи используется метрика IDS (ID Switches) [2], база была размечена вручную, что потребовалось в силу отсутствия данных для трекинга лица. Для тестирования работы алгоритмов реализована система трекинга (рисунок), использующая готовые решения,

реализованные в библиотеке OpenCV [5]. В качестве детектора используется метод Виолы-Джонса [6], а в качестве трекеров используются реализации на основе алгоритмов AdaBoost, MIL, MOSSE, TLD, Median Flow, а также фильтров KCF, CSRT.

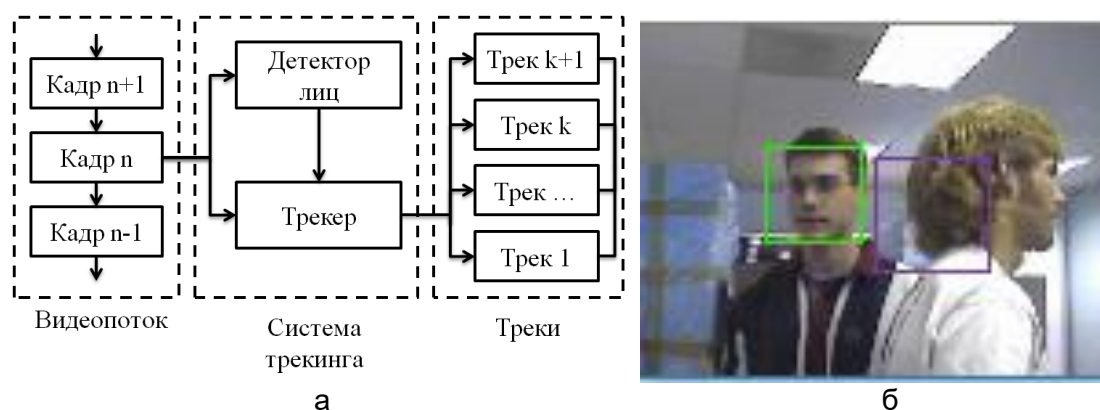


Рисунок. Обобщенная структура системы трекинга (а); пример работы системы трекинга (б)

Таблица. Результаты работы системы при использовании различных алгоритмов трекинга

Трекер	MOTA, %	MOTP, %	ID Switches
AdaBoost	-66,34	26,39	4
MIL	-67,33	37,36	7
MOSSE	52,48	26,67	6
TLD	-69,31	11,94	4
Median Flow	-18,81	36,95	3
KCF	41,59	23,05	5
CSRT	-132,67	28,61	5

Результаты исследования работы алгоритмов трекинга (таблица) показали, что на видеозаписи происходят частые смены идентификаторов трека, в особенности у трекеров MIL и CSRT. Низкие показатели MOTA у большинства трекеров вызваны тем, что в реализации трекеров не передавалась информация о потере объекта из поля зрения, из-за чего системой было допущено множество ошибок ложного срабатывания, что привело к отрицательному значению метрики. Полученные данные показывают, что проблема ошибок назначения трека при пересечении лиц или долгого отсутствия лица в видеопотоке актуальна и требует универсального решения.

Литература

1. Олейник А.Л. Применение бинарных дескрипторов для трекинга множества лиц в системах видеонаблюдения // Научно-технический вестник информационных технологий, механики и оптики. – 2016. – Т. 16. – № 4. – С. 670–677.
2. Hofmann M., Haag M., Rigoll G. Unified hierarchical multi-object tracking using global data association // IEEE International Workshop on Performance Evaluation of Tracking and Surveillance (PETS). – 2013. – P. 22–28.
3. Birchfield S. Elliptical head tracking using intensity gradients and color histograms // Proceedings. – 1998. – P. 232–237.
4. Bernardin K., Stiefelhausen R. Evaluating multiple object tracking performance: the CLEAR MOT metrics // Journal on Image and Video Processing. – 2008. – V. 2008. – P. 1.
5. OpenCV (Open source Computer Vision library) [Электронный ресурс]. – Режим доступа: <http://opencv.org/> (дата обращения: 11.02.2019).
6. Viola P., Jones M. Rapid object detection using a boosted cascade of simple features [Электронный ресурс]. – Режим доступа: <https://www.cs.cmu.edu/~srini/15-829/readings/ViJo01.pdf> (дата обращения: 06.01.2019).



Рюмин Дмитрий Александрович

Год рождения: 1991

Санкт-Петербургский институт информатики и автоматизации РАН,
мл.н.с.;

Университет ИТМО, факультет информационных технологий
и программирования, аспирант

Направление подготовки: 09.06.01 – Информатика и вычислительная
техника

e-mail: dl_03.03.1991@mail.ru



Аксёнов Александр Александрович

Год рождения: 1995

Санкт-Петербургский институт информатики и автоматизации РАН,
программист;

Университет ИТМО, факультет программной инженерии
и компьютерной техники, студент группы № Р4210

Направление подготовки: 09.04.04 – Информационно-вычислительные
системы

e-mail: a.aksenov95@mail.ru



Карпов Алексей Анатольевич

Год рождения: 1978

Санкт-Петербургский институт информатики и автоматизации РАН,
мл.н.с.;

Университет ИТМО, факультет информационных технологий
и программирования, д.т.н., профессор

e-mail: karpov@iias.spb.su

УДК 004.852

АВТОМАТИЧЕСКОЕ ОБНАРУЖЕНИЕ ЛИЦ ДЛЯ ЧЕЛОВЕКО-МАШИННОГО ВЗАИМОДЕЙСТВИЯ

Рюмин Д.А. (Санкт-Петербургский институт информатики и автоматизации РАН;
Университет ИТМО), **Аксёнов А.А.** (Санкт-Петербургский институт информатики
и автоматизации РАН; Университет ИТМО), **Карпов А.А.** (Санкт-Петербургский институт
информатики и автоматизации РАН; Университет ИТМО)

Научный руководитель – д.т.н., профессор Карпов А.А.
(Санкт-Петербургский институт информатики и автоматизации РАН; Университет ИТМО)

Работа выполнена при финансовой поддержке РФФИ № 18-07-01407 «Автоматическое
бимодальное распознавание естественных эмоций в русской речи», а также в рамках темы
НИР № 618278 «Синтез эмоциональной речи на основе генеративных состязательных сетей».

В работе предложена автоматическая система обнаружения лиц на основе сверточной нейронной сети.
Приведена подробная реализация системы с пояснением ключевых особенностей ее применения.
Проиллюстрированы результаты. Предложено дальнейшее применение системы в задачах биометрии,
компьютерном зрении, машинном обучении, автоматических системах распознавания лиц, речи и
элементов жестовых языков.

Ключевые слова: биометрия, распознавание лиц, компьютерное зрение, машинное обучение,
жестовый язык.

Введение. Наиболее сложным случаем постановки проблемы систем обнаружения лиц
является работа данных систем в сильно меняющихся средах с большим потоком входных
данных (городские улицы с интенсивным движением, метро, аэропорты, больницы, торговые

центры и т.д.). Для решения таких проблем необходимо использовать максимально доступную информацию для достижения приемлемых результатов детектирования лиц. Системы обнаружения лиц должны работать при разных условиях освещения, распознавать лица под разными углами и частичными перекрытиями областей лица. Для обеспечения таких возможностей требуется определенное оборудование – камеры с высоким разрешением и хорошей оптикой для захвата более точной информации о лицах. Кроме того, необходимо наличие размеченных баз данных с лицами, которые являются неотъемлемой частью любой системы распознавания лиц, которая базируется на современных алгоритмах детектирования и распознавания лиц.

Нынешнее исследование пытается частично решить описанные проблемы. Целью работы являлась разработка системы для надежного и точного детектирования лиц без ограничений в присутствии экстремальных изменений внешнего вида, освещения, макияжа, окклюзии и т.д. В работе использована собранная и размеченная база данных (БД), а также архитектура сверточной нейронной сети (СНС).

Описание автоматической системы. Сбор БД производился при помощи устройства Microsoft Kinect 2.0. Данное устройство позволяет осуществлять скелетное отслеживание до 6 человек с целью распознавания их действий на расстоянии от 1,2 до 3,5 м. Модель скелета человека разделяется на 25 суставов (элементов). Определение наличия пользователя в кадре производится на основе отслеживания головы и верхней части туловища человека [1]. Кроме того, Kinect 2.0 содержит инфракрасный излучатель, главное назначение которого – испускать инфракрасные лучи, которые отражаясь от предметов, попадают в инфракрасный приемник. Инфракрасный приемник собирает отраженные лучи с частотой 30 Гц и преобразует их в значения расстояния от сенсора до объекта или объектов. Таким образом, строится матрица расстояний для одного кадра, максимальное разрешение которого 512×424 пикселей. Помимо инфракрасного излучателя и приемника имеется цветная камера, используемая для захвата потока видео с углами обзора $43,5^\circ$ по вертикали и 57° по горизонтали и с максимальным разрешением 1920×1080 пикселей и частотой 30 Гц (15 Гц в условиях низкой освещенности).

Для хранения многомерных сигналов, полученных с сенсора Kinect 2.0 необходимо наличие многомодальной базы данных [2]. Наиболее оптимальным вариантом является БД с иерархической моделью (ИМ), которая может содержать структурированную информацию. Структурирование подразумевает явное выделение составных частей (элементов), связей между ними, а также типизацию элементов и связей, при которой с типом элемента (связи) соотносится определенная семантика и допустимые операции. С помощью ИМ все необходимые данные (в том числе и видеоданные) в совокупности можно представить в виде файловой системы, представленной на рис. 1, состоящей из корневого каталога и иерархии подкаталогов и файлов различного формата.

База данных состоит из набора видеозаписей воспроизведения определенных жестов различными дикторами. По каждому показанному демонстратором жесту в базе данных хранится следующая информация в виде отдельных файлов на диске:

- видеозаписи показанного жеста (оптическое разрешение для цветного формата 1920×1080 пикселей, для карты глубины и инфракрасного режима – 512×424 пикселей, частота кадров – 30 кадров в секунду);
- данные о координатах положения лиц на видеозаписях;
- изображения, выделенные покадрово из видеозаписей, необходимые для разметки.

Общее количество жестов составило 2132. Также в ходе исследования введен новый аннотированный набор видеоданных для детектирования лиц. Лицо человека локализуется с помощью ограничивающего прямоугольника (рис. 2), а прикрепленный ярлык соответствует лицу.

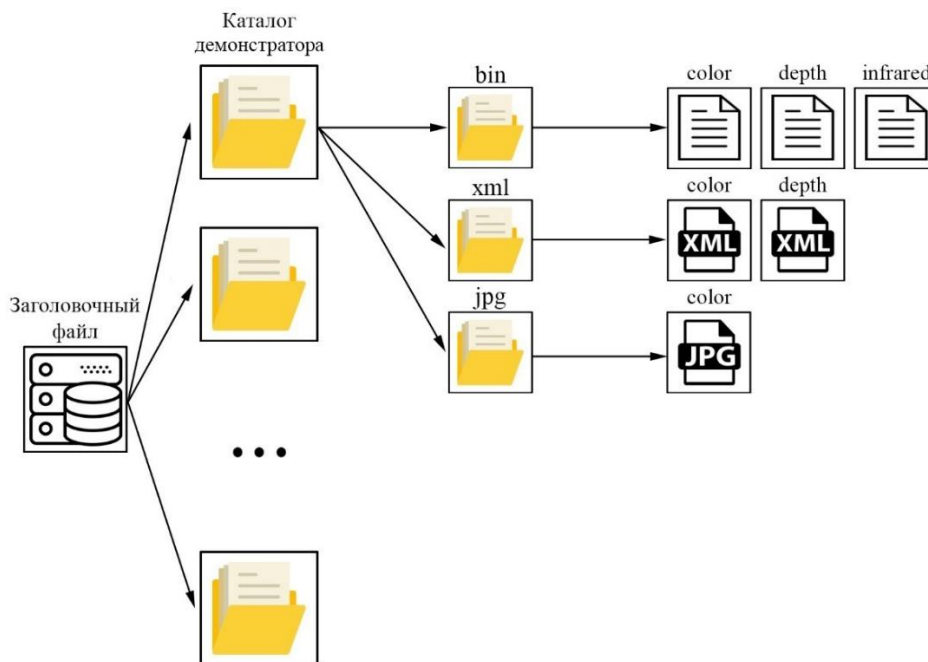


Рис. 1. Логическая структура базы данных



Рис. 2. Примеры цветных видеок кадров из базы данных

Анализ видеоинформации для обнаружения лиц осуществлялся за счет использования СНС, архитектура, которой представлена на рис. 3.

Современные детекторы на основе СНС, такие как Faster R-CNN [3], R-FCN [3], SSD [4], COCO [5] и YOLO [5], теперь достаточно хороши для развертывания в потребительских продуктах, а также достаточно быстры, чтобы работать на персональных компьютерах, мобильных устройствах и т.д. Однако на практике трудно решить, какая архитектура нейронной сети лучше всего подходит для анализа видеоинформации и дальнейшего детектирования лиц. Основываясь на стандартных показателях точности распознавания объекта, которым в данном исследовании выступает лицо человека, нельзя в полной степени оценить работу детектора, так как важное значение имеют скорость работы и использование памяти. Например, от автоматических систем анализа информации для определения и распознавания лиц требуется работа в режиме реального времени. Из описанного становится

понятно, что необходимо найти компромисс скорости по отношению к точности детектора для дальнейшего его использования в системе распознавания лиц.

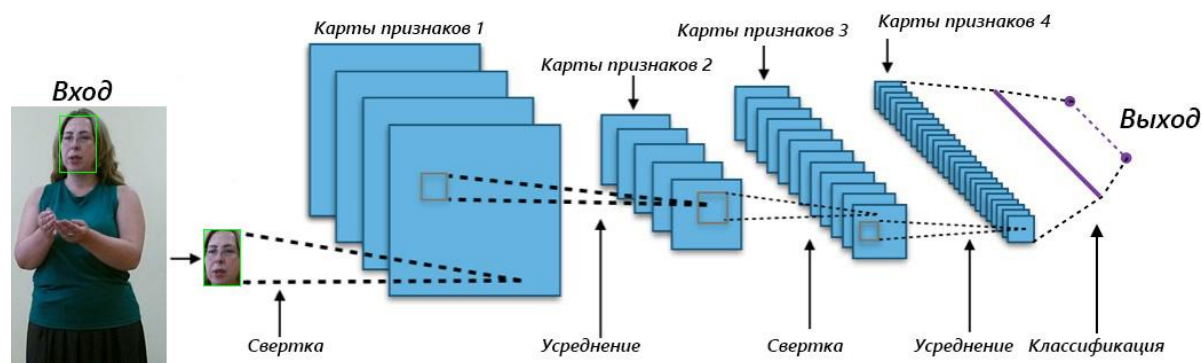


Рис. 3. Архитектура сверточной нейронной сети

На текущий момент времени существует подмножество детекторов объектов (Faster R-CNN, R-FCN, SSD, COCO и YOLO), которые позволяют учитывать скорость работы. Кроме того, они позволяют достичь определенной частоты кадров, а также предоставляют полную оценку, которая зависит от многих других факторов, таких как извлечение признаков области с лицом, размеры входного видеопотока данных, и т.п. Также данные детекторы схожи по общей методологии. Это дает возможность реализовать собственный метод и сравнить его с существующими сверточными нейронными сетями анализа видеoinформации для определения областей с лицами унифицированным образом. В частности, архитектуры детекторов Faster R-CNN, R-FCN, SSD, COCO и YOLO состоят из единой СНС с использованием смешанной регрессии и классификации.

Функциональная схема предлагаемого метода анализа видеoinформации для определения лиц представлена на рис. 4, а.

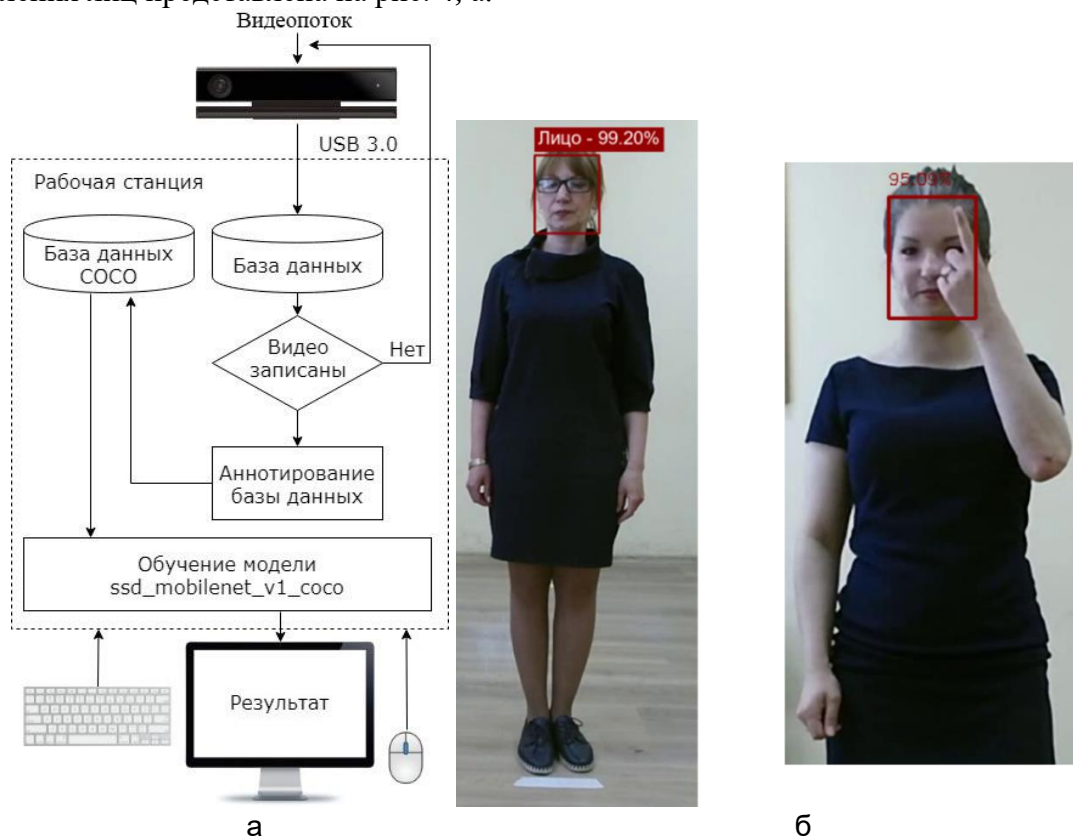


Рис. 4. Функциональная схема метода анализа видеoinформации для обнаружения лиц (а); результаты системы детекции лиц (б)

Результатом исследования является модель `ssd_mobilenet_v1_coco` обученная на БД СОСО, а также собственной записанной и аннотированной БД. Такое объединение БД позволило не только ускорить процесс обучения, но и построить систему анализа видеоинформации для определения лиц (рис. 4, б).

Заключение. Отличительная особенность предложенной автоматической системы – быстроедействие. Таким образом, описанная система позволяет определять в видеопотоке область лица. В дальнейшем за счет своей универсальности система может использоваться в задачах машинного обучения, компьютерном зрении, биометрии, автоматических системах распознавания лиц, речи и элементов жестовых языков. В текущих исследованиях модель используется в системе многомодального распознавания элементов русского жестового языка.

Литература

1. Ryumin D., Karpov A. Towards Automatic Recognition of Sign Language Gestures using Kinect 2.0 // In Proc. 19th International Conference on Human-Computer Interaction, HCII, Vancouver, Canada, Springer LNCS. – 2017. – V. 10278. – P. 89–104.
2. Ivanko D., Karpov A., Fedotov D., Kipyatkova I., Ryumin D., Ivanko Dm., Minker W., Zelezny M. Multimodal Speech Recognition: Increasing Accuracy using High Speed Video Data // Journal on Multimodal User Interfaces, Springer. – 2018. – V. 12. – № 4. – P. 319–328.
3. Ren S., He K., Girshick R., Sun J. Faster R-CNN: Towards real-time object detection with region proposal networks [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1506.01497.pdf> (дата обращения: 06.01.2019).
4. Liu W., Anguelov D., Erhan D., Szegedy C., Reed S., Fu C–Y., Berg A. SSD: Single Shot MultiBox Detector // European conference on computer vision, ECCV, Springer, Cham. – 2016. – P. 21–37.
5. Redmon J., Divvala S., Girshick R., Farhadi A. You only look once: Unified, real-time object detection // Proceedings of the IEEE conference on computer vision and pattern recognition. – 2016. – P. 779–788.

**Столбов Михаил Борисович**

Год рождения: 1952

Университет ИТМО, факультет информационных технологий
и программирования, к.т.н., доцент

e-mail: stolbov@speechpro.com

**Куан Чонг Тхе**

Год рождения: 1984

Университет ИТМО, факультет информационных технологий
и программирования, аспирантНаправление подготовки: 09.06.01 – Информатика и вычислительная
техника

e-mail: quantrongthe@corp.ifmo.ru

УДК 621.391.8

**АДАПТИВНЫЙ АЛГОРИТМ MVDR ДЛЯ ДВУХЭЛЕМЕНТНОЙ МИКРОФОННОЙ
РЕШЕТКИ****Столбов М.Б., Куан Ч.Т.****Научный руководитель – к.т.н., доцент Столбов М.Б.**

Работа выполнена в рамках темы НИР № 618278 «Синтез эмоциональной речи на основе генеративных состязательных сетей».

Рассмотрен адаптивный алгоритм минимума дисперсии шума для выделения речи из когерентного шума с использованием сигналов двухэлементных микрофонных решеток. Проанализировано шумоподавление в реальных случаях.

Ключевые слова: MVDR, микрофонные решетки, шумоподавление.

Введение. Двухэлементные микрофонные решетки (MP2) получили широкое применение благодаря своей простоте и возможности компактного размещения. Вопросы обработки сигналов MP2 рассмотрены во многих работах, например, [1–3]. В работе рассмотрена задача выделения речевых сигналов в присутствии широкополосного когерентного шума с использованием MP2. Работа посвящена исследованию адаптивного алгоритма (Minimum Variance Distortionless Response, MVDR) [4] и его связи с алгоритмом дифференциальных MP2 (ДМР2) [2].

Модель сигнала. Схема обработки сигналов MP2 в частотной области представлена на рис. 1. MP2 состоит из двух ненаправленных микрофонов, разнесенных на расстояние d . Направление прихода целевого сигнала и шума задаются углами θ_s , θ_v относительно оси, проходящей через микрофоны. В представлении кратковременного дискретного преобразования Фурье (ДПФ) сигнал $S(f, k)$ акустического целевого источника с направления θ_s и когерентный широкополосный шум $V(f, k)$ с направления $\theta_v(f)$ формирует на микрофонах вектор сигналов:

$$\mathbf{X}(f, k) = S(f, k)\mathbf{D}_s(f) + V(f, k)\mathbf{D}_v(f) = \mathbf{S}(f, k) + \mathbf{V}(f, k),$$

где f, k – индекс частоты и номера кадра; $\mathbf{D}(f, \theta_x) = \mathbf{D}_x(f) = [e^{+j\Phi_x}, e^{-j\Phi_x}]^T$ – вектор фазовых сдвигов сигналов микрофонов относительно центральной точки между ними, и $()^T$ – символ транспонирования, $\Phi_x(f)$ – фазовые сдвиги:

$$\Phi_x(f) = \pi d \cos(\theta_x) / \lambda = \pi f \tau_0 \cos(\theta_x),$$

где d – расстояние между микрофонами; c – скорость звука в воздухе; $\tau_0 = d/c$ – время прохождения звука по оси МР2 между микрофонами; θ_x – направление прихода сигнала.

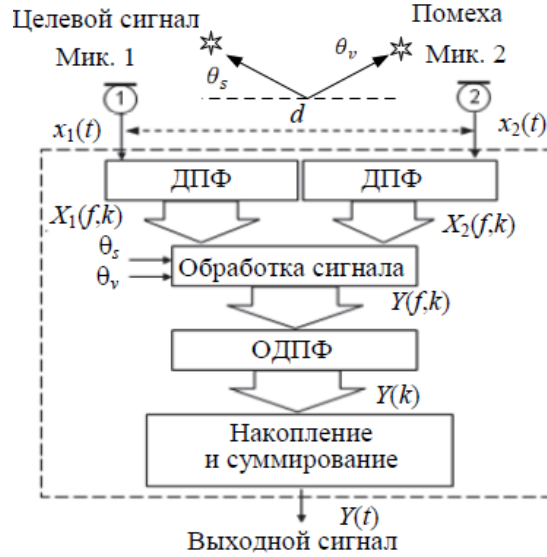


Рис. 1. Схема обработки сигналов МР2 в частотной области

Сигнал на выходе МР2 формируется взвешиванием входных сигналов комплексными весовыми коэффициентами:

$$Y(f, k) = \mathbf{W}^H(f) \mathbf{X}(f, k),$$

где $\mathbf{W}(f)$ – вектор коэффициентов; $(\cdot)^H$ – символ сопряжения Эрмита; $Y(f, k)$ – комплексные амплитуды спектра сигнала на k -м кадре.

Алгоритм MVDR основывается на условии неискаженного приема сигнала с целевого направления θ_s и минимизации мощности шума на выходе МР2. Оптимальные веса МР2 являются решением следующей оптимизационной задачи [1]:

$$\min_{\mathbf{W}} E\{|Y(f, k)|^2\} = \min_{\mathbf{W}} [\mathbf{W}^H(f) \mathbf{P}_{vv}(f) \mathbf{W}(f)],$$

при условии $\mathbf{W}^H(f) \mathbf{D}_s(f) = 1$, где $\mathbf{P}_{vv}(f)$ – ковариационная матрица спектров шума на микрофонах; $E\{\cdot\}$ – символ математического ожидания.

Ковариационная матрица шума оценивается для каждой частоты:

$$\mathbf{P}_{vv}(f) = E\{\mathbf{V}(f, k) \mathbf{V}^H(f, k)\}.$$

Решение оптимизационной задачи приводит к следующему соотношению для вектора оптимальных весов [1]:

$$\mathbf{W}_0(f) = \frac{\mathbf{P}_{vv}^{-1}(f) \mathbf{D}_s(f)}{\mathbf{D}_s^H(f) \mathbf{P}_{vv}^{-1}(f) \mathbf{D}_s(f)}.$$

В случае помехи, поступающей с направления θ_v , МР2 имеет следующий пространственный отклик [5]:

$$H_{MVDR}(f, \theta_x, \theta_s, \theta_v) = |\mathbf{W}_0^H(f, \theta_v, \theta_s) \mathbf{D}(f, \theta_x)| = \left| \frac{\sin(\Phi_x - \Phi_v)}{\sin(\Phi_s - \Phi_v)} \right|.$$

Отклик в направлении источника помехи равен нулю: $H_{dif}(f, \theta_x = \theta_v, \theta_v) = 0$, отклик в направлении источника целевого сигнала равен единице: $H_{MVDR}(f, \theta_s, \theta_v) = 1$.

Пространственно-частотный отклик MVDR можно представить в виде произведения передаточной функции эквалайзера и передаточной функции ДМР2:

$$H_{MVDR}(f, \theta_x, \theta_s, \theta_v) = H_{eq}(f, \theta_s, \theta_v) \times H_{dif}(f, \theta_x, \theta_v).$$

Адаптивный алгоритм MVDR. В адаптивном алгоритме MVDR ковариационная матрица и оптимальные веса оцениваются (в паузах целевого сигнала) по поступающим сигналам микрофонов:

$$\mathbf{P}_{vv}(f, k) = (1 - \beta)\mathbf{P}_{vv}(f, k - 1) + \beta\mathbf{X}(f, k)\mathbf{X}^H(f, k),$$

$$\mathbf{W}_0(f, k) = \frac{\mathbf{P}_{vv}^{-1}(f, k)\mathbf{D}_s(f)}{\mathbf{D}_s^H(f)\mathbf{P}_{vv}^{-1}(f, k)\mathbf{D}_s(f)},$$

где β – коэффициент сглаживания.

Сигнал на выходе вычисляется с использованием текущих значений коэффициентов:

$$Y(f, k) = \mathbf{W}_0^H(f, k)\mathbf{X}(f, k).$$

Вектор отсчетов сигнала $\mathbf{Y}(k)$ на k -м кадре вычисляется с помощью дискретного обратного преобразования Фурье (ОДПФ, IDFT):

$$\mathbf{Y}(k) = \text{IDFT}\{Y(f, k)\}.$$

Последовательность векторов отсчетов $\{\mathbf{Y}(k)\}$ преобразуется в выходную последовательность отсчетов сигнала на основе процедуры пересечения и суммирования (OverLap-and-Add, OLA): $y(t) = \text{OLA}\{\mathbf{Y}(k)\}$.

Эксперименты. Адаптивный алгоритм MVDR и алгоритм ДМР2 исследовались на сигналах, записанных в безэховой камере на микрофоны ($d=5$ см). Ноль ДМР2 формировался в направлении $\varphi_v \approx 30^\circ$. Широкополосная помеха с акустической колонки поступала на МР2 с направлений $\varphi_v(0^\circ - 90^\circ)$ относительно нормали МР2 с шагом 5° . Каждое положение акустической колонки озвучивалось речью диктора, находившегося приблизительно в направлении $\varphi_s = -30^\circ$ (рис. 2).

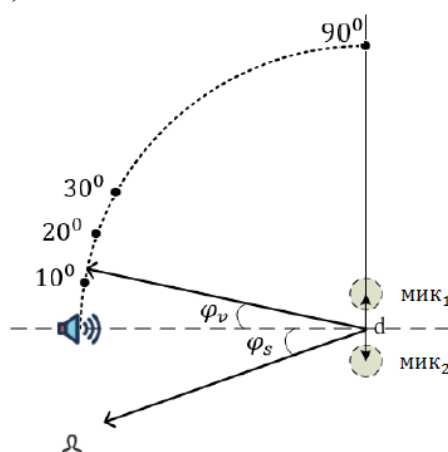


Рис. 2. Схема эксперимента в безэховой камере

Обработка проводилась со следующими параметрами: частота дискретизации $F_s=16$ кГц, размер кадра $N = 512$, $\beta = 0,5$. Результаты эксперимента представлены на рис. 3, 4.

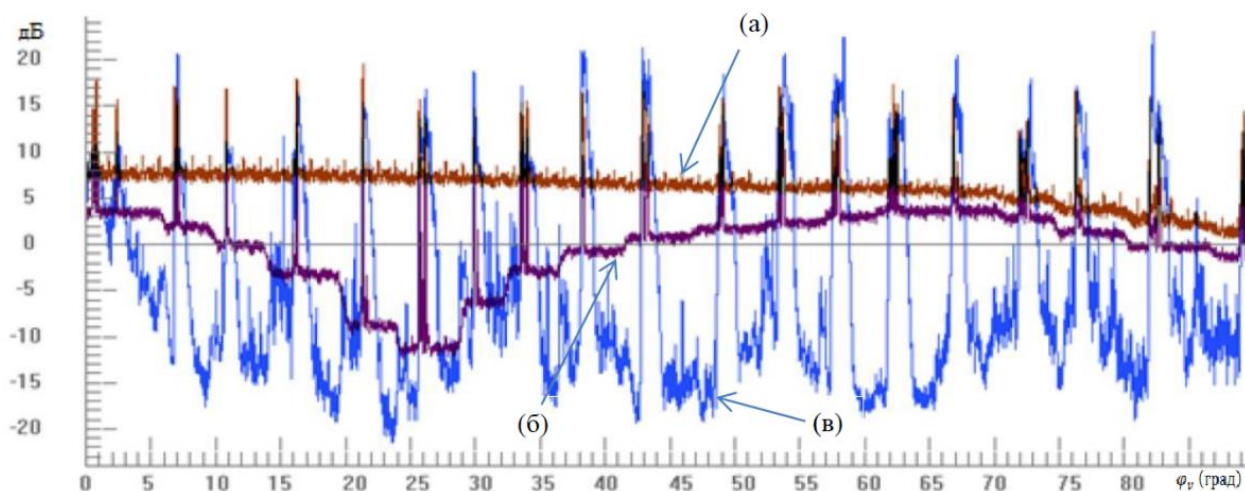


Рис. 3. Мощности сигналов: микрофона (а); ДМР2 (ноль $\varphi_v \approx 30^\circ$) (б); MVDR (в)

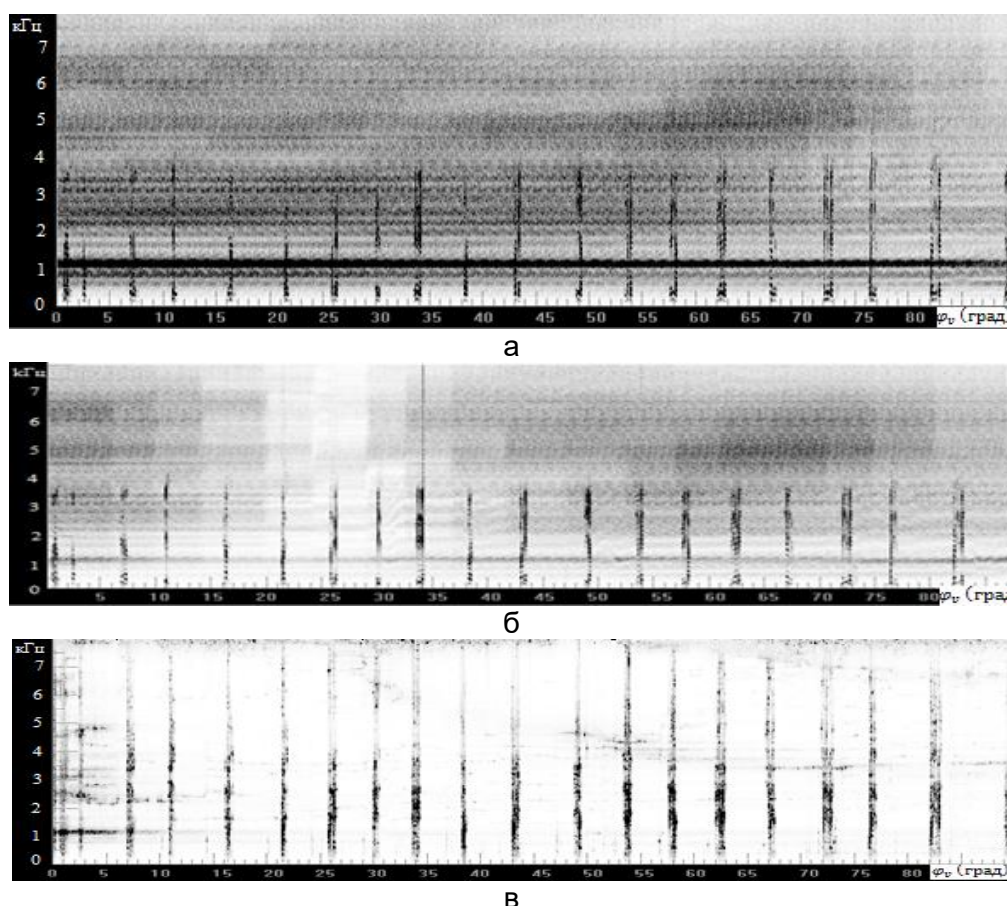


Рис. 4. Спектрограммы сигналов: микрофона (а); ДМР2 (б); МВДР (в)

Из эксперимента в безэховой камере можно сделать следующие выводы.

Адаптивный алгоритм MVDR позволяет отслеживать направление источника когерентного шума и подавлять этот шум. Степень подавления когерентного шума в алгоритме MVDR составила 18–20 дБ и ограничивалась, по всей вероятности, различием характеристик микрофонов. Степень подавления когерентного шума в алгоритмах MVDR и ДМР2 в направлении нуля пространственного отклика приблизительно совпадает. ДМР2 имеет мгновенный отклик на речевой сигнал диктора, в то время как алгоритм MVDR дает инерционный отклик, определяемый величиной коэффициента сглаживания.

Заключение. МР2 с адаптивным алгоритмом MVDR является обобщением ДМР2 для произвольного (отличного от 0°) направления целевого источника и различных направлений источников помех для разных частот.

Литература

1. Brandstein M., Ward D. Microphone Arrays: Signal Processing Techniques and Applications. – Springer, 2001. – 402 p.
2. Benesty J., Chen J. Study and Design of Differential Microphone Arrays. – Springer, 2013. – 184 p.
3. Benesty J., Chen J., Pan C. Fundamentals of Differential Beamforming. – Springer Science+Business Media, Singapore, Pte, Ltd, 2016. – 129 p.
4. Lockwood M. et al. Performance of time- and frequency-domain binaural beamformers based on recorded signals from real rooms // J. Acoust. Soc. Am. – 2004. – V. 115(1). – P. 379–391.
5. Столбов М.Б., Перелыгин С.В. Алгоритмы двухэлементной микрофонной решетки для выделения речевых сигналов в присутствии когерентных помех // Цифровая обработка сигналов. – 2017. – № 4. – С. 34–39.

**Фетисова Татьяна Анатольевна**

Год рождения: 1997

Санкт-Петербургский государственный университет,
факультет свободных искусств и наук, студентНаправление подготовки: 8.50.03.01 – Искусства и гуманитарные науки

e-mail: snowboat123@gmail.com

**Черных Герман Анатольевич**

Год рождения: 1967

Санкт-Петербургский государственный университет,
факультет свободных искусств и наук, к.ф.-м.н.

e-mail: german.ch@chgerman.com

УДК 004.89**РАЗРАБОТКА ДИАЛОГОВОГО АССИСТЕНТА ДЛЯ ПОЛУЧЕНИЯ ТЕКСТОВОЙ
ИНФОРМАЦИИ ИЗ СТРУКТУРИРОВАННЫХ БАЗ ЗНАНИЙ****Фетисова Т.А.** (Санкт-Петербургский государственный университет),**Черных Г.А.** (Санкт-Петербургский государственный университет)**Научный руководитель – к.ф.-м.н. Черных Г.А.**

(Санкт-Петербургский государственный университет)

Целью работы являлась разработка диалогового ассистента, предназначенного для автоматизации процесса получения информации из структурированных текстовых баз знаний.

Ключевые слова: диалоговый ассистент, чатбот, goal-oriented система, парсинг текста, нормализация текста, стемминг, поиск ключевых слов.

В настоящее время задача разработки автоматических диалоговых систем приобретает все большую актуальность. Широкое распространение диалоговых систем, работающих на основе методов обработки и анализа естественного языка (Natural Language Processing, NLP), происходит как в сфере обслуживания (контактные центры и службы поддержки), так и в сфере повседневной жизни современного человека (интеллектуальные помощники, диалоговые ассистенты).

Наиболее востребованы в настоящее время целе-ориентированные (или goal-oriented) сценарии использования диалоговых ассистентов [1, 2]. В отличие от чатботов (используемых в сценарии chit-chat), целью которых является поддержание диалога на естественном языке, goal-oriented диалоговые ассистенты предназначены для решения проблемы пользователя и получения конкретного результата. Например, к goal-oriented диалоговым помощникам можно отнести автоматизацию таких популярных в современном мире сценариев обслуживания, как: заказ такси, онлайн-покупка пиццы, бронирование билетов на поезд, получение информации о расписании занятий в вузе и другие целе-ориентированные задачи.

Основной сложностью при реализации автоматических диалоговых ассистентов для goal-oriented сценариев является невозможность решения проблемы пользователя и достижения целевого результата без получения информации из внешних баз знаний. Например, в качестве таких баз знаний могут выступать: продуктовый каталог, содержащий цены на различные виды пиццы в интернет магазине; расписание поездов с ценами на билеты разной категории и информацией о наличии свободных мест и т.д. Таким образом, основным

и необходимым шагом для разработки goal-oriented диалоговых помощников является интеграция с внешними базами знаний.

Настоящая работа решает задачу разработки диалогового ассистента, предназначенного для решения goal-oriented задачи – получения информации о расписании занятий в вузе. Разработанные методы были реализованы на языке программирования Python в интегрированной среде разработки программного обеспечения IntelliJ IDEA.

В ходе разработки были решены следующие научно-технические задачи:

- интеграция с внешней базой знаний, представленной в виде excel-таблиц, содержащих расписание занятий текущего семестра факультета свободных искусств и наук Санкт-Петербургского государственного университета;
- разработка текстового бота на платформе сервиса обмена мгновенными сообщениями (мессенджера) Telegram. Выбор платформы для разработки бота был обусловлен как статистикой его использования среди предполагаемых пользователей разрабатываемого диалогового ассистента, так и наличием в свободном доступе удобного в использовании API (Application Programming Interface) – программного интерфейса для разработки текстовых ботов.

Для решения поставленных задач были выполнены следующие шаги:

1. постановка проблемы и анализ возможности создания решения с реальной практической пользой;
2. исследование статистики использования различных мессенджеров среди потенциальных пользователей разрабатываемого диалогового ассистента (студентов вуза) и выбор наиболее популярного;
3. выбор и создание среды разработки:
 - установка мессенджера, создание канала бота (диалогового ассистента);
 - получение token – уникального ключа для доступа к созданному боту;
 - импорт необходимых python-библиотек в среду разработки:
 - Py Telegram Bot API – для создания бота внутри среды;
 - ruexcel – для чтения и редактирования информации в файлах типа .xlsx, в котором хранятся данные расписания;
 - pandas – для анализа данных;
4. разработка компонента, осуществляющего интеграцию с внешней базой знаний:
 - разработка формальной грамматики для последующего извлечения информации из базы знаний, содержащей расписание занятий. Составление грамматических паттернов для распознавания и категоризации данных с последующим извлечением таких элементов информации, как: дни недели, предметы, преподаватели, время и номер аудитории;
 - создание словаря синонимов, характерных для данной предметной области. Сформированный словарь синонимов устанавливает смысловое равенство таких синонимичных объектов, как, например, различное текстовое представление дней недели: «понедельник» и «пятница»;
 - парсинг текстовых данных, содержащихся в базе знаний, т.е. сопоставление последовательности слов или символов разработанной формальной грамматике и извлечение необходимой для поддержания диалога информации в структурированном виде;
5. разработка компонента, осуществляющего анализ текстовых запросов к диалоговому ассистенту и формирование ответа:
 - нормализация текста пользовательского запроса путем применения метода стемминга (извлечения грамматической основы) слов. Таким образом, осуществляется исключение грамматической информации с сохранением смысловой составляющей;

- анализ нормализованного текстового запроса пользователя: извлечение из запроса ключевых слов, соответствующих используемым в системе грамматическим паттернам (дни недели, предметы, преподаватели, время и номер аудитории);
- поиск в базе знаний, представленной в виде структурированных данных, сегментов текстовой информации, соответствующей извлеченным из пользовательского запроса ключевым словам. Возвращение найденных сегментов информации пользователю в качестве ответа диалогового ассистента (бота).

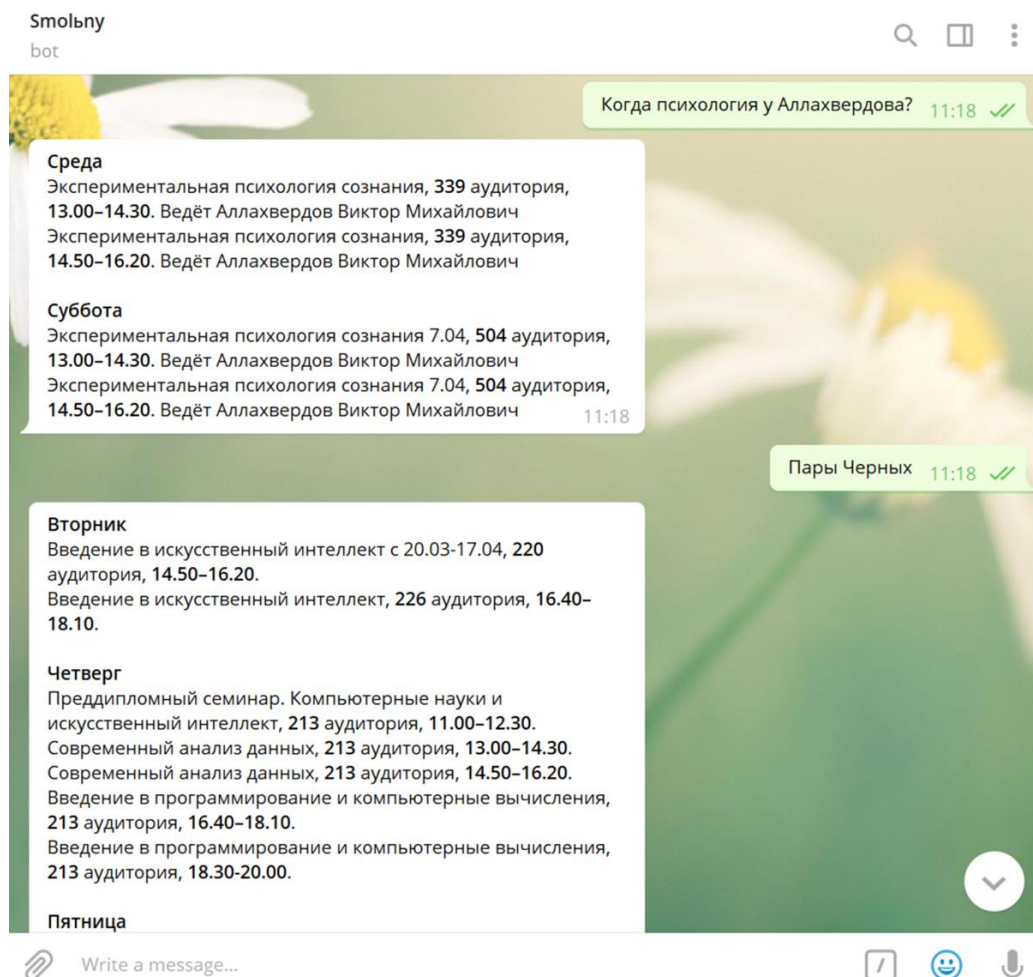


Рисунок. Пример использования разработанного диалогового ассистента

В результате проведенных работ был разработан диалоговый ассистент на базе популярного среди целевой аудитории мессенджера с удобным для использования интерфейсом (рисунок). Диалоговый ассистент позволяет формулировать запросы на естественном языке и предоставляет доступ только к необходимым пользователю сегментам информации, хранящейся в структурированной базе знаний. Применение разработанного диалогового ассистента, как студентами, так и преподавателями, позволило значительно ускорить процесс поиска целевой информации о расписании занятий в рамках факультета.

В качестве путей развития разработанного решения запланированы следующие работы: усовершенствование алгоритмов поиска информации (применение методов машинного обучения); повышение разнообразия диалоговых паттернов в рамках заданного сценария обслуживания; обеспечение большей гибкости при обработке запросов на естественном языке (в том числе за счет извлечения признаков на основе векторных представлений слов word2vec [3]); а также автоматическое оповещение пользователей системы (студентов) об изменениях информации в базе данных, таких как: отмена занятия, замена преподавателя, замена аудитории и т.п.

Литература

1. Crockett K., James O.S., Bandar Z. Goal orientated conversational agents: applications to benefit society // KES International Symposium on Agent and Multi-Agent Systems: Technologies and Applications. – 2011. – P. 16–25.
2. Akerkar R., Sajja P. Knowledge-based systems. – Jones & Bartlett Publishers, 2010. – 354 p.
3. Mikolov T. and etc. Efficient Estimation of Word Representations in Vector Space [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1301.3781.pdf> (дата обращения: 06.01.2019).

**Ховричев Михаил Аркадьевич**

Год рождения: 1995

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С4218сНаправление подготовки: 01.04.02 – Прикладная математика
и информатика

e-mail: mikhovr@niuitmo.ru

**Черных Ирина Александровна**

Год рождения: 1981

Университет ИТМО, факультет информационных технологий
и программирования, инженер

e-mail: chernykh-i@speechpro.com

**Нугманова Айгуль Айратовна**

Год рождения: 1995

Университет ИТМО, факультет информационных технологий
и программирования, аспирантНаправление подготовки: 09.04.02 – Информационные системы
и технологии

e-mail: nugmanova@speechpro.com

**Булусева Анна Владимировна**

Год рождения: 1981

ООО «ЦРТ-инновации», н.с.

e-mail: bulusheva@speechpro.com

**Кабаров Владимир Иосифович**

Год рождения: 1959

Университет ИТМО, факультет информационных технологий
и программирования, ст. преподаватель

e-mail: kabarov@speechpro.com

УДК 004.934.2

**ПРИМЕНЕНИЕ КОНТЕКСТНО-ЗАВИСИМЫХ МЕТОДОВ ДЛЯ ИЗВЛЕЧЕНИЯ
СИНОНИМОВ И КОНЦЕПТОВ ИЗ НЕСТРУКТУРИРОВАННЫХ ТЕКСТОВ
НА ЕСТЕСТВЕННОМ ЯЗЫКЕ****Ховричев М.А., Черных И.А., Нугманова А.А., Булусева А.В., Кабаров В.И.****Научный руководитель – Черных И.А.**

Работа выполнена при финансовой поддержке Министерства образования и науки Российской Федерации, соглашение 14.575.21.0178 (RFMEFI57518X0178).

В работе описана методология извлечения синонимичных рядов и концептов предметной области из корпуса текстов на естественном языке. Алгоритмы, лежащие в основе методологии, являются контекстно-зависимыми и не требуют использования размеченных данных и внешних ресурсов. Разработанные методы программно реализованы и применены к реальным данным, содержащим диалоги клиент-оператор из текстового чата службы поддержки крупного российского оператора связи.

Ключевые слова: обработка естественного языка, семантический анализ, извлечение информации.

Одним из этапов обработки текстов на естественном языке является семантический анализ. На этом этапе производится формирование семантических представлений и установление семантической связи между компонентами текста. Под семантической связью между отдельными словами понимаются отношения синонимии, антонимии, гипо- и гиперонимии и т.п. С точки зрения информационного поиска семантический анализ необходим для установления близости смысловых единиц, объединения слов, имеющих общее смысловое значение и обработки обобщений.

Предметная область настоящего исследования – телекоммуникации и связь, что определяет дальнейшую специфику работы с данными, представляющими из себя диалоги с операторами. Для конкретизации достаточно широких понятий «синоним» и «концепт» введем следующие определения:

- синонимы – текстовые единицы (слова или фразы), имеющие общее или близкое значение и употребляемые в схожем окружении в некоторой текстовой базе (дистрибутивное определение);
- концепты:
 1. группы синонимов (синсеты) в данной предметной области («сообщение, уведомление, предупреждение»);
 2. слоты (сущности) в данной предметной области (например, названия тарифов). Понятие концепта имеет некоторое соответствие лингвистическому термину «семантическое поле».

Для извлечения синонимических связей могут привлекаться внешние ресурсы, такие как WordNet [1]. Тем не менее, WordNet и подобные ресурсы содержат только общезыковые синонимы, но не содержат контекстуальные синонимы, характерные для конкретной предметной области или отдельно взятого корпуса текстов. Целью данной работы является разработка и применение ресурсно-независимых методов для извлечения синонимов и концептов из неструктурированных текстов, относящихся к одной предметной области. Рассмотренные в работе контекстно-зависимые методы для извлечения синонимов и концептов опираются на дистрибутивную гипотезу – идею о том, что слова, встречающиеся в схожих контекстах, близки по смыслу.

Основным объектом применяемых методов являются векторные представления слов из исследуемого корпуса текстов. Для получения векторных представлений необходимо задать алгоритм отображения слов на пространство $\mathbb{R}^n$. Наиболее популярным методом является Word2vec [2], также основанный на дистрибутивной гипотезе. Применение Word2vec предполагает использование одной из двух нейросетевых моделей – CBOW (Continue Bag of Words) и Skip-gram. Модель CBOW обучается предсказывать слово w_i по контексту $\{w_{i-k}, \dots, w_{i-1}, w_{i+1}, w_{i+k}\}$, где k – ширина контекстного окна. В модели Skip-gram вход и выход модели меняются местами (по слову предсказывается наиболее вероятный контекст). Использование любой из этих моделей в основе метода Word2vec позволяет получить подходящие для решения поставленной задачи векторные представления. Существуют и другие методы получения векторных представлений (fastText, GloVe, ELMo). В разработанных методах был использован Word2vec ввиду явного учета контекста векторной моделью.

Для получения групп синонимичных слов применен метод кластеризации векторных представлений слов. Выбор алгоритма кластеризации влечет за собой выбор параметров данного алгоритма. Например, популярный алгоритм k -средних предполагает, что количество кластеров известно заранее. На практике, как правило, неизвестно, сколько групп синонимов имеется в корпусе текстов. Это накладывает необходимость задания эвристик на параметры кластеризующих алгоритмов. Одним из направлений дальнейших исследований является автоматизация определения параметров алгоритмов кластеризации, в частности, количества кластеров. Ввиду возможного наличия многозначных слов в корпусе текстов имеет смысл использовать алгоритмы мягкой (нечеткой) кластеризации, такие как нечеткий метод s -средних [3]. В разработанных алгоритмах применялась как жесткая кластеризация методом k -средних, так и нечеткий метод s -средних.

В результате кластеризации и последующего отображения векторных представлений на словарь корпуса образуются группы слов. В случае использования нечеткой кластеризации пересечение этих множеств не пусто. Шумы представляют собой кластеры слов, не имеющих общей семантики (часто это служебные и редко используемые в тексте слова). Возможны следующие типы ошибок: в синсеты попадают служебные или лишние слова; словоформы одной и той же леммы попадают в разные кластеры; в шумы попадают слова, которые должны принадлежать синсетам. С целью минимизации ошибок производится постобработка полученных синсетов. Постобработка включает в себя:

1. выявление распределения частей речи слов, попавших в синсет;
2. исключение наиболее редких частей речи и служебных слов из синсетов, где выделяется доминирующая часть речи (более 50% синсета одной части речи).

Результирующие синсеты рассматриваются как концепты предметной области. В зависимости от множества слов, образующих концепт, рассматривается возможность подбора слотового имени концепта. Для извлечения из концептов слотового имени применяется следующие шаги:

1. стемминг синсета, образующего концепт, в результате получается множество стемов – основ слов;
2. получение BOW-вектора синсета и последующее TF-IDF взвешивание стемов;
3. фильтрация стемов по эвристически заданному порогу TF-IDF. Стемы, веса которых выше порога, являются кандидатами на слотовое имя;
4. сравнение максимального веса с оставшимися после фильтрации весами: слотовое имя образует либо стем с максимальным весом, либо несколько стемов с примерно равными весами ($\frac{w_i}{w_k} \in U_\varepsilon(1), \varepsilon \ll 1$).

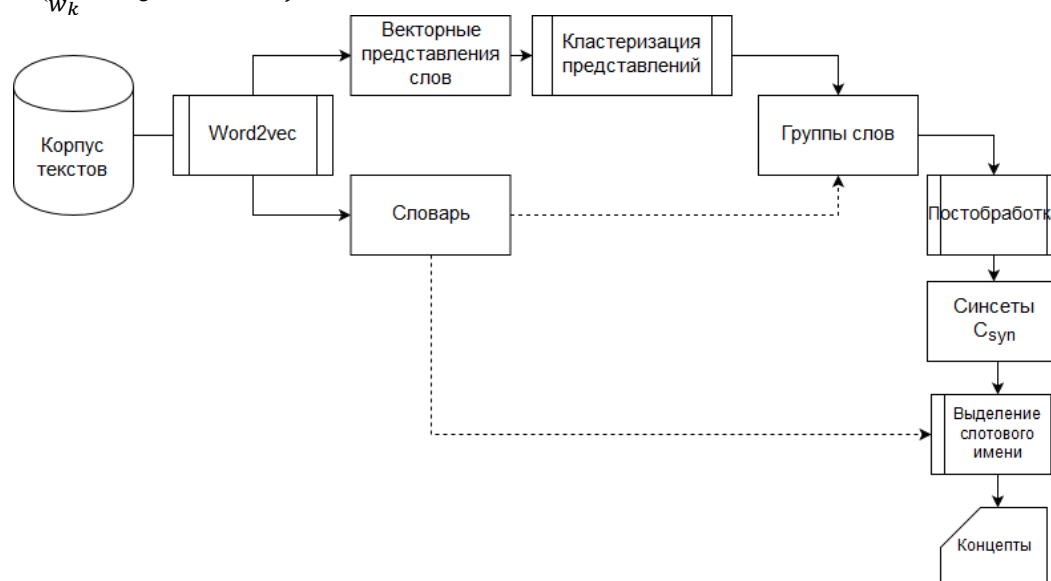


Рисунок. Схема процесса извлечения синонимов и концептов

Недостатком подхода является наличие эвристически задаваемого порога фильтрации. Также этот подход не позволяет выделить имя концепта среди синсета, состоящего из разных лексем без доминанты («сообщение, уведомление, предупреждение»). Тем не менее, описанный метод позволил слотовые имена для 48% концептов (среди синсетов, шумовые кластеры не рассматривались). Общая схема проведенных шагов при применении методов извлечения синонимов и концептов представлена на рисунке.

Для экспериментов использовался корпус неструктурированных текстов пользовательских обращений в текстовый чат службы поддержки крупного российского оператора связи. Сводные данные по экспериментам представлены в табл. 1.

Таблица 1. Информация по экспериментам

Размер корпуса, документов	1031041
Размер словаря, слов	40966
Тип и размерность векторной модели	CBOW, 400
Размер контекстного окна	6
Алгоритмы кластеризации	<i>k</i> -means, fuzzy <i>c</i> -means
Количество кластеров	300
Количество и доля C_{syn} среди всех кластеров	218 (74%)
Доля оставшихся после постобработки слов	0,6322 (2 итерации постобработки)
Выделено слотовых имен среди всех C_{syn}	96 (48%)

Каждый из извлеченных концептов оценивался вручную четырьмя независимыми ассессорами. Введена следующая категоризация: «шум» для кластеров, не являющихся синсетами, «сомнительный» для кластеров, содержащих как синсеты, так и не связанные семантикой слова, «1–2 корня» для синсетов, образованных словоформами 1–2 лексем, а также «концепт-сущность» для синсетов, формирующих концепт и «концепт-тема» для групп слов, описывающих одну тему предметной области. Результаты оценки представлены в табл. 2.

Таблица 2. Результаты оценивания 4 ассессорами

Категория кластера	Доля категории
шум	0,33
сомнительный	0,09
1–2 корня	0,165
концепт-сущность	0,254
концепт-тема	0,164

Полужирным шрифтом выделены доли категорий, которые дают положительный результат эксперимента.

В ходе проведенных экспериментов были разработаны и применены контекстно-зависимые методы извлечения синонимов и концептов из неструктурированных текстов на естественном языке, относящихся к одной предметной области. Разработанные методы являются ресурсно-независимыми, а также не требуют размеченных обучающих данных. Реализация является достаточно гибкой, так как позволяет использовать различные векторные модели и алгоритмы кластеризации, что дает возможность оптимизации методов. Дальнейшими направлениями исследований являются выбор и обоснование используемых эвристик, особенно количества кластеров. Было проведено оценивание, из которого следует, что 58% полученных кластеров либо сформировано из синонимов, либо являются концептами предметной области. На основе описанных методов разработана утилита поиска синонимов заданного слова.

Литература

1. Fellbaum C. (ed.). WordNet: An Electronic Lexical Database. – Cambridge, MA: MIT Press, 1998. – 447 p.
2. Mikolov T. et al. Distributed representations of words and phrases and their compositionality // Advances in neural information processing systems [Электронный ресурс]. – Режим доступа: <https://papers.nips.cc/paper/5021-distributed-representations-of-words-and-phrases-and-their-compositionality.pdf> (дата обращения: 06.01.2019).
3. Winkler R., Klawonn F., Kruse R. Fuzzy c-means in high dimensional spaces // IJFSA. – 2011. – V. 1(1). – P. 1–16.



Черных Герман Анатольевич

Год рождения: 1967

Санкт-Петербургский государственный университет,
факультет свободных искусств и наук, к.ф.-м.н.

e-mail: german.ch@chgerman.com



Латынина Виктория Олеговна

Год рождения: 1996

Санкт-Петербургский государственный университет,
факультет свободных искусств и наук, студент

Направление подготовки: 8.50.03.01 – Искусства и гуманитарные науки

e-mail: latynina.vika@mail.ru

УДК 004.89

СТАТИСТИЧЕСКИЙ АНАЛИЗ ЭМОЦИОНАЛЬНОЙ ДИНАМИКИ ЛИТЕРАТУРНЫХ ТЕКСТОВ НА РУССКОМ ЯЗЫКЕ

Черных Г.А. (Санкт-Петербургский государственный университет),

Латынина В.О. (Санкт-Петербургский государственный университет)

Научный руководитель – к.ф.-м.н. Черных Г.А.

(Санкт-Петербургский государственный университет)

Разработан метод анализа тональной динамики литературных текстов большого объема на русском языке по шести эмоциональным категориям: радость, гнев, печаль, страх, неприязнь и удивление. Метод позволяет проследить изменение корреляций между эмоциональными составляющими текста литературного произведения по мере развития сюжета и построить эмоциональный портрет целого произведения или его части, что может быть использовано для решения различных научно-прикладных задач, касающихся определения тональности.

Ключевые слова: тональность текста, эмоциональный портрет текста, дистрибутивная семантика, Word2Vec, WordNet-Affect, OpenCorpora, динамика корреляций тональностей.

Представленный в настоящей работе метод анализа динамики эмоций в русскоязычных текстах является развитием технологии оценивания тональности слов с использованием контекстной информации, предложенной в [1]. Технология позволяет получить количественные критерии эмоциональной окрашенности отдельных слов или словосочетаний текста, не являющихся непосредственно тональными, но относительно часто употребляемых в контексте последних. При этом слову или словосочетанию может быть сопоставлено одновременно несколько количественных критериев, относящихся к разным эмоциям. Поскольку процентное содержание эмоционально окрашенных слов в тексте становится значительно больше, чем в случае, когда используются только тональные слова, то появляется возможность проследивать динамику изменения эмоций в тексте, что, в свою очередь, позволяет говорить о временных тональных числовых рядах и применять к ним соответствующие методы анализа.

Процедура получения вышеупомянутых временных рядов состоит в следующем (более подробное описание приведено в работе [1]). В качестве источника слов, имеющих явно выраженную эмоциональную окраску, выбрана база WordNet-Affect [2], созданная на основе семантического лексикона английского языка WordNet. База представлена шестью

эмоциональными категориями: радость, страх, гнев, печаль, отвращение, удивление. Поскольку в WordNet-Affect слова содержатся в виде лемм (нормализованных основных форм слов), то возникает необходимость в наличии различных словоформ. Словоформы для лемм из базы WordNet-Affect получены с использованием морфологического словаря открытого корпуса русского языка OpenCorpora [3]. Далее шесть имеющихся групп тональных словоформ используются для обработки массива текстов большого объема (взят источник lib.ru [4]) с целью выделения и группировки в соответствии с принадлежностью к эмоциональным категориям предложений, имеющих в своем составе тональные словоформы. Таким образом, создается шесть текстовых баз (по количеству категорий WordNet-Affect), претендующих на соответствие отдельным типам тональности (рис. 1).



Рис. 1. Схема процедуры получения числовых рядов критериев тональности

Для каждой базы строятся векторные представления посредством технологии Word2Vec [5, 6]. При этом предлоги и устойчивые словосочетания, использующиеся в качестве предлогов, объединяются со следующими за ними по тексту словоформами. Метод вычисления количественного критерия эмоциональной окрашенности слова в зависимости от категории тональности состоит в следующем. С использованием векторных представлений находится n слов, ближайших к заданному слову, в смысле косинусного расстояния. Найденные слова нумеруются в порядке возрастания удаленности от исходного слова. Тогда величина критерия определяется выражением

$$\frac{w}{w_s} \sum_{k=1}^n \frac{\alpha_k}{k},$$

где $\alpha_k = 1$, если слово с номером k принадлежит соответствующей базе тональных словоформ, и $\alpha_k = 0$, в противном случае; w – суммарное количество слов в словаре Word2Vec, построенного по тональной текстовой базе, для которой вычисляется критерий; w_s – количество тональных словоформ соответствующей эмоциональной категории. В таблице приведены значения w и w_s для шести тональных баз и векторных представлений.

Таблица. Значения величин w_s и w для шести различных эмоциональных категорий

	удивление	радость	гнев	страх	печаль	отвращение
w_s	375	2003	868	842	858	225
w	13700	52054	16612	24141	24174	8136

Численное значение параметра n выбирается, исходя из эмпирической оценки максимального количества словоформ, которые могут образовывать группы семантически близких слов. Результаты получены для $n=20$. Слову, не имеющему векторного представления (не представленного в соответствующей Word2Vec-модели), присваивается нулевое значение. Основные этапы вычисления критериев тональности и получения числовых рядов (без учета нормализации текста и работы с предлогами) представлены на схеме рис. 1. Отметим, что формула вычисления критериев тональности модифицирована по сравнению с приведенной в работе [1] и теперь учитывает заметное различие в количестве тональных слов в базах.

По числовым рядам критериев тональности можно строить различные характеристики и использовать последние для решения разнообразных научных и прикладных задач. В качестве примера на рис. 2 приведены диаграммы, содержащие средние значения и дисперсии критериев тональности, а также усредненные коэффициенты корреляции рядов для двух произведений: романов Федор Достоевского «Преступление и наказание» и Ильи Ильфа и Евгения Петрова «12 стульев». Коэффициенты корреляции вычислялись по скользящему интервалу шириной 50 слов. Представленные диаграммы можно рассматривать как частный случай эмоционального портрета текста, где могут быть отражены различные глобальные характеристики произведения. Любопытно, что на представленном на рис. 2 примере у романа Достоевского в среднем выше уровень эмоциональной составляющей, включая дисперсионный разброс по сравнению с произведением Ильфа и Петрова, тогда как у последних выше уровень корреляций между рядами эмоциональных критериев.

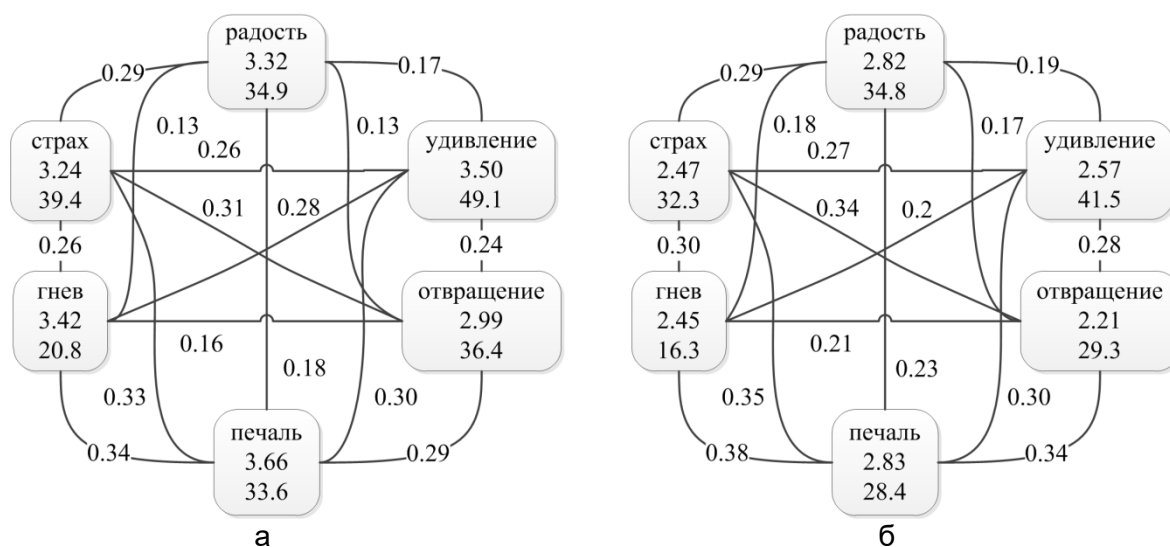


Рис. 2. Средние значения и дисперсии (указаны в узлах графов) и средние коэффициенты корреляции (указаны на дугах графов) для рядов критериев тональности романов «Преступление и наказание» (а) и «12 стульев» (б)

Необходимо отметить, что доступных текстовых баз, размеченных по шести эмоциональным категориям, не существует, поэтому объективная оценка качества предложенного метода в настоящее время затруднена. Однако ряд косвенных факторов, в частности, стабильность полученных данных при варьировании параметров или существенная разница в поведении критериев для субъективно разных по эмоциональной окрашенности произведений, свидетельствуют в пользу работоспособности алгоритма. Кроме того, метод уже позволяет решать задачу кластеризации текстов или их частей, а также может быть использован для ряда других прикладных задач анализа текстовых данных.

Литература

1. Черных Г.А. Оценка тональности текста на основе технологии Word2Vec // Современные проблемы физико-математических наук. Материалы III Международной научно-практической конференции. – 2017. – С. 369–373.
2. Лаборатория обработки естественного языка [Электронный ресурс]. – Режим доступа: http://lilu.fcim.utm.md/resourcesRoRuWNA_ru.html (дата обращения: 06.01.2019).
3. Открытый корпус русского языка [Электронный ресурс]. – Режим доступа: <http://www.opencorpora.org> (дата обращения: 06.01.2019).
4. Библиотека Максима Мошкова [Электронный ресурс]. – Режим доступа: <http://lib.ru> (дата обращения: 06.01.2019).
5. Mikolov T. and etc. Efficient Estimation of Word Representations in Vector Space [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1301.3781.pdf> (дата обращения: 06.01.2019).
6. Mikolov T. and etc. Distributed Representations of Words and Phrases and their Compositionality [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1310.4546.pdf> (дата обращения: 06.01.2019).

**НАПРАВЛЕНИЕ
ИНФОКОММУНИКАЦИОННЫЕ ТЕХНОЛОГИИ**

**Алексеев Николай Борисович**

Год рождения: 1996

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K41201Направление подготовки: 11.04.02 – Программное обеспечение
в инфокоммуникациях

e-mail: nklbrsch@yandex.ru

**Иванов Сергей Евгеньевич**

Год рождения: 1975

Университет ИТМО, факультет инфокоммуникационных технологий,
к.ф.-м.н., доцент

e-mail: 45is@mail.ru

УДК 004.738.5**ТЕХНОЛОГИИ И МЕТОДЫ ПОИСКОВОЙ ОПТИМИЗАЦИИ****Алексеев Н.Б.****Научный руководитель – к.ф.-м.н., доцент Иванов С.Е.**

Работа посвящена исследованию способов поднятия сайта в результатах выдачи поисковых систем и увеличению рейтинга посещаемости при помощи методов SEO. Объектами исследования работы являлись методы поисковых систем, оптимизирующих работу сайта.

Ключевые слова: поисковая оптимизация, SEO, продвижение сайта.

Поисковая оптимизация или SEO (Search Engine Optimization) – это действия, направленные на оптимизацию веб-страниц или целых сайтов с целью добиться авторитетности у поисковой системы, таким образом, поднимая позиции в результатах поиска.

Вначале стоит понять, как устроен процесс работы поисковой системы. Его можно разделить на пять этапов:

1. сканирование – процесс сбора новой информации и ее анализ, для определения нового контента возможного для предоставления пользователю. Для сканирования поисковые системы используют специальных роботов, называемых поисковыми роботами или пауками;
2. индексирование – процесс занесения на новых сайтах и содержащийся на них информации в базы данных поисковой системы;
3. обработка – сравнение строки поиска в поисковом запросе с индексированными страницами в базе данных;
4. ранжирование – выстраивание информации, занесенной в базу данных, по рангу, т.е. какую информацию поисковик будет показывать своим пользователям в первую очередь, а какую информацию помещать «рангом» ниже;
5. получение результатов – последний этап работы поисковой системы, извлекающей наиболее подходящие результаты, и их отображение в браузере.

Для решения задачи ранжирования поисковые системы учитывают довольно много свойств текста страницы и всего сайта – факторов оптимизации, которые можно условно разделить на две группы: внутренняя и внешняя [1, С. 173].

Внутренняя оптимизация – кропотливая и долгая работа над сайтом, подразумевающая улучшение контента и структуры, подбор и правильное размещение ключевых слов, изменение заголовков, соответствующих каждой странице сайта, и многое другое.

Внешняя оптимизация – увеличение входящего трафика путем ссылочного продвижения, продвижения в социальных сетях, написания статей и т.п.

Цель внутренней и внешней оптимизации любого сайта одна – повысить позиции сайта в результатах выдачи поисковых систем. А вот методы достижения этой цели бывают разные, так называемые: белая, серая и черная оптимизации, или как их еще называют white hat, grey hat и black hat.

Белая SEO. Белая оптимизация – метод продвижения сайта естественным путем, с учетом требований поисковых систем и их алгоритмов, т.е. белая SEO подразумевает работу над сайтом и его внутренней оптимизацией, в соответствии с рекомендациями поисковых систем. При использовании данного метода не следует никаких санкций со стороны поисковых систем по отношению к сайту. За счет чего данный метод является наиболее безопасным способом увеличения трафика на ресурсе [2, С. 8].

SEO можно считать белой, если она соответствует следующим пунктам:

- она соответствует руководящим принципам поисковой системы;
- она не связана с каким-либо обманом;
- она гарантирует, что контент на веб-странице создан для пользователя, а не для поисковой системы;
- она обеспечивает высокое качество веб-страниц;
- она обеспечивает доступность полезного контента на веб-страницах.

Черная SEO. Черная оптимизация подразумевает продвижение сайта за счет обмана поисковых систем и нахождения слабых мест в их алгоритмах. Методы черной оптимизации обычно ориентированы на поисковые боты, а не на пользователя. Такая стратегия несет в себе высокий риск быть заблокированным поисковой системой.

К наиболее распространенным способам относят: клоакинг, создание фальшивых сайтов – дорвеев, свопинг, скрытый текст или ссылки и разные виды ссылочного спама [3, С. 231]:

- клоакинг (Cloaking) – контент на странице для пользователя отличается от контента на странице для поискового робота. Так для поискового робота можно отобразить страницу с большим количеством ключевых слов, а пользователю – просто изображение или Flash-приложение;
- создание дорвеев (Doorway) – автоматическое перенаправление пользователя с одного ресурса на другой. Так создаются страницы, заточенные под определенные ключевые слова, после перехода на которые пользователь перенаправляется на другой ресурс;
- свопинг – смена контента после достижения страницей высоких позиций;
- скрытый текст или ссылки – мелкий шрифт или одинаковый с фоном цвет текста, который видит поисковой робот, но не видит пользователь.

Серая SEO. Серая SEO – оптимизация, находящаяся на границе разрешенных и запрещенных способов, но являющихся разрешенными.

К серой SEO можно отнести такие методы как:

- покупка ссылок у авторитетных партнеров;
- покупка подписчиков в социальных сетях;
- добавление в текст контента большого количества ключевых фраз;
- создание дорвеев без автоматического перенаправления – основное отличие от метода черной оптимизации в том, что здесь дорвей содержит актуальную для пользователя информацию.

Разобравшись в принципах и методах поисковой оптимизации, остается понять какими технологиями и инструментами можно достичь желаемого результата. У основных в России поисковых систем Google и Yandex [4] существуют свои инструменты для сбора статистики и анализа необходимой для SEO, основные функциональные возможности которых примерно одинаковы, поэтому разберем инструменты, предоставленные поисковой системой Google [5].

Google Webmaster Tool – данный инструмент помогает веб-мастерам лучше контролировать взаимодействие Google с веб-сайтом и получать полезную информацию. Так, например, в панели можно посмотреть внутренние и внешние ссылки на сайт, узнать статус и посмотреть ошибки индексации, узнать поисковые запросы, приносящие трафик, просмотреть ошибки индексации, а также иметь ряд других показателей, которые позволяют улучшить сайт:

- Google Analytics Tool – инструмент для отслеживания трафика и его анализа. Google Analytics дает глубокое представление о таких вещах как: Как пользователи приходят и ведут себя на сайте; Какой контент является наиболее популярен; Помогает оценить влияние сделанной оптимизации на сайт;
- Google Optimiz – инструмент для тестирования сайта и его оптимизации, позволяющий сравнивать эффективность различных вариантов страниц или экранов приложения, показывая их пользователям, которые выбираются случайным образом;
- Google Trends – инструмент, позволяющий оценить статистику запросов по определенным ключевым словам или фразам за весь период времени, оценить актуальность тех или иных тем, а также, в каких географических регионах люди искали этот термин больше всего.

Для того чтобы добиться высоких результатов в сети, необходимо прибегать к технологиям поисковой оптимизации сайта, поскольку на сегодняшний день поисковые системы являются одним из основных источников всего трафика. Методы внутренней и внешней оптимизации равнозначно важны для улучшения позиции в поисковой выдаче, в то время как из трех технологий SEO, белая оптимизация является наилучшим выбором, поскольку является наиболее долгосрочной и безопасной, но менее эффективной по сравнению с черной и серой. Алгоритмы ранжирования регулярно обновляются, что затрудняет работу оптимизаторов, но, в то же время, делая поисковые службы наиболее удобными для пользователей, такая тенденция позволяет сфере интернет-маркетинга не стоять на месте и стабильно развиваться.

Литература

1. Ашманов И., Иванов А. Оптимизация и продвижение сайтов в поисковых системах. – 3-е изд. – СПб.: Питер, 2011. – 464 с.
2. Иванов И. SEO: Поисковая оптимизация от А до Я. – Интернет-издание, 2012. – 537 с.
3. Яковлев А., Ткачев В. Раскрутка сайтов: основы, секреты, трюки. – Изд-во: БХВ-Петербург, 2015. – 352 с.
4. Рейтинг поисковых систем на 2018 год [Электронный ресурс]. – Режим доступа: <http://gs.seo-auditor.com.ru/sep/2018/> (дата обращения: 15.01.2019).
5. Севостьянов И.О. Поисковая оптимизация. Практическое руководство по продвижению сайта в Интернете. – Изд-во: Питер, 2016. – 272 с.



Амосова Анастасия Федоровна

Год рождения: 1984

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № К4220

Направление подготовки: 11.04.02 – Программное обеспечение
в инфокоммуникациях

e-mail: amosova@outlook.com



Осетрова Ирина Станиславовна

Год рождения: 1961

Университет ИТМО, факультет инфокоммуникационных технологий,
ст. преподаватель

e-mail: irina@ifmo.spb.ru

УДК 004.652

ОБЗОР СОВРЕМЕННЫХ ПЛАТФОРМ ДЛЯ РАЗРАБОТКИ СИСТЕМ БИЗНЕС-АНАЛИТИКИ

Амосова А.Ф.

Научный руководитель – Осетрова И.С.

В работе рассмотрена структура платформ бизнес-аналитики, их основные возможности и особенности, тенденции дальнейшего развития и перспективные направления в указанной области.

Ключевые слова: системы бизнес-аналитики, хранилища данных, Teradata, Microsoft SQL Server.

Целью работы являлось рассмотрение возможностей современных платформ для разработки решений в сфере бизнес-аналитики и перспективных направлений их развития. Существует достаточно большое число таких платформ. Для данного обзора были выбраны платформы крупнейших разработчиков: Microsoft и Teradata, – поскольку их функционал наиболее универсален и полон, а построенные на их основе системы Business Intelligence (BI) в настоящее время с успехом внедряются не только в крупных корпорациях, но и в сравнительно небольших компаниях.

В целом, BI – это совокупность технологий, программного обеспечения и методов, направленных на достижение бизнес-целей посредством наиболее эффективного применения имеющихся данных [1]. Функционал BI-систем можно разделить на четыре основных направления:

1. хранение данных;
2. интеграция данных;
3. анализ данных;
4. представление данных.

Данные, используемые для бизнес-анализа, организуются в специальные хранилища – предметно-ориентированные, интегрированные, стабильные, поддерживающие хронологию наборов данных, используемые для поддержки принятия управленческих решений. Основными требованиями к хранилищам являются:

- обеспечение консистентности данных;
- обеспечение надлежащего качества данных;
- управление жизненным циклом данных;
- высокая скорость работы хранилища данных.

Интеграция подразумевает наличие BI-инфраструктуры: в рамках платформы должны использоваться общие данные и единая объектная модель. Кроме того, необходим набор инструментов для создания BI-приложений, а также средства для обмена информацией внутри команды, такие как, например, корпоративные порталы, позволяющие назначать задания и отслеживать их выполнение.

Для всестороннего анализа данных в современных BI-системах используются оперативная аналитическая обработка данных – так называемые OLAP-кубы, предиктивное моделирование и интеллектуальный анализ данных, интерактивные графики и схемы, а также карты ключевых показателей эффективности и, конечно же, встроенная система поддержки принятия решений.

Представление информации должно осуществляться как посредством отчетов, форматированных или интерактивных, так и с помощью информационных панелей. Необходима поддержка стандартного пакета Microsoft Office, а также мобильных устройств. Кроме того, система должна иметь встроенный поисковик, поддерживающий, нетиповые уникальные запросы.

К системам хранения данных в бизнес-аналитике предъявляются особые требования. Поскольку объемы данных, генерируемых, обрабатываемых и помещаемых в хранилища, растут стремительно и неуклонно, система хранения данных остается наиболее затратным компонентом BI-решений.

Роль хранилищ данных в BI-проектах велика не только с учетом финансовых соображений. Корпоративное хранилище является поистине фундаментальным элементом BI-системы, так как именно оно отвечает за бесперебойную поставку данных приложениям бизнес-аналитики из различных источников.

В платформе данных Microsoft в роли пространства для интеграции данных выступает SQL Server, обеспечивающий высокую доступность и продвинутую аналитику как структурированных, так и неструктурированных данных, включая большие данные.

SQL Server 2019 представляет собой целостную платформу для искусственного интеллекта и машинного обучения, поскольку содержит основные элементы озера данных – распределенную файловую систему Hadoop, Spark и инструменты анализа [2]. Кластеры больших данных SQL Server включают встроенные службы гибридного развертывания и управления.

Сервер отчетов Power BI позволяет создавать интерактивные отчеты, а также использовать возможности формирования отчетов SQL Server Reporting Services.

PolyBase используется для виртуализации данных, которая является альтернативой ETL и объединяет данные из разных источников без их репликации или перемещения, формируя единый «виртуальный» уровень данных – дата-хаб. Он позволяет пользователям запрашивать данные из многих источников через единый унифицированный интерфейс.

Облачная часть платформы Microsoft построена на базе инфраструктуры Azure и включает в себя весь спектр необходимых компонентов: виртуальные машины, облачные хранилища, OLAP-службы, озера и фабрики данных и всевозможные инструменты для разработки AI и BI.

Teradata – американская корпорация, специализирующаяся на разработке и поставке аппаратно-программных комплексов для обработки и анализа данных.

Teradata UBC объединяет интегрированное хранилище данных, систему многофункционального анализа данных Aster Analytics и технологию Hadoop в единое решение, а QueryGrid предоставляет пользователям непосредственный доступ к данным и процессам аналитики в различных системах. Данное решение сводит к минимуму необходимость в перемещении и копировании, позволяя обрабатывать данные там, где они хранятся.

Возможности аналитической платформы Teradata доступны даже в небольшом центре обработки данных благодаря Teradata UDA Appliance – полностью конфигурируемой экосистеме, где все программное обеспечение размещается в одном блоке.

В прошлом году компания Teradata представила новую платформу Vantage, которая осуществляет доступ к хранилищу консолидированных данных с помощью исполнительного механизма Teradata SQL Engine, предоставляющего широкий ассортимент встроенных инструментов бизнес-аналитики.

В состав Vantage входит исполнительный механизм машинного обучения, с помощью которого исследователи данных могут создавать выразительные и мощные алгоритмы машинного обучения [3]. Teradata также предлагает более 180 готовых BI-инструментов.

Vantage предоставляет и встроенный исполнительный механизм для обработки графики, позволяющий анализировать графические изображения, идентифицируя и измеряя отношения между людьми, продуктами и процессами.

В заключение хотелось бы подчеркнуть, что ключевым элементом при построении BI-систем является именно платформа по управлению данными, поэтому к ее выбору необходимо подходить взвешенно. Если верить мировой статистике, около половины проектов хранилищ данных завершаются провалом, поскольку их создатели ориентировались не на запросы бизнес-пользователей, а на доступные виды данных [4].

Следует учесть и программное окружение проекта для обеспечения бесперебойной и, по возможности, бесшовной интеграции решения с имеющимися информационными системами.

В целом, функционал, предлагаемый различными платформами, схож, однако, его исполнение существенно различается. Так, в крупной корпорации, оперирующей поистине «большими» данными, предпочтительным представляется использование платформы Teradata, так как последняя является программно-аппаратным комплексом на базе массово-параллельной архитектуры и выигрывает по производительности у сугубо программных систем. Кроме того, крупная корпорация может позволить себе организацию обучения сотрудников для повышения эффективности использования указанной платформы. Для небольшой организации платформа Microsoft выглядит привлекательнее как за счет меньших вложений в аппаратную часть, так и в связи с большим количеством обучающих материалов, предоставляемых Корпорацией Microsoft бесплатно.

Литература

1. Ronthal A., Edjlali R., Greenwald R. Magic Quadrant for Data Management Solutions for Analytics // Gartner. – 2018. – 76 p.
2. Microsoft SQL Server 2019. Technical white paper // Microsoft Corporation. – 2018. – 24 p.
3. Choose the Analytic Solution to Meet Business Requirements. White Paper // Teradata Corporation. – 2018. – 8 p.
4. Logan V., Herschel G., Schlegel K., Sicular S., Ronthal A., Hare J., Davis M., Heizenberg J. Harnessing the Pervasive Nature of Domain Data and Analytics // Gartner. – 2018. – 13 p.

**Бойко Наринэ Карапетовна**

Год рождения: 1995

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K4232Направление подготовки: 27.04.03 – Системный анализ

и интеллектуальное управление в бизнесе

e-mail: narinearzumanyan500@gmail.com

Панова Анастасия Юрьевна

Год рождения: 1965

Университет ИТМО, факультет инфокоммуникационных технологий, к.э.н., доцент

УДК 004.04**ПРИМЕНЕНИЕ VI-СИСТЕМЫ ДЛЯ АВТОМАТИЗАЦИИ ФОРМИРОВАНИЯ
ОПЕРАТИВНОЙ ЭКОНОМИЧЕСКОЙ ОТЧЕТНОСТИ****Бойко Н.К.****Научный руководитель – к.э.н., доцент Панова А.Ю.**

В работе рассмотрено применение VI-системы для автоматизации формирования оперативной управленческой отчетности компании нефтегазовой отрасли ПАО «Газпром нефть». В связи с тем, что оперативная отчетность должна характеризоваться высокой скоростью сбора и обработки данных для отражения текущего положения компании, особенно важной является задача автоматизации формирования этой отчетности. На ее сбор вручную тратится много ресурсов и времени сотрудников компании, а для принятия эффективных управленческих решений необходимо иметь возможность постоянного мониторинга показателей, что и позволит сделать внедрение соответствующей VI-системы.

Ключевые слова: VI-системы, Business Intelligence, аналитические системы, управленческая отчетность, автоматизация формирования отчетности.

Для контроля ведения бизнеса, анализа эффективности деятельности предприятия и принятия управленческих решений, безусловно, необходима отчетность. По законодательству Российской Федерации и международным стандартам существует много обязательных форм отчетности, часть из которых акционерные общества предоставляют в открытом доступе. Для внутреннего использования существуют также и другие формы отчетности, в частности, формы оперативной отчетности, автоматизация которой была рассмотрена в данной работе [1, 2].

Оперативная отчетность должна отражать ключевые показатели ведения бизнеса и характеризоваться скоростью сбора. Иначе говоря, фактически она должна отражать текущее состояния в режиме «онлайн». В крупных компаниях с большим количеством дочерних обществ (ДО), точек продаж и других объектов управления, невозможно добиться сбора информации в таком режиме, в связи с чем делается послабление на время передачи информации – оно должно быть минимально возможным от факта совершения операции до появления информации в отчетности.

Внутри каждой компании самостоятельно определяются ключевые показатели оперативной управленческой отчетности, но, безусловно, есть показатели, характерные для отрасли в целом. В качестве примера рассматривается компания нефтегазовой отрасли ПАО «Газпром нефть» и анализ деятельности сбытового блока компании и его ДО.

Разные ДО ведут учет хозяйственной деятельности в разных учетных системах, что затрудняет формирование отчетности из-за необходимости приведения к единому виду и объединения данных. В связи с этим для полноты картины, упрощения контроля и анализа данных информацию о продажах необходимо консолидировать в одну систему. Для решения

данного класса задач существуют аналитические системы, авторами рассмотрены преимущества от внедрения BI-системы (Business Intelligence system).

Термин «Business Intelligence» (BI) имеет много определений, в зависимости от контекста и глубины погружения ему дают различные трактовки. Исходно BI считался инструментом для анализа информации, более глубокого понимания данных с целью дальнейшего более точного и эффективного принятия решений бизнес-пользователями, а также был средством выявления в большом объеме данных ценной для бизнеса информации. В данной работе наиболее подходящим представлена трактовка в виде процесса трансформации данных в информацию и ценные знания о бизнесе для дальнейшего облегчения принятия решений, а также сами методы и средства сбора данных, объединения информации и обеспечения доступа бизнес-пользователей к знаниям.

Компания Gartner, предоставляющая исследовательские и консалтинговые услуги на рынке информационных технологий, ежегодно анализирует рынок BI-систем. В феврале 2018 они опубликовали отчет в виде Магического Квадранта Gartner для платформ Аналитики и Бизнес Аналитики [3], в котором представили основных вендоров BI. Квадрат состоит из 4 секторов – leaders (лидеры), challengers (претенденты в лидеры), visionaries (провидцы – дальновидные производители программного обеспечения), niche players (нишевые игроки – отстающие в отрасли). К лидерам в 2018 аналитики Гартнера отнесли Microsoft, Tableau, Qlik. Основными продуктами Qlik являются QlikSense и QlikView, последний из которых и был выбран для разработки BI-приложения с целью автоматизации формирования оперативной управленческой отчетности. Ключевым преимуществом QlikView является возможность реализации сложных преобразований данных из большого числа источников для последующего централизованного использования BI бизнес-пользователями, что и требовалось в рассматриваемой задаче.

Аналитические системы осуществляют многомерный анализ бизнес-данных по необходимому срезу. Для ПАО «Газпром нефть» при анализе продаж можно выделить следующие ключевые аналитики: период, сценарий, организация, регион, канал реализации нефтепродуктов. В качестве вспомогательных (углубленных) аналитик можно выделить: номенклатурные группы, подразделения, контрагенты. Результатом работы системы являются достоверные и непротиворечивые данные, но с учетом специфики оперативной отчетности (режим, близкий к режиму «онлайн»), допускаются небольшие расхождения в данных, связанные в основном с человеческим фактором.

Преимуществами BI-системы являются возможность объединения информации из большого числа системных и несистемных источников, предварительная обработка полученных данных и их приведение к единым метрикам. Таким образом, может быть произведена унификация справочной информации, что позволяет полноценно сравнивать данные различных систем, корректно объединять их для последующего анализа. Также имеется возможность произвести предрасчет данных в BI-системе без внесения каких-либо изменений в системы-источники. С точки зрения разработки аналитических приложений существенно дешевле производить преобразование данных на стороне аналитической системы, чем во всех системах-источниках.

В крупных проектах количество систем-источников измеряется десятками, а до внедрения аналитической системы в основном происходит ручное формирование отчетности в Excel. Для формирования оперативной управленческой отчетности сбытового блока ПАО «Газпром нефть» используется множество систем, но часть данных все же формируется в Excel. В связи с этим требуются колоссальные трудозатраты, а также существует дополнительный риск наличия ошибок в виду человеческого фактора.

Визуализация, которую возможно разработать в BI-системах, позволяет наглядно представлять информацию, акцентировать внимание на проблемах и достижениях. Часто используют так называемые «светофоры»: например, если оперативные показатели больше, чем эти же показатели прошлых периодов, или средних значений, или неких плановых

показателях, то индикатор окрашивается в зеленый цвет, меньше – в красный. Причем все эти показатели измеряются в нужной детализации, добиться такого с использованием средств Excel невозможно в силу очень большого объема данных [4, 5].

Таким образом, создание аналитического приложения для формирования оперативной информации о продажах нефтепродуктов дочерними обществами ПАО «Газпром нефть» оправдано. Разработка и дальнейшее использование BI-приложения сократит трудозатраты по формированию отчетности, увеличит скорость сбора информации, ее точность, а также позволит производить анализ данных в углубленной детализации и наглядно представлять основные показатели для прозрачности деятельности компании и оперативного выявления проблем с целью их последующего решения.

Литература

1. Духонин Е.Ю., Исаев Д.В., Мостовой Е.Л. и др. Управление эффективностью бизнеса. Концепция Business Performance Management / Под. ред. Г.В. Генса. – М.: Альпина Бизнес Букс, 2005. – 269 с.
2. Kisielnicki J., Misiak A.M. Effectiveness of Agile Implementation Methods in Business Intelligence Projects from an End-user Perspective [Электронный ресурс]. – Режим доступа: <http://inform.nu/Articles/Vol19/ISJv19p161-172Kisielnicki2706.pdf> (дата обращения: 06.01.2019).
3. Магический Квадрант Gartner для платформ Аналитики и Бизнес Аналитики [Электронный ресурс]. – Режим доступа: <https://fbconsult.ru/qlik-nazvan-liderom-gartner-magic-quadrant-2018> (дата обращения: 06.01.2019).
4. Abai N.H.Z., Yahaya J., Deraman A. Business Intelligence and Analytics in Managing Organizational Performance: The Requirement Analysis Model // Journal of Advances in Information Technology. – 2016. – V. 7. – № 3. – P. 208–213.
5. Ong L., Siew P.H. An empirical analysis on business intelligence maturity in Malaysian organizations // International J. Inf. Syst. Eng. – 2013. – V. 1. – № 1. – P. 1–10.

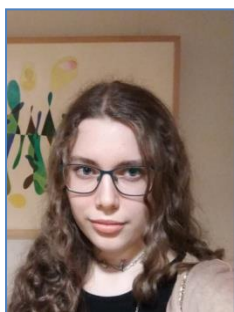


Говоров Антон Игоревич

Год рождения: 1991

Университет ИТМО, факультет инфокоммуникационных технологий,
ассистент

e-mail: antongovorov@gmail.com



Слизень Елена Витальевна

Год рождения: 1996

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K4140

Направление подготовки: 09.04.02 – Информационные системы
и технологии

e-mail: helen.slizen@gmail.com

УДК 004.09

ПРОБЛЕМАТИКА СИНТАКСИЧЕСКОГО АНАЛИЗА SQL-КОДА

Говоров А.И., Слизень Е.В.

Научный руководитель – Говоров А.И.

В работе рассмотрен процесс анализа SQL-кода в рамках оптимизации SQL-запросов со стороны системы управления базами данных. В работе представлен порядок обработки SQL-запросов и плана их оптимизации, рассмотрен принцип синтаксического анализа SQL-кода и построения дерева разбора. Обоснована важность грамотного обучения составлению оптимизированных запросов для последующей квалифицированной деятельности сотрудников и необходимость разработки методов определения оптимальности написания SQL-запросов.

Ключевые слова: SQL-запросы, базы данных, СУБД, оптимизация SQL-запросов.

Проблема производительности системы управления базами данных (СУБД) на данный момент остается актуальной, и встает вопрос о настройке оптимизации. В настоящее время под термином «оптимизация» в СУБД часто понимается именно аспект оптимизации SQL-запросов [1]. Под оптимизацией SQL-запросов, согласно Сорокину В.Е., подразумевается «способ выполнения запросов по процедурному плану, наиболее оптимальному при существующих в БД управляющих структурах и распределении данных» [2, С. 48]. С целью описания процесса оптимизации запросов в СУБД представлен процесс обработки запросов в рамках их выполнения в СУБД (рис. 1) [3].

Куклин А. в статье [4] рассматривает, как происходит выполнение SQL-запросов сервером Microsoft SQL с помощью процессора запросов. Автор выделяет его два основных компонента:

1. оптимизатор запросов;
2. исполнитель запросов.

На примере оператора SELECT автор перечисляет следующие шаги, которые выполняются при обработке инструкции данного оператора с использованием компонентов оптимизатора запросов [4]:

1. разбиение инструкции на логические единицы (ключевые слова, выражения, операторы) с помощью синтаксического анализатора;
2. семантический анализ: проверка существования объектов базы данных, указанных в запросе, корректность использования операторов и выражений;

3. построение дерева разбора запроса с описанием необходимых для преобразования исходных данных к желаемому результату логических шагов;
4. анализ способов, дающих возможность обратиться к исходным таблицам, и выбор оптимизатором ряда шагов, с помощью которых результат вернется с наименьшими временными и ресурсными затратами.

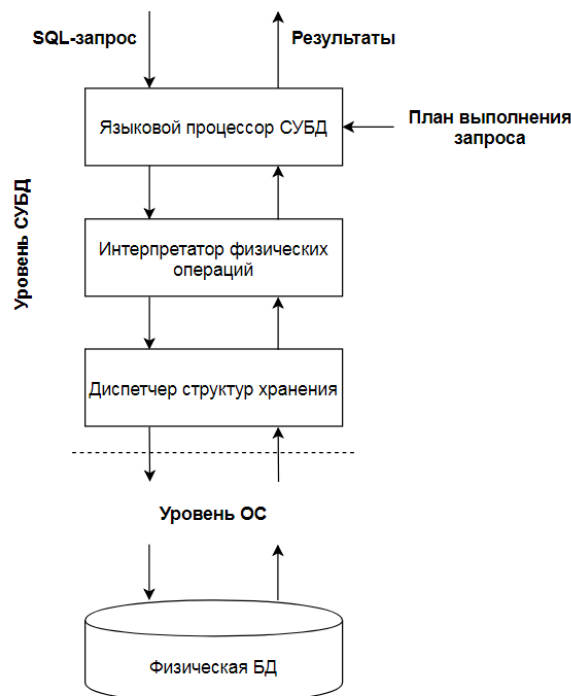


Рис. 1. Структурная схема обработки SQL-запроса

В реляционных СУБД оптимальным планом выполнения запроса можно считать «последовательность применения операторов реляционной алгебры к исходным и промежуточным отношениям, которое для конкретного текущего состояния базы данных может быть выполнено с минимальным использованием вычислительных ресурсов» [2, С. 49]. Отсюда следует, что основной идеей оптимизации запросов является сокращение числа логических операций чтения, которое позволяет уменьшить время выполнения запроса [1].

Общая схема плана оптимизации запросов представлена на рис. 2 [5].

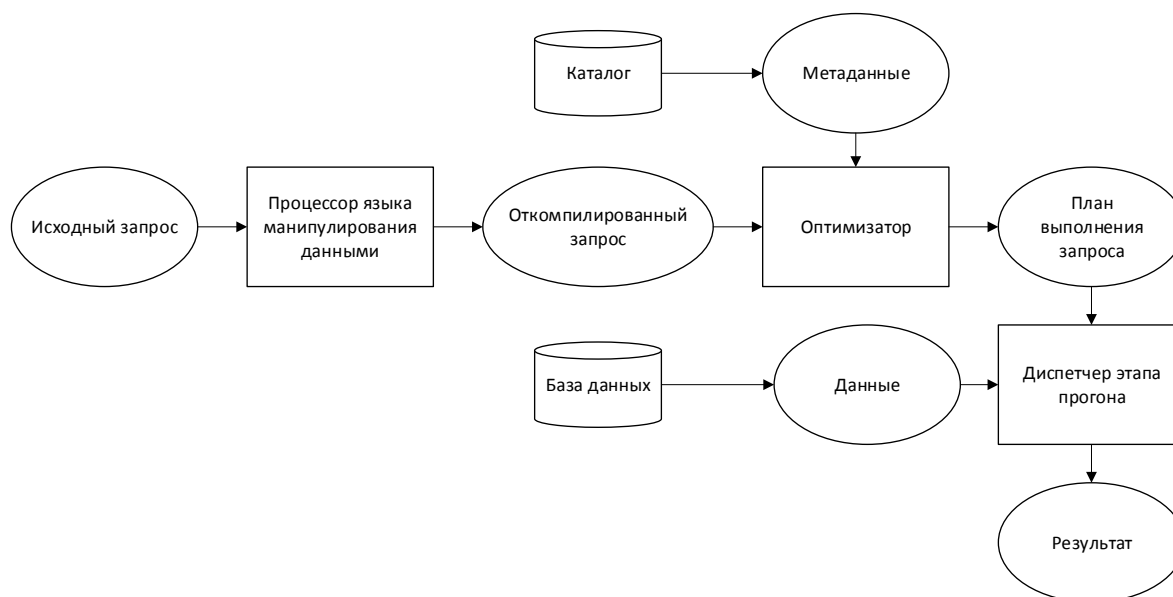


Рис. 2. Общая схема плана оптимизации запросов

План оптимизации можно разделить на этапы следующим образом:

1. лексический и синтаксический анализ, во время которого вырабатывается внутреннее представление запроса, отражающее его структуру и содержащее информацию, которая характеризует объекты базы данных, упомянутые в запросе (отношения, поля, константы);
2. логическая и семантическая оптимизация, во время которой могут применяться различные преобразования начального представления запроса (эквивалентные или семантические);
3. выбор плана выполнения запроса на основе информации, которой располагает оптимизатор (статистическая информация о состоянии базы данных или информация о механизмах реализации различных путей доступа) [6];
4. генерация процедурного плана выполнения запроса.

Таким образом, можно отметить, что синтаксический анализ запроса является составляющей первой фазы, и поэтому он был выбран для начального этапа исследования. Схема процесса синтаксического анализа представлена на рис. 3 [7].

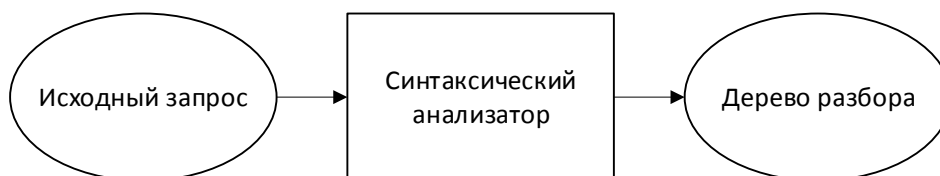


Рис. 3. Схема процесса синтаксического анализа

Дерево разбора, являющееся результатом синтаксического анализа SQL-запроса, содержит узлы двух типов: атомы и синтаксические категории. Атомами являются лексические элементы: ключевые слова, имена атрибутов или отношений, константы, скобки, операторы и др. Синтаксические категории – это имена семейств, представляющих часть запроса. К базисным синтаксическим категориям можно отнести атрибут (attribute), отношение (relation) и паттерн (pattern). К базисным категориям также относится <SFW>, являющееся сокращением для определения запроса типа «SELECT-FROM-WHERE».

Пример построения дерева разбора при синтаксическом анализе для запроса вида:

```

SELECT <SelList1>
FROM <FromList1>
WHERE
<Attribute1> = <Pattern1>;
    
```

отражен на рис. 4 [7].

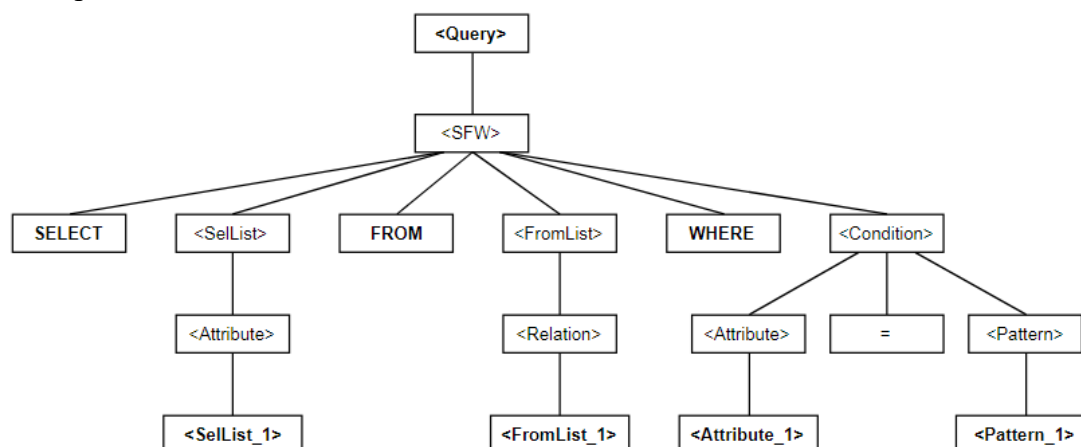


Рис. 4. Пример построения дерева разбора

На основе построенного дерева разбора осуществляется логический план выполнения запроса, представляющий собой дерево, в узлах которого стоят реляционные операции, а

листьями являются отношения. Логический план однозначно соответствует выражению реляционной алгебры.

Так как SQL является достаточно гибким языком программирования, один и тот же результат можно получить различными способами, что влияет на построение дерева разбора и, как следствие, на производительность СУБД. Соответственно, возможны ситуации, когда возвращающие одни и те же данные запросы, написанные по-разному, выполняются за разные промежутки времени. Это означает, что «эффективность выполнения SQL-запроса определяется, в том числе и тем, насколько качественно он написан» [3, С. 277].

Средства СУБД не регулируют степень оптимальности написанного пользователем запроса, при этом оптимизация запросов является важным навыком для «разработчиков SQL и администраторов баз данных» [6, С. 1], который позволяет специалистам понимать работу оптимизатора запросов и используемые им методы, чтобы выбрать надлежащий путь доступа и подготовить соответствующий план выполнения запроса [8].

Следовательно, актуальными задачами являются качественное обучение сотрудников написанию оптимальных SQL-запросов и разработка методов определения оптимальности пользовательских SQL-запросов.

Литература

1. Беляков С.Л., Диденко Д.А. Оптимизация запросов к кадастровой базе данных // Известия ЮФУ. Технические науки. – 2008. – № 10(87). – С. 178–182.
2. Сорокин В.Е. Метод искусственного соответствия SQL-запросов индексам реляционных баз данных // Программные продукты и системы. – 2013. – № 2. – С. 47–54.
3. Шустова Л.И., Тараканов О.В. Базы данных. – М: Инфра-М, 2016. – 304 с.
4. Куклин А. Кэш планов и параметризация запросов. Часть 1. Анализ кэша планов. [Электронный ресурс]. – Режим доступа: <https://club.directum.ru/post/1125> (дата обращения: 15.01.2019).
5. Карпова И.П. Основы базы данных. Учебное пособие. – Изд-во: Московский Государственный институт электроники и математики, 2007. – 75 с.
6. Grant Fritchey. SQL Server Execution Plans. – Second edition. – SQL Handbooks, 2012. – 322 p.
7. Соколинский Л.Б. Обработка запросов в системах баз данных [Электронный ресурс]. – Режим доступа: <https://sok.susu.ru/courses/QueryProc/lectures/Slidesrus/02%20Query%20analysis.pdf> (дата обращения: 14.01.2019).
8. Kumari N. SQL Server query optimization techniques - tips for writing efficient and faster queries // International Journal of Scientific and Research Publications. – 2012. – V. 2. – № 6. – P. 1–4.



Гордиенко Антон Александрович

Год рождения: 1995

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K4220

Направление подготовки: 11.04.02 – Программное обеспечение
в инфокоммуникациях

e-mail: yukiko.kun@yandex.ru



Иванов Сергей Евгеньевич

Год рождения: 1975

Университет ИТМО, факультет инфокоммуникационных технологий,
к.ф.-м.н., доцент

e-mail: 45is@mail.ru

УДК 004.428.4

АНАЛИЗ СОВРЕМЕННЫХ ДОПОЛНЕНИЙ И РАСШИРЕНИЙ В БРАУЗЕРАХ

Гордиенко А.А.

Научный руководитель – к.ф.-м.н., доцент Иванов С.Е.

В работе проведен обзор расширений, рассмотрена их структура в различных имплементациях для разных браузеров. Также проанализирована безопасность расширений и их производительность. Рассмотрены проблемы поддержания расширений для старых версий (обратная совместимость) и портирования их для других браузеров.

Ключевые слова: браузерные расширения, браузер, анализ расширений, аддон, плагин.

Расширениями (плагинами, аддонами) для браузера пользуется большинство пользователей. Они помогают автоматизировать многие монотонные задачи, позволяя более удобно пользоваться браузером. Но с их установкой возникают некоторые проблемы – некоторые расширения работают только под определенными браузерами, другие же значительно снижают производительность. В данной работе были проанализированы расширения – их структура, влияние на производительность и возможность портирования между различными браузерами.

Первые расширения были представлены в 1998 году в браузере Mozilla Firefox [1]. Разработчики решили дать возможность пользователю самому выбрать дополнительные функции браузера, поэтому, несмотря на то, что во встроенной комплектации Firefox было не так много возможностей, браузер легко расширялся за счет установки необходимых пользователю аддонов. Такой подход сделал браузер легковесным (что в те времена было большим преимуществом), легко поддерживаемым и стабильным. Также это позволяло легко расширять функционал сайтов – благодаря скриптам, которые встраивались в страницу, можно было расширять возможности на страницах (например, скачивать медиаконтент при нажатии на одну кнопку и т.д.).

Впоследствии расширения стали поддерживать и другие браузеры, и в настоящее время все современное программное обеспечение для просмотра веб-страниц имеет возможность устанавливать расширения, хоть и иногда под другим названием (плагины, аддоны и т.д.). В данной работе был использован термин «Расширения».

Расширения – это программа, которая встраивается в браузер и расширяет его функциональные возможности. Расширения могут модифицировать браузер (добавлять новые

элементы управления с определенным функционалом), его настройки (например, управление закладками или проху-серверами), а также модифицировать страницы, открытые в данном браузере (добавляя новый функционал на страницы и/или меняя контент на ней).

Расширения для современных браузеров используют следующие технологии – JavaScript, HTML, CSS и XML. Эти технологии – стандартный стек для клиентской части веб-приложения. HTML и CSS используются для отображения HTML-документа (в поп-ап окнах или на странице настроек расширения), JavaScript же отвечает за логику расширения. При этом скрипты могут быть двух видов.

Контент-скрипты встраиваются в страницу, открытую в браузере, и используются для ее модификации, используя JavaScript. Бэкграунд-скрипты запускаются при открытии расширения и работают до момента его отключения, выполняясь постоянно на специальной фоновой странице. Бэкграунд-скрипты и контент-скрипты могут взаимодействовать друг с другом, обеспечивая достаточно легкую имплементацию переноса данных с одной страницы на другую.

Основная структура расширения со ссылками на скрипты, HTML-страницы, CSS-стили и медиаконтент (иконки, картинки, музыка и т.д.) представлена в файле `manifest.json`, который присутствует в любом расширении. Ссылки на различные файлы приложения определяются с помощью специальных ключевых слов. Эти ключевые слова могут быть как определенными для каждого браузера (нативный синтаксис) или определенные стандартом WebExtension API (Application Programming Interface).

Для совместимости между различными браузерами был создан стандарт WebExtension API. Основные функции данного API описаны в документе, составленном W3C Community Group [2]. Этот интерфейс реализован в большинстве современных браузеров. Но не все расширения используют данное API, иногда полностью или частично полагаясь на предоставляемое разработчиками браузера API, что делает такие приложения несовместимыми с некоторыми другими браузерами. При этом некоторые разработчики браузеров предоставляют полезные статьи [3] для относительно несложного портирования приложений (со сравнением одинаковых функций, имеющих разные имена в разных браузерах, и определенными нюансами портирования). Ниже представлена таблица совместимости между различными браузерами.

Таблица. Совместимость расширений для различных браузеров

	Mozilla Firefox	Google Chrome	Opera	Yandex Browser
Mozilla Firefox	X	–	–	–
Google Chrome	Частично совместимы	X	+	+
Opera	Частично совместимы	+	X	+
Yandex Browser	Частично совместимы	+	+	X

Из-за того, что браузеры Google Chrome, Opera и Yandex Browser используют chromium как движок для отображения страниц, то многие расширения, написанные для какого-либо из этих браузеров, легко открываются и в одном из других браузеров. Firefox же используется собственный движок gecko, который имеет определенные особенности, что затрудняет простой перенос приложений. Но если приложение было написано с использованием WebExtension API, то оно должно работать во всех браузерах, поддерживающих его (его поддерживают 4 браузера).

Со стороны безопасности браузеры ограничивают доступные для исполнения функции в скриптах и запрещают выполнение встроенных скриптов (через функцию `eval()` или напрямую в HTML-странице). Также работают политики безопасности, применяемые и для стандартных скриптов на сайтах (запрет на доступ к кросс-домен контенту). В файле

manifest.json разработчикам необходимо указать права расширения, и пользователь, при установке приложения, будет знать, к каким его данным расширение будет иметь доступ.

Расширения ухудшают работу браузера, в среднем замедляя запуск браузера на 10% с каждым расширением [4]. Также они увеличивают нагрузку на ЦПУ и оперативную память, и могут приводить к утечкам памяти. При этом необходимо понимать, что многие расширения используют такие базы данных, как SQLite, что будет негативно сказываться как на потреблении оперативной памяти, так и на потреблении места на жестком диске. Помимо этого, продвинутое расширение, например блокировщики рекламы (adBlock, uBlock и т.д.) используют проверку запросов браузера на низком уровне, замедляя время отклика. Поэтому целесообразно использовать только необходимые расширения и удалять или отключать неиспользуемые.

Разработчики браузеров изначально старались поддерживать обратную совместимость, но в настоящее время некоторые браузеры изменяют API, что приводит к неработоспособности предыдущих расширений на новых версиях браузеров. Например, Mozilla Firefox, относительно недавно, прекратил поддержку legacy-расширений.

После анализа расширений были сделаны следующие выводы – расширения имеют продуманную структуру, которая позволяет легко и удобно модифицировать его контент. Установка множества расширений значительно замедлит работу браузера, поэтому следует устанавливать только часто используемые расширения. Расширения слабо совместимы между разными браузерами, если они написаны на разных движках и не используют WebExtension API. Также нет какого-либо инструмента для портирования расширения с нативного синтаксиса браузера на WebExtension API. На создание такого инструмента и будут направлены дальнейшие исследования.

Литература

1. Жарков А.А. Сравнительный анализ интернет-браузеров [Электронный ресурс]. – Режим доступа: <https://files.scienceforum.ru/pdf/2014/7629.pdf> (дата обращения: 06.01.2019).
2. Browser Extensions [Электронный ресурс]. – Режим доступа: <https://browserext.github.io/browserext/#api-availability> (дата обращения: 06.01.2019).
3. Chrome Incompatibilities [Электронный ресурс]. – Режим доступа: https://developer.mozilla.org/en-US/docs/Mozilla/Add-ons/WebExtensions/Chrome_incompatibilities (дата обращения: 06.01.2019).
4. Improving Add-on Performance [Электронный ресурс]. – Режим доступа: <https://blog.mozilla.org/addons/2011/04/01/improving-add-on-performance/> (дата обращения: 06.01.2019).

**Долгачев Руслан Ферузович**

Год рождения: 1995

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K41201Направление подготовки: 11.04.02 – Программное обеспечение
в инфокоммуникациях

e-mail: dolgachev9@gmail.com

**Осипов Никита Алексеевич**

Год рождения: 1972

Университет ИТМО, факультет инфокоммуникационных технологий,
к.т.н., доцент

e-mail: nikita@ifmo.spb.ru

УДК 004. 93'14

**АНАЛИЗ ПРОЦЕССА КЛАССИФИКАЦИИ ИЗОБРАЖЕНИЯ С ПОМОЩЬЮ
МАШИННОГО ОБУЧЕНИЯ****Долгачев Р.Ф.****Научный руководитель – к.т.н., доцент Осипов Н.А.**

В работе проведено исследование метода классификации изображений с помощью машинного обучения. Для этого выполнен аналитический обзор метода машинного обучения с использованием сверточных нейронных сетей.

Ключевые слова: машинное обучение, компьютерное зрение, сверточная нейронная сеть, классификация, функция активации.

Введение. В наше время классификация изображений достаточно активно развивается и уже применяется в различных сферах, например, в системах организации информации (индексация баз данных изображений), в системах управления процессами (автономные транспортные средства), в системах видеонаблюдения и многих других [1]. Классификация изображений – это задача машинного обучения, для решения которой необходимо определить набор целевых классов (объектов, которые нужно идентифицировать на изображениях) и обучить модель распознавать их, используя маркированные примеры фотографий [2]. Во время этого могут возникнуть ожидаемые проблемы, для решения которых нужно знать этапы построения модели.

Классификация изображений. Ранние модели компьютерного зрения полагались на необработанные данные пикселей в качестве входных данных для модели [3]. Однако эти данные не могут обеспечить достаточно стабильного представления, чтобы охватить множество вариантов объекта, захваченного на изображении. Положение объекта, фон за объектом, окружающее освещение, угол камеры и фокус камеры – все это может привести к колебаниям в необработанных пиксельных данных. Для более гибкого моделирования объектов в классические модели компьютерного зрения добавили новые элементы, основанные на данных пикселей, таких как цветовые гистограммы, текстуры и формы. Недостатком этого подхода является огромное количество входных данных.

Сверточные нейронные сети. Прорыв в построении моделей для классификации изображений произошел после того, как научились использовать сверточную нейронную сеть (CNN) для постепенного извлечения более высокого уровня представлений содержимого изображения. Вместо предварительной обработки данных для получения таких элементов, как текстуры и формы, CNN принимает в качестве входных данных только необработанные пиксельные данные изображения и «учится» извлекать эти элементы и, в конечном итоге, выводить, какой объект они составляют.

Для начала CNN получает карту входных объектов: трехмерную матрицу, где размер первых двух измерений соответствует длине и ширине изображений в пикселях. Размер третьего измерения равен трем (соответствует трем каналам цветного изображения: красному, зеленому и синему). CNN содержит стек модулей, каждый из которых выполняет три операции.

Свертка извлекает плитки входной карты объектов и применяет к ним фильтры для вычисления новых объектов, создания выходной карты объектов или свернутых объектов (которые могут иметь другой размер и глубину, чем карта входных объектов). Свертки определяются двумя параметрами:

- размер извлекаемых плиток (обычно 3×3 или 5×5 пикселей);
- глубина выходной карты объектов, которая соответствует количеству примененных фильтров.

Во время свертки фильтры (матрицы того же размера, что и размер элемента мозаичного изображения) эффективно проходят по сетке входной карты объектов по горизонтали и вертикали, по одному пикселю за раз, выделяя каждый соответствующий элемент мозаичного изображения [4].

Для каждой пары элементов мозаичного изображения CNN выполняет поэлементное умножение матрицы фильтра и матрицы элементов мозаичного изображения, а затем суммирует все элементы результирующей матрицы, чтобы получить одно значение. Каждое из этих результирующих значений для каждой пары «фильтр-мозаика» затем выводится в свернутой матрице признаков.

Во время обучения CNN «запоминает» оптимальные значения для матриц фильтра, которые позволяют ему извлекать значимые элементы (текстуры, края, фигуры) из входной карты объектов. По мере увеличения количества фильтров (глубины карты выходных объектов), применяемых к входным данным, увеличивается и количество объектов, которые CNN может извлечь [5]. Однако компромисс заключается в том, что фильтры составляют большую часть ресурсов, затрачиваемых CNN, поэтому время обучения также увеличивается по мере добавления большего количества фильтров. Кроме того, каждый фильтр, добавляемый в сеть, обеспечивает меньшее добавочное значение, чем предыдущий, поэтому инженеры стремятся создавать сети, использующие минимальное количество фильтров, необходимых для извлечения функций, необходимых для точной классификации изображений.

После каждой операции свертки CNN применяет преобразование выпрямленной линейной единицы (ReLU, rectified linear unit) к свернутой функции, чтобы ввести нелинейность в модель. Функция ReLU, $F(x) = \max(0, x)$, возвращает x для всех значений $x > 0$ и возвращает 0 для всех значений $x \leq 0$.

После ReLU наступает этап объединения, на котором CNN сокращает число свернутых объектов (чтобы сэкономить время обработки), сокращая количество измерений карты объектов, сохраняя при этом наиболее важную информацию об элементах. Общий алгоритм, используемый для этого процесса, называется максимальным объединением (max pooling).

Максимальное объединение работает аналогично свертке. Мы проходим по карте объектов и извлекаем плитки указанного размера. Для каждой плитки максимальное значение выводится на новую карту объектов, а все остальные значения отбрасываются. Операции принимают два параметра:

1. размер фильтра максимального объединения (обычно 2×2 пикселя);
2. шаг – расстояние в пикселях, отделяющее каждую извлеченную плитку. В отличие от свертки, когда фильтры проходят по карте объектов попиксельно, при максимальном объединении шаг определяет места, где извлекается каждая плитка. Для фильтра 2×2 шаг 2 указывает, что операция максимального объединения будет извлекать все неперекрывающиеся фрагменты 2×2 из карты объектов.

В конце сверточной нейронной сети находятся один или несколько полностью связанных уровней (когда два уровня «полностью связаны», каждый узел первого уровня связан с каждым узлом второго уровня). Их работа заключается в выполнении классификации на основе признаков, извлеченных из свертки. Как правило, последний полностью связанный уровень содержит функцию активации softmax, которая выводит значение вероятности от 0 до 1 для каждой метки классификации, которую модель пытается предсказать.

Заключение. Таким образом, можно сделать вывод, что, как и в любой модели машинного обучения, ключевой проблемой при обучении сверточной нейронной сети является переобучение: модель оказывается настолько настроенной на специфику обучающих данных, что ее невозможно применить на новые примеры. Используются два метода предотвращения переобучения при создании CNN:

- увеличение данных: искусственное увеличение разнообразия и количества обучающих примеров путем выполнения случайных преобразований существующих изображений для создания набора новых вариантов. Дополнение данных особенно полезно, когда исходный набор обучающих данных относительно мал;
- регуляризация отсева: случайное удаление единиц из нейронной сети во время шага градиента обучения.

Литература

1. Компьютерное зрение [Электронный ресурс]. – Режим доступа: https://ru.wikipedia.org/wiki/Компьютерное_зрение (дата обращения: 06.01.2019).
2. Машинное обучение [Электронный ресурс]. – Режим доступа: https://ru.wikipedia.org/wiki/Машинное_обучение (дата обращения: 06.01.2019).
3. ML Practicum: Image Classification [Электронный ресурс]. – Режим доступа: <https://developers.google.com/machine-learning/practica/image-classification/> (дата обращения: 06.01.2019).
4. Наглядное введение в нейросети на примере распознавания цифр [Электронный ресурс]. – Режим доступа: <https://proglab.io/p/neural-network-course/> (дата обращения: 06.01.2019).
5. Simple Image classification using deep learning-deep learning series 2 [Электронный ресурс]. – Режим доступа: <https://medium.com/intro-to-artificial-intelligence/simple-image-classification-using-deep-learning-deep-learning-series-2-5e5b89e97926> (дата обращения: 06.01.2019).



Загряжская Наталья Ильинична

Год рождения: 1996

Университет ИТМО, факультет инфокоммуникационных технологий, студент группы № К4220

Направление подготовки: 11.04.02 – Программное обеспечение в инфокоммуникациях

e-mail: nzagryajskaya@gmail.com



Иванов Сергей Евгеньевич

Год рождения: 1975

Университет ИТМО, факультет инфокоммуникационных технологий,

к.ф.-м.н., доцент

e-mail: 45is@mail.ru

УДК 004.6

МЕТОДЫ ПОИСКА СИНОНИМОВ ПРИ ОБРАБОТКЕ ЕСТЕСТВЕННОГО ЯЗЫКА

Загряжская Н.И.

Научный руководитель – к.ф.-м.н., доцент Иванов С.Е.

Работа посвящена исследованию методов поиска синонимов, и предполагается новый подход к разрешению проблемы синонимии для словосочетаний. Экспериментально применен метод нахождения векторного представления слов и оценки их схожести.

Ключевые слова: обработка естественного языка, синоним, проблема синонимии, word2vec, обработка текста.

В работе затронута проблема выбора наиболее подходящего метода для решения проблемы синонимии, ведь одно понятие может быть выражено различными словами. В результате этого явления релевантные документы, в которых используются синонимы понятий, указанных пользователем в запросе, могут быть не обнаружены системой. Главная цель работы была в том, чтобы рассмотреть основные техники, которые можно использовать для извлечения синонимичных понятий в корпусах для обучения моделей поисковых систем и классификаторов.

Для исследования и оценки методов поиска синонимов проведем эксперименты на базе данных Wikipedia. Общая информация об исходных данных: 3122 текстов на английском языке, более 2 миллионов слов. Исходный корпус данных содержит 13 колонок, из них наиболее интересные для нас это title, text, lang.

Большинство методов генерации синонимов решают эту задачу как двухэтапную проблему [1]. Первым этапом является Candidate Generation (нахождение кандидатов в синонимы). На этом шаге создаются все возможные кандидаты, которые могут быть синонимами этого слова. В зависимости от конкретного приложения вот два самых распространенных подхода:

1. Word embeddings: способ тренировать векторы слов для корпуса, а затем находить синонимы для текущего слова, используя ближайших соседей или определяя некоторое понятие «сходства» (в основном с помощью similarity matrix) [2];
2. Historical user data: также можно проследить историю поведения пользователя, чтобы понять, что лично для него может являться схожими понятиями;
3. Lexical synonyms: это грамматические синонимы, определенные правилами языка.

Для исследования был использован word embeddings, а именно функционал, реализованный библиотекой word2vec. Word2Vec – это векторный алгоритм, в котором оператор в основном берет последовательности текста, используя скользящее окно фиксированной длины, скользит по тексту и тренирует векторное представление каждого слова, используя векторные представления окружающих его слов. Ключевым моментом является то, что каждое слово встраивается в вектор фиксированной длины, и этот вектор определяется минимизацией ошибки с помощью линейных комбинаций того, что вокруг него в тексте. Доказана эффективность алгоритма для отдельных слов, в основной части работы расширены возможности алгоритма для двух слов (словосочетания).

Второй этап генерации синонимов – это собственно их обнаружение (Synonym detection). Когда есть набор кандидатов-синонимов для данного слова, необходимо выяснить, какие из них на самом деле синонимы.

Построим собственную модель для извлечения синонимов из тестового корпуса. Для этого необходимо:

1. выстроить similarity matrix между векторными представлениями слов;
2. найти близкие понятия (выберем для оценки косинусное расстояние);
3. найти вектор словосочетания, и попробовать найти близкие для него понятия.

Многие современные системы и методы нейролингвистического программирования рассматривают слова как атомарные единицы – понятия сходства между словами нет, так как они представлены в виде индексов в словаре. Предварительная обработка текста для дальнейшего анализа включает в себя:

1. токенизацию (Tokenization), т.е. разбиение текста на отдельные слова. Эта задача нетривиальна для некоторых азиатских языков, включая китайский, но для английского языка такой проблемы нет;
2. нормализация (Token normalization), что в зависимости от поставленной задачи может быть, как простым складыванием (отбрасыванием информации о регистре букв), так и сложным, как лемматизация (полный морфологический анализ в синтетических, и особенно флективных, языках, таких как чешский);
3. исправление ошибок в правописании (Spelling correction).

После базовой предобработки текста начинается подготовка к обучению модели. Тренируем корпус, основываясь на подходе Радима Рехурека (Radim Rehurek) [3], а именно обучим векторную модель и проведем оценку, основываясь на матрицах схожести (similarity matrix):

```
model = gensim.models.Word2Vec(documents,size=150>window=10,min_count=2,workers=10)
```

```
model.train(documents,total_examples = len(documents),epochs=10)
```

За кулисами авторы фактически тренируют простую нейронную сеть с одним скрытым слоем. Цель на этом этапе состоит в том, чтобы узнать вес скрытого слоя. Эти веса, по сути, являются векторами слов, которые необходимо выучить. Для улучшения качества результата и экономии оперативного запоминающего устройства (ОЗУ) основная нагрузка переносится на KeyedVectors, и не будет использоваться вся модель при каждом запросе на извлечение синонима. Особенность KeyedVectors в том, что им не нужно хранить состояние модели, которое позволяет обучение. Несколько процессов могут повторно использовать одни и те же данные, сохраняя только одну копию в ОЗУ с помощью mmap.

Рассмотрим качество построенной модели для отдельных слов (табл. 1).

Таблица 1. Результаты экспериментов для одного слова

Слово	Найденные синонимы
Good	great best
System	model process

Слово	Найденные синонимы
Python	Scala Java JavaScript Kotlin programming

Данные результаты получены на нахождении матрицы однородности для векторного представления каждого слова по отдельности. Далее переучим модель с более широким окном и для словосочетаний попробуем найти численные представления, сохраняющие семантическую связь с остальными словами. Для этого, кроме настройки модели на большее количество эпох обучения и расширенное окно, нужно предусмотреть, что исходная матрица все равно будет дробить словосочетания на отдельные векторы [4].

Попробуем для подхода:

1. суммирование векторов с сохранением контекста;
2. среднее значение вектора с повторным переобучением.

Результаты эксперимента 1 (табл. 2).

Таблица 2. Результаты экспериментов для словосочетания в эксперименте 1

Словосочетание	Найденные синонимы
golden gate	crossing band
white swan	black rat
new website	existing brand

Для проведения эксперимента 2 напомним функции для вычисления среднего двух векторов, а именно: функции сложения векторов, нахождения их суммы, скалярное произведение, и используем их в функции `vector_mean(word1, word2)` [5].

`f = vector_mean([model['good'], model['time']])`

Результатом будет 'great moment'. Продолжим настройку модели, и приведем примеры результатов (табл. 3).

Таблица 3. Результаты экспериментов для словосочетания

Словосочетание	Найденные синонимы
program performance	deployment reliability
Python language	programming language Java framework
language processing	speech recognition
credit card	debit card cash

Как можно видеть, качество предоставленных синонимов во втором эксперименте значительно лучше, и более того, с человеческой точки зрения это действительно близко стоящие понятия.

В данной работе рассмотрены основные существующие методы решения проблемы синонимии, и авторы попробовали расширить применение векторной интерпретации связей между словами и на поиск синонимичных словосочетаний. Можно сделать вывод, что в итоге синонимы выглядят логично и читаемо, что является хорошим поводом для продолжения исследования.

Литература

1. Bird S., Klein E., Loper E. Natural Language Processing with Python: Analyzing Text with the Natural Language Toolkit. – O'Reilly Media, 2009. – 504 p.

2. Mikolov T., Sutskever I., Chen K., Corrado G., Dean J. Distributed Representations of Words and Phrases and their Compositionality [Электронный ресурс]. – Режим доступа: <https://arxiv.org/pdf/1310.4546.pdf> (дата обращения: 06.01.2019).
3. [Электронный ресурс]. – Режим доступа: https://radimrehurek.com/phd_rehurek.pdf (дата обращения: 06.01.2019).
4. Клиот-Дашинский М.И. Алгебра матриц и векторов. – СПб.: Лань, 2001. – 160 с.
5. Лутц М. Программирование на Python. – Т. 2. – Изд-во: Символ-плюс, 2011. – 992 с.



Зенцова Елена Сергеевна

Год рождения: 1996

Университет ИТМО, факультет программной инженерии
и компьютерной техники, студент группы № K41401

Направление подготовки: 09.04.02 – Информационные системы
и технологии

e-mail: zencova1996@mail.ru



Иванов Сергей Евгеньевич

Год рождения: 1975

Университет ИТМО, факультет инфокоммуникационных технологий,
к.ф.-м.н., доцент

e-mail: 45is@mail.ru



Лавская Лина Владимировна

Год рождения: 1996

Университет ИТМО, факультет программной инженерии
и компьютерной техники, студент группы № P41601

Направление подготовки: 09.04.02 – Информационные системы
и технологии

e-mail: lavskayalv@yandex.ru

УДК 004.02

ТЕХНОЛОГИИ И СРЕДСТВА РАЗРАБОТКИ МЕДИЦИНСКОЙ ИНФОРМАЦИОННО-АНАЛИТИЧЕСКОЙ СИСТЕМЫ

Зенцова Е.С., Иванов С.Е., Лавская Л.В.

Научный руководитель – к.ф.-м.н., доцент Иванов С.Е.

В работе рассмотрены вопросы анализа архитектуры системы «РЕГИЗ», представлен сравнительный анализ функциональных возможностей медицинских информационных систем, производится выбор показателя, на основе которого будет формироваться мониторинг. В рамках работы описан анализ математических методов, как средств проверки результатов разработанных отчетов.

Ключевые слова: информационная система, медицинская информационная система, РЕГИЗ, МИС, информационно-аналитическая система, SQL.

В настоящее время технологии в медицинских системах здравоохранения выходят на новый уровень. Так, уже применяются онлайн-запись на прием к врачу, телемедицина, запущена национальная система, ведутся электронные больничные листы, а также мониторинги по наиболее важным показателям.

Благодаря внедрению информационных технологий (ИТ) в медицину открылось большое количество возможностей как для пациентов, так и для медицинского персонала. Яркими примерами являются возможности обмена данными пациентов между разными учреждениями, дистанционного обучения персонала, ведение электронного документооборота, телеконсультация пациентов и персонала посредством интернета и различных гаджетов, контроль над проведением хирургических вмешательств в реальном времени, анализ больших объемов данных, онлайн-запись на прием к врачу и др.

Без ИТ уже сложно представить большинство сфер нашей жизни, в том числе и здравоохранение, поэтому существует необходимость в постоянном совершенствовании уже имеющихся систем.

Таким образом, была поставлена следующая глобальная цель – создание отчета в виде сложных SQL для мониторинга показателей объема выполненных услуг по амбулаторным талонам и выпискам из стационара.

В рамках первого этапа работы были выполнены следующие задачи:

1. анализ архитектуры системы «РЕГИЗ», на базе которой будет производиться дальнейшая разработка;
2. сравнительный анализ функциональных возможностей медицинских информационных систем (ИС);
3. выбраны показатели, на основе которых будет формироваться мониторинг;
4. анализ математические методов, как средства проверки результатов разработанных отчетов;
5. выделены предложения по улучшению мониторинга.

Анализ архитектуры медицинской ИС «РЕГИЗ». РЕГИЗ обеспечивает информационную поддержку функционирования системы здравоохранения Санкт-Петербурга для органов управления здравоохранением, работников медицинских организаций и жителей города, выполняя в Санкт-Петербурге функции регионального фрагмента Единой государственной ИС в сфере здравоохранения.

Сервисы и подсистемы РЕГИЗ интегрированы с федеральными сервисами «Федеральная электронная регистратура» и «Интегрированная электронная медицинская карта», а также с подавляющим большинством медицинских ИС медицинских организаций Санкт-Петербурга.

Конфигурация РЕГИЗ позволяет подключаться любым ИС, удовлетворяющим протоколам обмена и требованиям безопасности и защиты персональных данных.

В состав РЕГИЗ входят несколько десятков подсистем, обеспечивающих автоматизацию следующих основных процессов:

- запись на прием к врачу в электронном виде;
- интегрированная электронная медицинская карта;
- управление очередями на плановую госпитализацию, консультативный прием, исследования КТ/МРТ;
- обмен данными лабораторных исследований и др. [1].

РЕГИЗ имеет клиент-серверную архитектуру, которая состоит из следующих компонентов:

- сервер баз данных, управляющий хранением данных, доступом и защитой, резервным копированием, отслеживающий целостность данных в соответствии с бизнес-правилами и, самое главное, выполняющий запросы клиента (в РЕГИЗ данную позицию занимает Microsoft SQL Server);
- клиент, предоставляющий интерфейс пользователя, выполняющий логику приложения, проверяющий допустимость данных, посылающий запросы к серверу и получающий ответы от него;
- сеть и коммуникационное программное обеспечение, осуществляющее взаимодействие между клиентом и сервером посредством сетевых протоколов.

В рамках анализа архитектуры системы «РЕГИЗ» были выделены основные преимущества:

- разделение программного кода;
- пониженные требования к машинам-клиентам;
- гибкость системы.

Также были выделены недостатки:

- большие затраты на серверное оборудование;

– компетентные кадры, обслуживающие сервер.

В качестве заключения стоит явно акцентировать внимание на том, что архитектура клиент-сервер не делит машины на только клиент или только сервер, а скорее позволяет распределить нагрузку и разделить функционал между клиентской частью и серверной [2].

Также в рамках исследования был проведен сравнительный анализ функционала медицинских информационных систем (МИС).

В качестве конкурентных систем были взяты четыре наиболее известных МИС российского производства, описание и скриншоты которых есть в открытом доступе. Среди фирм разработчиков есть предприятия как с частной формой собственности, так и с государственной.

В качестве критериев выявления качества МИС были выбраны наиболее востребованные на сегодняшний день функции, которые ранее были определены Киреевым В.С., Агамовым Н.А. в статье [3]. По представленным критериям была также оценена государственная ИС «РЕГИЗ». По максимальному количеству критериев были проставлены баллы каждой из представленных МИС. Результаты сравнительного анализа представлены в таблице.

Таблица. Сравнительный анализ функционала МИС

Функционал	МИС Инфоклиника	ТеКо Мед	МИС ЕМИАС	ГИС «РЕГИЗ»
1. Регистратура поликлиники и/или приемного отделения стационара: – регистрация и хранение персональных (паспортных) данных обслуживаемых пациентов; – создание первичных медицинских документов и печать этих документов; – учет прикрепления, открепления, перерегистрации обслуживаемых граждан, анализ движения прикрепленного контингента; – печать различных документов о согласии пациента (на хранение и обработку данных, оперативные вмешательства и т.д.); – запись на прием к врачу; – направление пациентов в отделения.	6	5	5	5
2. Электронная медицинская карта (ЭМК) пациента: – ведение/регистрация/хранение документации врачебных осмотров; – учет случаев обращений пациента; – возможность распечатки стандартного листа назначений; – формирование направлений на получение медицинской помощи; – автоматизированный подбор медико-экономических стандартов лечения по показаниям и назначение лечения по стандарту; – планирование и учет результатов оперативных вмешательств;	10	9	10	9

Функционал	МИС Инфоклиника	ТеКо Мед	МИС ЕМИАС	ГИС «РЕГИЗ»
– формирование листов наблюдений; – наличие интерфейса просмотра динамики лабораторных показателей без разбивки по случаям; – возможность экспорта ЭМК на внешний носитель, а также в формате HL7.				
3. Управление взаиморасчетами за оказанную медицинскую помощь: – учет вида финансирования пациентов; – учет для каждой услуги номенклатуры видов финансирования; – ведение прейскурантов цен на услуги; – настройка импорта цен на услуги; – ведение перечня контрагентов и договоров на оказание медицинских услуг; – формирование реестров счетов за услуги.	6	5	6	4
4. Анализ деятельности и формирование отчетности: – подготовка произвольных аналитических отчетов о деятельности организации; – подготовка утвержденной государственной статистической отчетности; – предварительный просмотр сформированного отчета, печать отчетов; – экспорт отчетов в офисные приложения, включая MS Office, OpenOffice; – экспорт отчетов в TIFF, PDF, HTML и т.д.	5	5	3	3

Все эти показатели развернуты максимально широко для ознакомления с основными возможностями современной медицинской ИС.

Максимальное количество баллов в рамках сравнительного анализа функционала известных медицинских ИС набрала МИС Инфоклиника. Каждая из представленных систем имеет схожие модули, которые, безусловно, еще можно и нужно усовершенствовать.

МИС Инфоклиника, исходя из сравнительного анализа, может послужить примером для дальнейшего усовершенствования ГИС «РЕГИЗ». Благодаря тому, что есть к чему стремиться и что совершенствовать, было принято решение о создании отчета «Мониторинг качества ведения данных по амбулаторным талонам и выпискам из стационара», как способа улучшения общего функционала системы «РЕГИЗ» [4, 5].

Выбор показателей, на которых будет основываться мониторинг. После изучения системы «РЕГИЗ» в целом, сравнении ее функционала с основными российскими МИС, было принято решение по разработке отчета в виде сложных SQL для мониторинга качества ведения данных по амбулаторным талонам и выпискам из стационара, как дополнительного модуля отчетности.

Для получения информации о качестве ведения данных по амбулаторным и стационарным случаям лечения на основании анализа данных, полученных из специализированной ИС, был разработан основной показатель – качество ведения данных,

который определяется соотношением количества заполненных полей по представленным показателям и общего количества записей.

Анализ математических методов, как средства оценки результатов. В рамках первого этапа работы был также проведен анализ научной литературы по теме применения математических методов, а также методов статистики по направлениям анализа затрат и эффективности использования ресурсов медицинских организаций. Для повышения эффективности деятельности медицинских организаций, а также повышения качества оказания медицинской помощи целесообразно использовать специальные инструментальные средства.

В рамках работы проведен сравнительный анализ четырех математических методов:

- статистический анализ;
- численные методы;
- аналитический метод;
- графический метод.

Исходя из сравнительного анализа, был выбран статистический метод, так как он наиболее подходит для медицинской области. Показатели в медицине достаточно специфичны, обладают большой вариабельностью, их характеристики меняются во времени и пространстве в зависимости от многих факторов, а также существенно отличаются друг от друга [6].

В завершение были выделены предложения по улучшению мониторинга. В дальнейшей перспективе планируется осуществлять оценку качества ведения данных с применением дополнительных показателей для более точного мониторинга качества ведения данных медицинскими организациями, например, введение показателя заполнения полей по конкретным группам талонов, в конкретных учреждениях и по конкретным заболеваниям для дальнейшей оптимизации работы учреждений.

Литература

1. Подсистемы ГИС «РЕГИЗ» [Электронный ресурс]. – Режим доступа: <http://spbmiac.ru/ehlektronnoe-zdravookhranenie/podsistemy-gis-regiz/> (дата обращения: 12.01.2019).
2. Гайнанова Р.Ш., Широкова О.А. Создание клиент-серверных приложений // Вестник Казанского технологического университета. – 2017. – С. 79–84.
3. Киреев В.С., Агамов Н.А. Сравнительный обзор медицинских информационных систем, представленных на российском рынке // Международный научно-технический журнал «Теория. Практика. Инновации». – 2017. – С. 184–193.
4. Инфоклиника. Базовая версия системы [Электронный ресурс]. – Режим доступа: <http://www.sdsys.ru/modules/26/> (дата обращения: 12.01.2019).
5. Шевелева О. ЕМИАС – Здравоохранение будущего // Врач и информационные технологии. – 2012. – С. 75–77.
6. Румянцев П.О., Саенко В.А., Румянцева У.В., Чекин С.Ю. Статистические методы анализа в клинической практике. – Обнинск: ГУ РМНЦ РАМН, 2009. – 46 с.

**Карлов Борис Александрович**

Год рождения: 1977

Университет ИТМО, факультет инфокоммуникационных технологий, студент группы № K41201

Направление подготовки: 11.04.02 – Программное обеспечение в инфокоммуникациях

e-mail: bkarlov@gmail.com

**Ананченко Игорь Викторович**

Год рождения: 1968

Университет ИТМО, факультет инфокоммуникационных технологий, к.т.н., доцент

e-mail: igor@anantchenko.ru

УДК 004.891.2

ИССЛЕДОВАНИЕ ТЕХНОЛОГИИ СЛИЯНИЯ ДАННЫХ**Карлов Б.А., Ананченко И.В.****Научный руководитель – к.т.н., доцент Ананченко И.В.**

В работе рассмотрена технология слияния данных: история возникновения и развития, области применения, сравнение со смежными технологиями, актуальность в настоящее время, возможность использования в проектах программы «Умный город» в Санкт-Петербурге.

Ключевые слова: слияние данных, разнородные источники, Data Fusion.

Введение. Слияние данных – процесс объединения источников данных для получения более согласующейся, точной и полезной информации, чем информация от одного отдельного источника [1] (рис. 1).

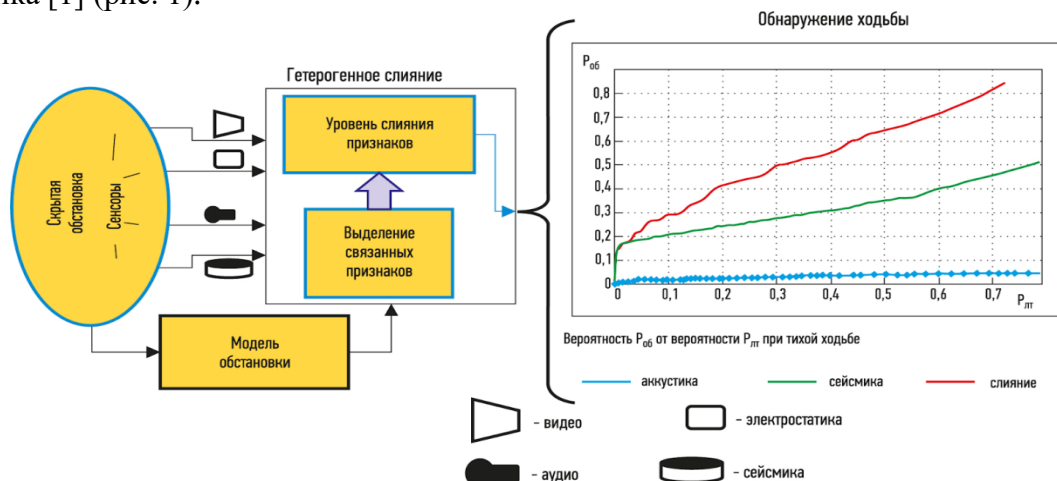


Рис. 1. Повышение вероятности обнаружения нарушителей за счет слияния информации от датчиков [2]

Задача слияния данных появилась достаточно давно – более полувека назад появились первые математические алгоритмы, а сама идея обработки информации с нескольких совместно используемых датчиков возникла в 1950-х годах. Тогда технология была применима большей частью в военных разработках и требовала серьезных вычислительных мощностей, сейчас же она снова стала востребованной – с появлением проектов умных

городов, с развитием интернета вещей, с популярностью персональной бытовой электроники с мониторингом здоровья и функциями виртуальной и дополненной реальности [3].

Как уже упоминалось ранее, изначально технология нашла свое применение в военных целях, для автоматического определения целей, уточнения боевой обстановки и задач навигации и контроля автономных объектов. В то же время сам человек функционирует по схожим принципам: именно комбинируя информацию от разных органов чувств, мозг принимает те или иные решения. В современной жизни можно встретить множество других примеров – от автомобильного бортового компьютера до персонального фитнес-трекера.

Интеграция и слияние данных. Интеграцию данных можно рассматривать как комбинацию множеств, при которой большее множество сохраняется, в то время как слияние является техникой сокращения множества с улучшением надежности [1].

Для решения задачи интеграции данных существует немало программных продуктов. Достаточно упомянуть таких гигантов, как Oracle, Microsoft, IBM, SAP, при этом не всегда они лидеры рынка [4]. Однако задача интеграции данных, которая обычно стоит перед предприятиями, несколько отличается от слияния. Поэтому применение существующих решений требует некоторой адаптации и определения граничных условий.

Актуальность. В настоящее время интерес к технологии слияния данных возрастает. В практическом плане увеличивается число комплексных программно-аппаратных решений. При этом в центре внимания чаще всего находится решение различных ситуативных прикладных задач [5].

Только лишь по данным научной электронной библиотеки eLIBRARY.RU (рис. 2) за последние 15 лет, количество публикаций, связанных с технологией Data Fusion, увеличилось вдвое.

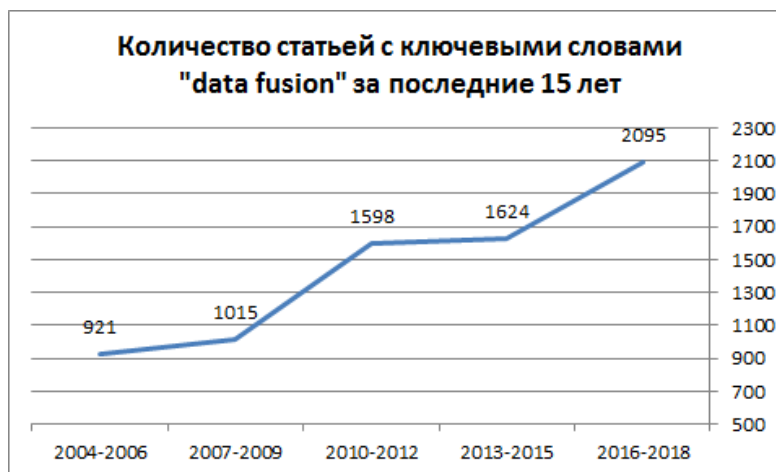


Рис. 2. Динамика количества тематических публикаций по данным научной электронной библиотеки eLIBRARY.RU

В настоящее время компания АО «ЭР-Телеком Холдинг» реализует в Санкт-Петербурге проекты, в которых может быть применена технология data/sensor fusion: исполнение программы «Безопасный Город» и запуск сети промышленного интернета вещей (IIoT) на базе сети LoRaWAN.

В каждом из этих проектов есть разнородные источники данных, разнородные сетевые устройства. Например, в «Безопасном Городе» есть видеоданные или результаты видеоаналитики и информация о пользователях бесплатной сети Wi-Fi, а в проекте IIoT – различные датчики одного объекта. Очевидно, изучение полученной информации при совместной работе всех источников может и должно придать проектам большую значимость и полезность.

Заключение. Целью следующей работы предполагается выбор подходящей модели слияния данных и применимого программного решения для одного из проектов, его разработка и, возможно, апробация на практике.

Литература

1. Википедия [Электронный ресурс]. – Режим доступа: <http://ru.wikipedia.org/> (дата обращения: 21.10.2018).
2. Буймистряк Г. Технологии слияния сенсорной информации для управления в критических ситуациях [Электронный ресурс]. – Режим доступа: <https://controlengrussia.com/teoriya/sliyaniye-sensornoj-informatsii/> (дата обращения: 02.12.2018).
3. Морри Голдман (перевод: В. Гавриков). Взросление технологии слияния данных Sensor Fusion [Электронный ресурс]. – Режим доступа: <https://spb.terraelectronica.ru/news/5488> (дата обращения: 02.12.2018).
4. Data Integration Tools [Электронный ресурс]. – Режим доступа: <https://www.trustradius.com/data-integration> (дата обращения: 23.12.2018).
5. Ананченко И.В., Гайков А.В., Мусаев А.А. Технологии слияния гетерогенной информации из разнородных источников (Data Fusion) [Электронный ресурс]. – Режим доступа: <https://elibrary.ru/item.asp?id=20164725> (дата обращения: 21.10.2018).



Кузьмина Ольга Юрьевна

Год рождения: 1995

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K4132

Направление подготовки: 27.04.03 – Системный анализ и управление

e-mail: kuzminaoljusya@mail.ru



Дюкина Татьяна Олеговна

Университет ИТМО, факультет инфокоммуникационных технологий,
к.э.н., доцент

e-mail: dto@mail.ru

УДК 330.42

ПРИМЕНЕНИЕ ДЕКОМПОЗИЦИИ ПРИ ПЛАНИРОВАНИИ ФИНАНСОВЫХ ПОКАЗАТЕЛЕЙ

Кузьмина О.Ю.

Научный руководитель – к.э.н., доцент Дюкина Т.О.

Работа посвящена такому инструменту как декомпозиция в планировании финансовых показателей предприятия; рассмотрена финансовая модель Дюпона и ее модификации; также представлено применение модели Дюпона для расчета показателей рентабельности активов и собственного капитала на примере предприятия в пивоваренной отрасли.

Ключевые слова: декомпозиция, финансовая модель, модель Дюпона, рентабельность собственного капитала.

Экономическая сфера в современном мире занимает одну из главных ниш в жизни общества, посредством управления которой создаются богатства, способные обеспечить необходимые условия для комфортного существования каждого человека. Финансовая устойчивость предприятия во многом определяется его способностью поддерживать баланс материально-вещественных ресурсов на каждом этапе функционирования, а также умением противостоять неблагоприятным обстоятельствам в условиях рыночной конкуренции. Для того чтобы предотвратить всевозможные критические ситуации на предприятии применяют такой инструмент, как планирование, применение которого позволяет оценить ключевые финансовые показатели предприятия на определенный промежуток времени, а также разработать основные задачи предприятия и пути их достижения.

Одной из важных задач для компании является проведение качественного анализа, который бы охватывал все основные финансовые показатели. Во многом успешный результат финансовой проверки зависит от выбора методик и инструментов. Основная цель инструмента, применяемого в финансовом анализе – это получение наиболее полной информации о состоянии предприятия: его прибыли и убытков, изменений в структуре активов и пассивов, а также его положения на рынке.

Термин «декомпозиция» означает метод, позволяющий детально исследовать финансовые показатели, тем самым значительно облегчив весь процесс их планирования и упростив процесс выявления сильных и слабых сторон субъекта хозяйствования. Во время

проведения финансового анализа экономисты могут столкнуться с рядом проблем, но некоторые из них можно решить с помощью декомпозиции.

Во-первых, процесс финансового анализа зачастую сводится к выявлению темпов изменения показателей, в то же время обнаружение причины увеличения или уменьшения финансового коэффициента является более важной задачей для аналитика.

Во-вторых, существует огромное множество финансовых коэффициентов, что затрудняет процесс их планирования.

Применение декомпозиции позволит не только определить причинно-следственную связь между финансовыми показателями, но и сконцентрировать внимание аналитика на ключевых коэффициентах для каждого аспекта деятельности компании.

Одним из частных примеров декомпозиции является модель Дюпона. На сегодняшний день конкретная финансовая модель достаточно распространена и довольно часто применима при планировании и прогнозировании финансовых показателей. Данная модель была представлена еще в 30-х годах XX в. Дональдсом Брауном и впоследствии стала одной из первых многофакторных моделей рентабельности [1].

Модель Дюпона основана на анализе соотношений такого показателя, как рентабельность собственного капитала. В результате математических преобразований модель декомпозируется и обретает вид двухфакторной, трехфакторной или пятифакторной, причем каждый ее компонент – самостоятельный содержательный финансовый показатель. Для оценки финансовой эффективности компании часто применяют относительные показатели, позволяющие измерить доходность компании с различных сторон. Ключевой показатель модели Дюпона – рентабельность собственного капитала, позволяющая судить, что денежные средства акционеров вложены целесообразно и приносят прибыль. Кроме того, коэффициент дает возможность аналитику соотнести положительный результат своей компании с возможным доходом от альтернативных вложений. Также более детальное представление переменных Дюпона позволяет оценить воздействие на рентабельность, уровень использования активов и характер привлеченных денежных средств [2].

Двухфакторная модель Дюпона дает возможность увидеть взаимосвязь таких коэффициентов, как рентабельность активов (ROA), рентабельность продаж (ROS) и оборачиваемость активов (K_{OA}). Увязка перечисленных факторов представляется в виде детерминированной зависимости, когда преобразование происходит путем сокращения: деления числителя и знаменателя на одинаковый показатель. В результате появляется модель с новым набором факторов. Представить данную взаимосвязь можно следующим образом:

$$ROE = \frac{ЧП}{EBT} \cdot \frac{EBT}{EBIT} \cdot \frac{EBIT}{B} \cdot \frac{B}{A} \cdot \left(1 + \frac{D}{E}\right), \quad (1)$$

где $ROA = ROS \left(= \frac{ЧП}{B}\right) \cdot K_{OA} \left(= \frac{B}{A}\right)$.

Основным показателем в данной модификации модели Дюпона является коэффициент рентабельности активов, который показывает эффективность расходования ресурсов предприятия. Анализируя составляющие двухфакторной модели, можно определить от чего зависит увеличение активов компании. Так повышение продаж и источников прибыли является причиной прироста ключевого показателя (ROA) данной модификации модели Дюпона. Поскольку одновременный рост рентабельности продаж и коэффициента оборачиваемости активов, как правило, невозможен, то одной из задач для аналитика является достижение равновесия между ними. Отрицательным моментом в двухфакторной модели – это невозможность проследить, за счет каких инвестиционных вложений происходит рост активов компании [3].

Устранить недостатки первой модификации модели Дюпона можно, изменив ее на трехфакторную, в которой рентабельность собственного капитала будет распадаться на три составляющие: рентабельность продаж, оборачиваемость активов и мультипликатор капитала.

Именно последний показатель дает возможность охарактеризовать инвестиционную структуру капитала компании. Другими словами, мультипликатор капитала есть ничто иное, как эффект финансового рычага, под которым понимают отношение заемного капитала к собственному. Для каждого акционера компании важно эффективное использование собственных активов хозяйствующего субъекта. Использование заемных средств при реализации дорогостоящих проектов представляет собой некий риск для акционеров компании, но в то же время способно обеспечить более высокую отдачу и операционную прибыль, что непосредственно увеличит доходность акционеров компании. Представленная ниже формула иллюстрирует взаимосвязь описанной трехфакторной модели [4].

$$ROE = ROS \cdot K_{OA} \cdot \left(1 + \frac{D}{E}\right) = ROA \cdot \left(1 + \frac{D}{E}\right). \quad (2)$$

Представление модели Дюпона в виде модифицированной пятифакторной формулы позволяет еще более детально рассмотреть взаимозависимость рентабельности собственного капитала и эффективности результата от операционной, финансовой и инвестиционной деятельности компании. Связь между показателем ROE и деятельностью компании можно увидеть через призму коэффициента налогового бремени, представляющего отношение чистой прибыли (ЧП) к прибыли до налогообложения (EBT), а отношение последнего финансового индикатора (EBT) к операционной прибыли представляет собой коэффициент процентного бремени. Далее представим пятифакторную модель Дюпона наглядно:

$$ROE = \frac{ЧП}{EBT} \cdot \frac{EBT}{EBIT} \cdot \frac{EBIT}{B} \cdot \frac{B}{A} \cdot \left(1 + \frac{D}{E}\right). \quad (3)$$

Разница между пятифакторной и трехфакторной моделями Дюпона заключается в более детальном рассмотрении рентабельности продаж (ROS), что дает возможность оценить уровень влияния уплачиваемых налогов и процентов на ключевой показатель модели Дюпона – ROE [5].

Для того чтобы продемонстрировать взаимосвязь между всеми описанными выше факторами, автором были рассмотрены финансовые показатели «Пивоваренной компании «Балтика». Данное предприятие является крупнейшим в пивоваренной отрасли России. Были проделаны следующие расчеты на основе показателей финансовой отчетности компании, которые представлены в таблице.

Таблица. Показатели рентабельности активов и собственного капитала
ООО «Пивоваренной компании «Балтика» за 2015–2017 гг.

Показатель	2015	2016	2017
Коэффициент чистой рентабельности, $RONA$, %	18,5	18,5	23,4
Коэффициент оборачиваемости активов, раз	0,8	0,7	0,7
ROA , %	11,2	11,7	15,07
ЭФР	2,93	2,61	3,24
ROE , %	14,9	15,03	19,12

Прослеживается положительная динамика коэффициента рентабельности собственного капитала, так значение показателя в 2017 году составило 19,12% и численно превосходит значение этого же показателя в 2015 году.

Также стоит заметить, что во многом снижение или увеличение рентабельности собственного капитала зависит от эффекта финансового рычага (ЭФР), чем выше значение этого показателя, тем больше доля заемных средств компании.

В заключение хочется сказать, что модель Дюпона, а именно декомпозиция ее ключевых финансовых индикаторов, дает возможность определить и продемонстрировать сравнительную характеристику основных источников изменений тех или иных показателей. Также модель Дюпона позволяет не только выявить причины отрицательных результатов, но и контролировать достижение запланированных стратегических целей компании.

Литература

1. [Электронный ресурс]. – Режим доступа: <https://cyberleninka.ru/article/v/analiz-rentabelnosti-kompanii-na-osnovanii-modeli-dyupon> (дата обращения: 20.01.2019).
2. Павлова И.А. Анализ и прогнозирование финансовой устойчивости организации с учетом жизненного цикла на основе интегрального показателя // Финансы и кредит. – 2017. – № 23. – С. 73–76.
3. Теплова Т.В. Планирование в финансовом менеджменте. – М.: Высшая школа экономики, 2018. – 75 с.
4. Башмаков А.И. Теория систем и системный анализ: учебное издание. – М.: Башмаков А.И., 2015. – 199 с.
5. Li M., Nissim D., Penman S.H. Profitability decomposition and operating risk [Электронный ресурс]. – Режим доступа: <https://pdfs.semanticscholar.org/d158/77845b84dcb1da8eab73b1a0a51a10588bca.pdf> (дата обращения: 06.01.2019).



Купина Анна Павловна

Год рождения: 1996

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K41201

Направление подготовки: 11.04.02 – Программное обеспечение

в инфокоммуникациях

e-mail: anna-kupina@yandex.ru



Осипов Никита Алексеевич

Год рождения: 1972

Университет ИТМО, факультет инфокоммуникационных технологий,
к.т.н., доцент

e-mail: nikita@ifmo.spb.ru

УДК 004.414.22

ОБЗОР ПРОБЛЕМ ИНТЕГРАЦИИ СУЩЕСТВУЮЩИХ МЕССЕНДЖЕРОВ В ЕДИНУЮ СИСТЕМУ ОБРАБОТКИ ОБРАЩЕНИЙ

Купина А.П.

Научный руководитель – к.т.н., доцент Осипов Н.А.

В работе произведен обзор проблем интеграции существующих мессенджеров в единую систему обработки обращений. Описаны ключевые моменты интеграции, с которыми могут возникнуть затруднения, а также сами затруднения, специфичные для описываемых мессенджеров. На основе полученных в процессе обзора данных были сделаны выводы об эффективности интеграции с каждым описываемым мессенджером.

Ключевые слова: мессенджер, интеграция, обработка обращений, мультиканальность, омниканальность, WebAPI.

В современном мире повсеместно распространено использование различных мессенджеров (программ, мобильных приложений или веб-сервисов для мгновенного обмена сообщениями [1]). По некоторым оценкам в начале 2015 года количество активных пользователей мессенджеров превысило количество активных пользователей социальных сетей и продолжает активно расти.

Существует множество мессенджеров, как популярных, так и используемых в определенных кругах. По результатам опросов наиболее популярными мессенджерами в России являются [2] (в порядке убывания, рис. 1 [2]):

1. WhatsApp;
2. Viber;
3. V Kontakte Messenger;
4. Skype;
5. Telegram.

Если говорить о качестве клиентского сервиса с точки зрения коммуникации с клиентами, наиболее прогрессивным способом коммуникации является омниканальная система (рис. 2 [3]). Омниканальность – это объединение всех каналов коммуникации вместе с единой историей обращений и покупок клиентов [3]. В омниканальной системе при «передаче» клиента от одного специалиста к другому клиент не испытывает неудобства и необходимости сменить вид связи. И наоборот: если клиент решит сменить канал связи, то он

будет обслуживаться по тому же стандарту, что и в другом канале связи, а также не будет теряться его история как клиента [4]. Исходя из ранее сказанного про мессенджеры, одна из важных составляющих омниканальной системы – это интеграция с популярными мессенджерами.

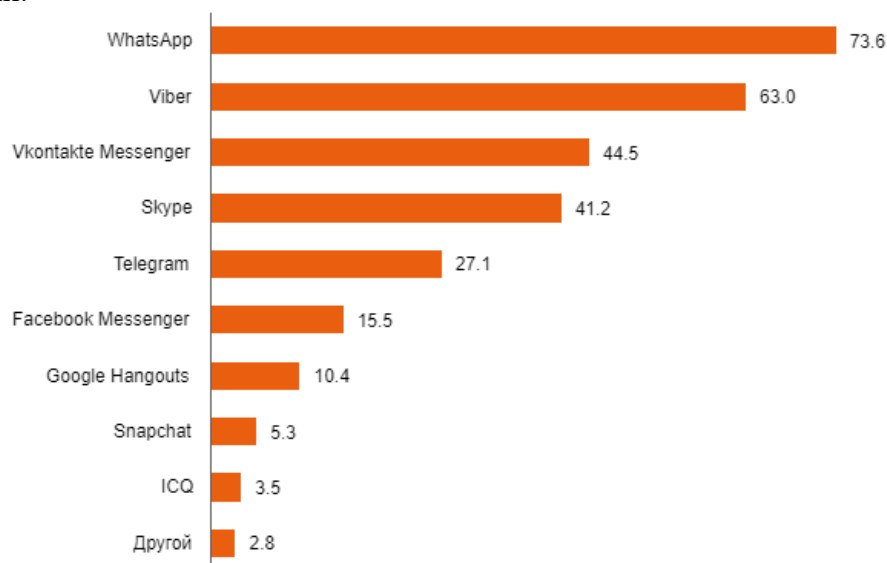


Рис. 1. Статистика использования мессенджеров жителями России (процент респондентов, у которых на смартфоне установлен мессенджер)



Рис. 2. Эволюция коммуникации бизнеса с клиентами

Для каждого мессенджера были определены следующие ключевые моменты, с которыми могут возникнуть проблемы в ходе интеграции. Наличие или отсутствие проблем по каждому пункту влияют на «привлекательность» и эффективность каждого мессенджера с точки зрения интеграции в омниканальную систему:

1. популярность мессенджера;
2. наличие WebAPI (Application Programming Interface или интерфейс программирования приложений – метод получения данных из запросов [5]) или другого механизма доступа к сервисам мессенджера не через обычный пользовательский интерфейс;
3. открытость WebAPI;
4. наличие интеграционной документации и ее актуальность;
5. удобство тестирования интеграции;
6. наличие контроля и фильтрации интеграций со стороны мессенджера;
7. наличие юридических проблем с использованием мессенджера.

В таблице представлен обзор наиболее популярных в России мессенджеров (WhatsApp, Viber, V Kontakte Messenger, Skype, Telegram) по вышеописанным моментам.

Таблица. Обзор особенностей популярных мессенджеров

	WhatsApp [6]	Viber [7]	Vkontakte Messenger [8]	Skype [9]	Telegram [10]
Популярность (% респондентов, у которых установлен)	73,6	63,0	44,5	41,2	27,1
Наличие API	Есть	Есть	Есть	Есть	Есть
Открытость API	Доступно для аккаунтов WhatsApp в Facebook Business Manager	Доступно для публичных аккаунтов фирм	Доступно для любых сообществ ВКонтакте	Доступно для ботов, созданных на Microsoft Bot Framework	Доступно для любых ботов Telegram
Наличие актуальной документации	Да, на сайте developers.facebook.com	Да, на сайте developers.viber.com	Да, на сайте vk.com/dev/	Да, на сайте dev.skype.com	Да, на сайте core.telegram.org/api
Удобство тестирования	Невозможно начать до активации бизнес-аккаунта	Невозможно начать до активации публичного аккаунта	Можно начать разработку и тестирование сразу после создания сообщества (не нужно дополнительных действий для активации сообщества)	Есть быстрый тестовый режим для не более чем 100 пользователей	Тестирование можно начать сразу и без ограничений
Валидация и контроль интеграций	Необходим бизнес-аккаунт, который должен быть и оплачен на определенный период (заявка на формирование т.н. «кредитной линии» формируется заблаговременно)	Необходим публичный аккаунт, который должен пройти предварительную проверку со стороны Viber	Со временем блокируются сообщества, нарушающие правила сайта	Более чем 100 пользователей доступно после валидации бота и проверки его соответствия «certification requirements», предъявляемых со стороны Skype	Отсутствует
Наличие юридических проблем	Отсутствуют	Отсутствуют	Отсутствуют	Отсутствуют	Данный мессенджер и сайт с документацией заблокированы на территории России

Выводы. Можно заметить, что каждый мессенджер обладает своими недостатками, которые способны замедлить процесс интеграции и повлиять на ее итоговую эффективность и прибыльность. Наиболее эффективной с точки зрения охвата аудитории является мессенджер WhatsApp, так как он является наиболее популярным мессенджером в России. Но при этом удобство интеграции с данным мессенджером достаточно низкое из-за сложного доступа к сервисам WhatsApp не через обычный интерфейс, а также из-за неудобства продолжительного тестирования интеграции для многих публичных аккаунтов, так как некоторые из них блокируются из-за подозрения на спам по причине чересчур активного тестирования. Наиболее удобным с точки зрения разработки является V Kontakte Messenger: данный мессенджер имеет открытое и понятное API, понятную и постоянно обновляющуюся документацию, не имеет проблем с созданием публичного аккаунта-лица компании, не ограничивается в использовании законом и не имеет особых ограничений в тестировании. Чуть менее удобным, но наиболее эффективной является интеграция с Viber, так как она обладает всеми теми же удобствами, что и интеграция с V Kontakte Messenger, за исключением легкости создания публичного аккаунта компании – для его создания необходимо заранее написать заявку, чтобы ваш потенциальный аккаунт проверили на добросовестность. Но этот недостаток не так велик по сравнению с тем преимуществом, что Viber является вторым по популярности мессенджером в России.

Литература

1. Обзор мессенджеров. Лучшие и популярные интернет мессенджеры [Электронный ресурс]. – Режим доступа: <http://www.voipoffice.ru/tags/messendzhery> (дата обращения: 06.01.2019).
2. Эпоха мессенджеров: какие из них самые популярные и почему? [Электронный ресурс]. – Режим доступа: <https://www.sostav.ru/publication/kak-vladeltsy-smartfonov-ispolzuyut-messendzhery-30288.html> (дата обращения: 06.01.2019).
3. Зачем нужна омниканальность? [Электронный ресурс]. – Режим доступа: <https://livetex.ru/blog/2017/10/zachem-nuzhna-omnikanalnost/> (дата обращения: 06.01.2019).
4. Омниканальность и мультиканальность: что в имени твоём? [Электронный ресурс]. – Режим доступа: <http://russia.blog.genesys.com/2017/03/29/omnichannel-versus-multichannel-whats-name/> (дата обращения: 06.01.2019).
5. Введение в Web API [Электронный ресурс]. – Режим доступа: <https://metanit.com/sharp/mvc/12.1.php> (дата обращения: 06.01.2019).
6. Facebook Business [Электронный ресурс]. – Режим доступа: <https://www.facebook.com/business/> (дата обращения: 06.01.2019).
7. Viber REST API [Электронный ресурс]. – Режим доступа: <https://developers.viber.com/docs/api/rest-bot-api/> (дата обращения: 06.01.2019).
8. VK Developers [Электронный ресурс]. – Режим доступа: <https://vk.com/dev/> (дата обращения: 06.01.2019).
9. Skype for developers [Электронный ресурс]. – Режим доступа: <https://dev.skype.com/> (дата обращения: 06.01.2019).
10. Telegram APIs [Электронный ресурс]. – Режим доступа: <https://core.telegram.org/api> (дата обращения: 06.01.2019).



Минич Никита Станиславович

Год рождения: 1996

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K41201

Направление подготовки: 11.04.02 – Программное обеспечение
в инфокоммуникациях

e-mail: minichn96@mail.ru



Иванов Сергей Евгеньевич

Год рождения: 1975

Университет ИТМО, факультет инфокоммуникационных технологий,
к.ф.-м.н., доцент

e-mail: 45is@mail.ru

УДК 004.054

ТЕХНОЛОГИИ И МЕТОДЫ АВТОМАТИЗИРОВАННОГО ТЕСТИРОВАНИЯ КОДА ПРИ РАЗРАБОТКЕ РАСПРЕДЕЛЕННЫХ СИСТЕМ

Минич Н.С.

Научный руководитель – к.ф.-м.н., доцент Иванов С.Е.

Работа посвящена обзору методов и технологий тестирования распределенных систем. Объектом исследования данной работы являются методы поисковых систем, оптимизирующих работу сайта.

Ключевые слова: тестирование, методы, распределенные системы, функциональные требования, отказоустойчивость.

Распределенная система – это набор независимых компьютеров, представляющий их пользователям единой объединенной системой [1]. Основная задача, которую пытаются решить с помощью распределенных систем – обеспечение максимально простого доступа к возможно большему количеству ресурсов и как можно большему числу пользователей.

Как и при разработке любого другого программного обеспечения (ПО), распределенным системам необходимо тестирование с целью наглядной демонстрации работы, либо с целью определения возможных изъянов разработанного продукта.

Для тестирования распределенных систем характерно использование классических методов тестирования (рисунков), таких как:

- модульное тестирование. Модульное тестирование – это метод, применяемый отдельными разработчиками ПО для обеспечения того, чтобы мельчайшие, автономные фрагменты кода функционировали так, как задумано, и обеспечивали правильные результаты [2]. Поскольку ручное модульное тестирование требует больших затрат времени и усилий, существует множество инструментов для автоматического запуска модульных тестов;
- интеграционное тестирование. Интеграционное тестирование проверяет согласованность работы отдельных компонентов и подсистем, в том виде, в каком они предоставляются разработчикам приложения. Выделяют два уровня интеграционного тестирования: компонентный и системный. Работа ведется только с компонентами одной системы, в системном ведется работа с различными подсистемами [3];
- системное тестирование. Системное тестирование – это тестирование законченного и полностью интегрированного программного продукта. Обычно ПО является лишь одним из элементов более крупной компьютерной системы. В конечном счете, ПО

взаимодействует с другими программно-аппаратными системами. Системное тестирование на самом деле представляет собой серию различных тестов, единственной целью которых является использование всей системы. Основной задачей системного тестирования является проверка как функциональных, так и не функциональных требований к системе в целом. При этом выявляются дефекты, такие как неправильное использование ресурсов системы, непредусмотренные комбинации данных на пользовательском уровне, несовместимость с окружением, непредусмотренные сценарии использования, отсутствующая или неверная функциональность, неудобство использования и т.д. [4].

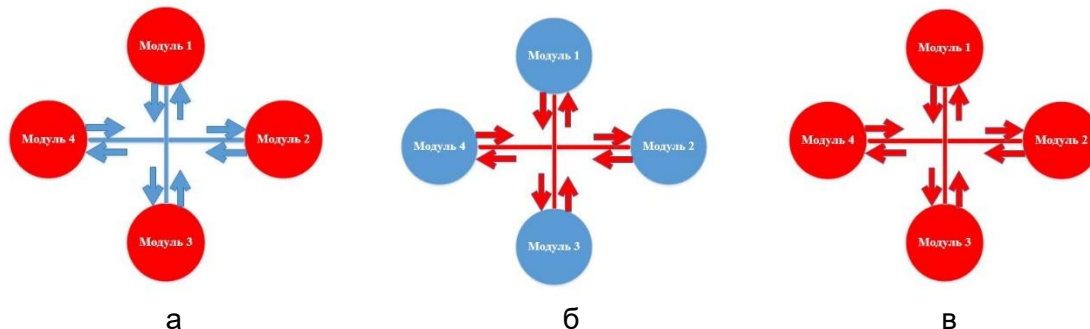


Рисунок. Наглядное представление: модульное тестирование (а); интеграционное тестирование (б); системное тестирование (в)

Тестирование внедрением сбоя (fault injection). Во время работы системы при помощи специальных программ и механизмов добавляются сбои: сбой дисков или целых машин, может быть, сетевые сбои, сбои внутренних компонентов тестируемой системы. По сути, это тестирование на отказоустойчивость, поскольку это одно из важнейших нефункциональных требований к распределенным системам [5]. Это гарантирует, что система способна обрабатывать и восстанавливать данные после сбоя или ошибки, а также выявляет слабые стороны проекта, в которых один сбой потенциально может перерасти в серьезную ошибку или системный сбой. Имеются несколько подходов к тестированию с внедрением сбоя:

- **Chaos Engineering.** При разработке ПО способность конкретной системы ПО выдерживать сбои, в то же время, обеспечивая надлежащее качество обслуживания, которое часто обобщается как отказоустойчивость, обычно указывается в качестве требования. Тем не менее, команды разработчиков часто не выполняют это требование из-за таких факторов, как короткие сроки или отсутствие знаний в этой области. Chaos Engineering – это подход, проверяющий требования к устойчивости системы. Ярким примером Chaos Engineering является Chaos Monkey. Chaos Monkey – это программный инструмент, разработанный инженерами Netflix для проверки устойчивости и восстанавливаемости их веб-сервисов Amazon (AWS). ПО имитирует сбои экземпляров служб, работающих в рамках групп автоматического масштабирования (ASG), путем выключения одной или нескольких виртуальных машин [6]. Данный подход использует рабочую систему и требует определенной зрелости продукта;
- **Jepsen,** который является модулем языка Clojure позволяет эмулировать сбои в тестовой среде. При запуске системы управляющий узел раскручивает набор логически однопоточных процессов, каждый со своим клиентом для распределенной системы. Генератор генерирует новые операции на выполнение для каждого процесса. Затем процессы применяют эти операции к системе, используя своих клиентов. Начало и конец каждой операции записывается в историю. При выполнении операций специальный процесс Немезиды вносит сбои в систему, которые также планируются генератором. Jepsen использует средство проверки для анализа истории теста на правильность, а также для генерации отчетов, графиков и т.д. Jepsen используется для проверки всего: от непротиворечивых коммутативных баз данных до линейризуемых систем координации и распределенных планировщиков задач. Он также может генерировать графики

производительности и доступности, помогая определить, как система реагирует на различные сбои [7].

Формальные методы (FM). Формальные методы состоят из набора методов и инструментов, основанных на математическом моделировании, и логике, которые используются для спецификации и верификации требований для проектирования компьютерных систем и ПО. Использование FM в проекте может принимать различные формы, от случайных математических обозначений, встроенных в английские спецификации, до полностью формальных спецификаций, использующих спецификации языков с точной семантикой.

- **формальная спецификация.** Формальная спецификация – математическое описание программы или оборудования, которое может быть использовано для разработки реализации. В ней описывается, что должна делать система, но не (обязательно) указывается как. Имея такую спецификацию, можно, используя технику формальной верификации, продемонстрировать, что предложенный проект системы является правильным, в отношении спецификации. TLA+ – это язык формальной спецификации. Это инструмент для разработки систем и алгоритмов, а затем программной проверки того, что в этих системах нет критических ошибок. Это программный эквивалент проекта;
- **формальная верификация.** Формальная проверка вычислительной системы влечет за собой математическое доказательство, показывающее, что система удовлетворяет желаемым свойствам или спецификации. Для этого нужно использовать некоторую математическую структуру для моделирования интересующей системы и получения ее желаемых свойств как теоремы о структуре. Verdi framework – это платформа для реализации и формальной верификации распределенных систем в Coq. Verdi формализует различную семантику сети с разными ошибками, и разработчик выбирает наиболее подходящую модель ошибок при проверке их реализации [8].

Литература

1. Таненбаум Э., Ван Стеен М. Распределенные системы. Принципы и парадигмы. – СПб.: Питер, 2003. – 23 с.
2. Dalton J. Great Big Agile. – Apress, Berkeley, CA, 2019. – P. 269–270.
3. Данилов А.Д., Фёдоров А.И. Иерархическая структура процесса тестирования сложного программного обеспечения [Электронный ресурс]. – Режим доступа: <https://cyberleninka.ru/article/v/ierarhicheskaya-struktura-protsessa-testirovaniya-slozhnogo-programmnogo-obespecheniya> (дата обращения: 06.01.2019).
4. Пивень А.А., Скорин Ю.И. Тестирование программного обеспечения // Системы обработки информации. – 2012. – Вып. 4(1). – С. 56–58.
5. Тестирование распределенных систем [Электронный ресурс]. – Режим доступа: <https://habr.com/ru/company/jugru/blog/313908/> (дата обращения: 06.01.2019).
6. Chaos Monkey [Электронный ресурс]. – Режим доступа: <https://whatis.techtarget.com/definition/Chaos-Monkey/> (дата обращения: 06.01.2019).
7. Jepsen [Электронный ресурс]. – Режим доступа: <https://github.com/jepsen-io/jepsen/> (дата обращения: 06.01.2019).
8. Wilcox J.R., Woos D., Panchekha P., Tatlock Z., Wang X., Ernst M.D., Anderson T. Verdi: A Framework for Implementing and Formally Verifying Distributed Systems // In Proceedings of the 36th ACM SIGPLAN Conference on Programming Language Design and Implementation. – 2015. – P. 357–368.

**Пац Карина Михайловна**

Год рождения: 1995

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № К4245сНаправление подготовки: 09.04.03 – Прикладная информатика

e-mail: karina.m.pats@gmail.com

**Сергушичев Алексей Александрович**

Год рождения: 1991

Университет ИТМО, факультет информационных технологий
и программирования, к.т.н., доцент

e-mail: alserg@corp.ifmo.ru

УДК 004.94**СРАВНЕНИЕ МЕТОДОВ РАСЧЕТА СВОБОДНОЙ ЭНЕРГИИ ГИББСА
ДЛЯ КАТАЛИЗАТОРОВ НА ОСНОВЕ ОДНОВАЛЕНТНОГО МАРГАНЦА****Пац К.М.****Научный руководитель – к.т.н. Сергушичев А.А.**

Работа выполнена в рамках темы НИР № 618269 «Автоматизированный поиск переходных состояний в химических реакциях».

Работа посвящена сравнительному анализу методов расчета свободной энергии Гиббса – важнейшего параметра, характеризующего химические реакции. Рассмотрены расчеты с помощью методов квантовой химии и полуэмпирической квантовой химии. Проведено сравнение геометрии полученных в результате оптимизации структур, а также значений их свободных энергий с поправкой на энергию основного состояния. Сделаны выводы о применимости результатов, полученных с помощью полуэмпирических расчетов, в качестве отправной точки для подбора параметров силового поля для расчетов методами молекулярной механики.

Ключевые слова: компьютерная химия, компьютерное моделирование, квантовая химия, каталитическая химия, молекулярная механика, энергия Гиббса.

Свободная энергия Гиббса – важнейший параметр, определяющий возможность и направление протекания химических реакций. Энергия Гиббса может быть рассчитана *ab initio* методами, а именно путем квантовохимических вычислений, или же оценена с помощью полуэмпирических методов расчета или методов молекулярной механики.

Для высокопроизводительного анализа превращений в каталитических путях молекулярная механика является методом выбора в силу ее высокой производительности и приемлемой точности расчетов. Однако затруднения при использовании данного метода вызывает тот факт, что для таких расчетов необходимо силовое поле, параметризованное для конкретной системы (рис. 1). При этом в литературе данных о таком силовом поле нет.

Для разработки силового поля необходимо иметь данные расчетов методами квантовой химии. Такие данные были получены из работы [1], опубликованной сотрудниками лаборатории SCAMT Университета ИТМО. Расчеты производились методом теории функционала плотности (Density Functional Theory, DFT) в программном обеспечении (ПО) Gaussian. Однако данные расчеты крайне трудоемки (приблизительно 24 ч на систему, состоящую из 70 атомов), и воспроизвести их не представляется возможным. Поэтому в

дополнение к исходным данным энергия Гиббса была рассчитана также и с помощью полуэмпирических методов, которые все еще базируются на принципах квантовой химии, но в силу эмпирической составляющей требуют гораздо меньше времени для вычисления (2–3 мин для идентичной системы), хотя и теряют в точности результатов. Данные расчеты проводились методом PM6-D3H4X [2] в ПО МОРАС2016.

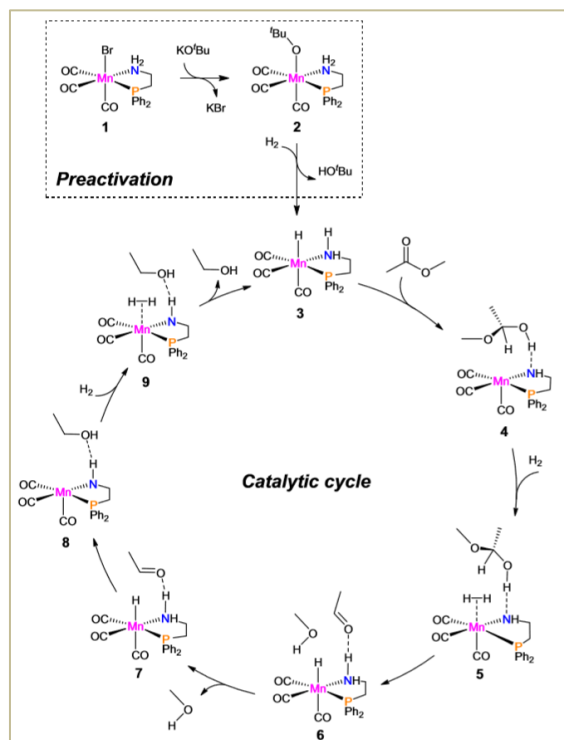


Рис. 1. Схематическое изображение каталитического цикла, для которого необходимо произвести расчет свободной энергии Гиббса [1]

В результате сравнения результатов квантовохимических и полуэмпирических расчетов между ними была обнаружена высокая корреляция (рис. 2). Коэффициент корреляции Пирсона составил 0,951 и 0,997 для компонентов двух каталитических реакций.



Рис. 2. Сравнение значений энергии Гиббса, полученных различными методами вычислений (вверху – квантовая химия, внизу – полуэмпирические расчеты)

Однако данная корреляция была обнаружена только при расчетах энергии Гиббса для марганцевых комплексов. Для малых молекул, участвующих в реакциях, коэффициент Пирсона составил 0,113, что говорит об отсутствии корреляции между результатами расчетов.

Также было проведено сравнение геометрии комплексов после их оптимизации различными методами (табл. 1). Несмотря на более высокое значение разницы в длине связи между азотом и марганцем, все отличия находятся в пределах допустимой погрешности (0,05 Å для ковалентных связей и 0,1 Å для донорно-акцепторных связей).

Таблица 1. Сравнение значений длин связей, полученных оптимизацией геометрии различными методами, на примере одного из соединений каталитического цикла

Связь	Длина l_1 , Å (DFT)	Длина l_2 , Å (PM6-D3H4X)	l_1-l_2 , Å
C ₁ -N	1,472	1,513	0,041
C ₁ -C ₂	1,523	1,527	0,004
N-Mn	2,114	1,998	0,116
C ₁₅ -Mn	1,784	1,805	0,021
C ₁₆ -Mn	1,779	1,798	0,019
C ₁₇ -Mn	1,775	1,784	0,009
C ₂ -P	1,842	1,910	0,068
C ₃ -P	1,826	1,866	0,040
C ₉ -P	1,829	1,887	0,058
Mn-P	2,355	2,297	0,058

Также для сравнения такого параметра, как dG реакции (рассчитывается как разница между энергией Гиббса продуктов и реагентов) была введена поправка на энергию основного состояния (Zero-point energy, ZPE) (табл. 2). Ее значение также было получено в результате проведенных расчетов. Данная поправка позволяет проводить сравнения между разными методами расчетов, поскольку для каждого из них характерна своя шкала измерений, абсолютные значения которых сравнивать не представляется возможным (это также можно увидеть на рис. 2). Согласно полученным результатам, разница параметра dG , рассчитанного методами DFT и PM6-D3H4X находится в пределах погрешности, данные о которой приводятся в литературе [3].

Таблица 2. Сравнение значений свободной энергии Гиббса с поправкой на энергию основного состояния, полученных различными методами вычислений

Реакция	DFT		PM6-D3H4X		$dG_{z1}-dG_{z2}$, кДж/моль
	dG_1 , кДж/моль	dG_{z1} , кДж/моль	dG_2 , кДж/моль	dG_{z2} , кДж/моль	
3+MeCHO→7	19,86	54,31	59,31	40,61	13,69
7→8	-6,50	6,18	-5,07	8,51	-2,33
8+H ₂ →9	78,01	31,03	53,55	31,78	-0,75
9→3+EtOH	-111,02	-52,93	-71,94	-67,59	14,66
3a→4c+MeCHO	16,58	-94,83	-52,17	-49,03	-45,79
3a→6	61,42	27,89	35,53	20,03	7,85
4c+H ₂ →3+MeOH	13,95	-28,41	-29,07	-33,40	4,99
2+H ₂ →3+HOtBu	-10,68	-33,58	-38,56	-44,68	11,10
4c+KOtBu→2+KOMe	32,18	6,01	2,81	8,55	-2,54
3+MeOAc→3a	20,83	59,76	55,58	53,24	6,52
3a→4b	44,55	0,29	4,98	2,03	-1,73
4b→4c+MeCHO	-27,97	-51,40	-57,15	-51,06	-0,34

Таким образом, полученные расчеты обладают достаточной точностью и могут использоваться в качестве отправной точки для моделирования каталитических реакций с помощью методов молекулярной механики, а именно в первую очередь для подбора параметров силового поля для системы катализаторов на основе одновалентного марганца.

Литература

1. Liu C. et al. Computational insights into the catalytic role of the base promoters in ester hydrogenation with homogeneous non-pincer-based Mn-P,N catalyst // *Journal of Catalysis*. – 2018. – № 363. – P. 138–143.
2. Dobeš P.S. et al. Quantum Mechanical Scoring: Structural and Energetic Insights into Cyclin-Dependent Kinase 2 Inhibition by Pyrazolo[1,5-a]pyrimidines // *Curr. Comput.-Aid. Drug.* – 2013. – № 9(1). – P. 118–129.
3. Patra N., Ioannidis E.I., Kulik H.J. Computational investigation of the interplay of Substrate positioning and reactivity in catechol O-methyltransferase // *PLoS ONE*. – 2016. – № 11(8). – P. e0161868.

**Попов Николай Сергеевич**

Год рождения: 1996

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K41211Направление подготовки: 11.04.02 – Программное обеспечение
в инфокоммуникациях

e-mail: i_nikser@mail.ru

**Осипов Никита Алексеевич**

Год рождения: 1972

Университет ИТМО, факультет инфокоммуникационных технологий,
к.т.н., доцент

e-mail: nikita@ifmo.spb.ru

УДК 004.735

**АНАЛИЗ ФРАКТАЛЬНЫХ СВОЙСТВ СЕТЕВОГО ТРАФИКА
В ИНФОКОММУНИКАЦИОННЫХ СИСТЕМАХ****Попов Н.С.****Научный руководитель – к.т.н., доцент Осипов Н.А.**

В работе проводился анализ фрактальных свойств сетевого трафика в инфокоммуникационных системах. Для этого было рассмотрено определение фрактал и методы определения фрактальной размерности. Выполнен аналитический обзор существующих методов рассмотрения сетевого трафика.

Ключевые слова: фрактал, самоподобие, сетевой трафик, фрактальные свойства, геометрические объекты.

Введение. Считалось, что характер сетевого трафика соответствует процессу Пуассона. Сетевой трафик – это передаваемый объем данных через компьютерную сеть за определенное время [1]. С течением времени численность измеряемых и исследованных характеристик сетевого трафика возросла. В результате было подмечено, что не всегда возможно моделировать трафик пакетов в локальной или глобальной сети с помощью процесса Пуассона.

Многочисленные исследования процессов в информационных сетях показали, что поведение сетевого трафика хорошо моделируется с использованием, так называемого самоподобного (фрактального) процесса [2].

Термин «фрактал» был впервые использован в работе Бенуа Мандельброта. Слово фрактал происходит от латинского фразиса, означающего «сломанный» или «разбитый». Мандельброт использовал термин «фрактал» для геометрических объектов, которые имеют сильно фрагментированную форму и могут обладать свойством самоподобия. Фракталы – бесконечно сложные формы, которые являются самоподобными в разных масштабах. Фракталы создаются путем повторения простого процесса снова и снова в непрерывном цикле обратной связи. Фракталами, управляемыми рекурсией, является изображение динамических систем – изображения Хаоса. С одной стороны фракталы могут быть сложными, содержащие бесконечно много элементов, а с другой стороны они могут быть построенные по очень простым законам [3]. Примеры фракталов в природе многочисленны. Например: деревья, горы, кроны деревьев, листья растений и т.д. Если внимательно присмотреться к листку папоротника, то можно заметить, что каждый маленький лист – часть более крупного – имеет

ту же форму, что и весь лист папоротника. Можно сказать, что лист папоротника является самоподобным.

Фрактальные свойства сетевого трафика. Свойство самоподобия ассоциируется с одним из фрактальных типов, т.е. когда меняется масштаб, корреляционная структура фрактального процесса остается неизменной. Фрактальные процессы, описывающие явления в сетевом трафике, обладают рядом свойств. Фрактальные (самоподобные) свойства были найдены в локальных и глобальных сетях, в таких как Ethernet, ATM, TCP, IP, VoIP и приложений видеотрансляции. Фрактальные свойства в инфокоммуникационных потоках, передаваемых клиентами, имеет большую значимость на производительность распределения. Особенно важную роль он играет для услуг, предоставляемых отправку мультимедиа и трафика в реальном времени.

Часто в инфокоммуникационных системах, измерения характеристик наборов данных трафика показывают стохастические фрактальные свойства. Теория фрактальных стохастических процессов не так развита, как теория пуассоновских процессов. Однако учитывая, что фрактальные модели (фигуры) более отлично характеризуют сетевой трафик, чем модели пуассоновского процесса, создание инструментов для познания фрактальных процессов и синтеза искусственного сетевого трафика стало приоритетной задачей, отображающей основные характеристики этих процессов. Самоподобие в трафике – это адаптивность трафика в сети. Многие факторы участвуют в создании этой характеристики. Такое представление является удобным по сравнению с теорией, используемой в большинстве современных публикаций, в которой предполагается, что самоподобие трафика основывается исключительно на распределении файлов по размеру (с большими хвостами).

Многие экспериментальные данные имеют фрактальную статистику, моделирование и анализ которой может быть получено с использованием методов фрактального анализа. Существует несколько методов определения фрактальной размерности для временного ряда:

1. определение фрактальной размерности методом клеточного покрытия временного ряда, когда график покрывается несколькими сетками, а фрактальная размерность определяется так же, как и для геометрических фракталов. Примерами геометрических фракталов являются: снежинка Коха, Т-квадрат, треугольник Серпинского и т.д.;
2. изучение фрактальных временных рядов, представленных Бенуа Мандельбротом, создавший множество Мандельброта, основанных на исследованиях, проведенных английским исследователем Херстом, и называется методом нормированного размаха или R/S-анализом. Он построен на анализе диапазона параметров (наибольшего и наименьшего значения в исследуемом промежутке) и среднеквадратичного отклонения. R/S-анализ сложный процесс, так как надо переработать большое количество данных;
3. метод, основанный на изменении длины кривой в зависимости от масштаба. Если кривая близка к фрактальной, то при уменьшении масштаба длина кривой будет увеличиваться степенным образом [4].

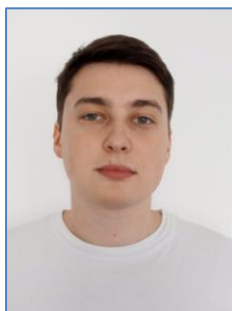
Исключительная значимость фрактального анализа временных рядов заключается в том, что учитывается поведение системы не только в течение периода измерения, но и его предыстория.

Заключение. В настоящее время отсутствуют систематизированные исследования для оценивания влияния фрактальных свойств мультимедийных видеоприложений. Измерение и управление трафиком важны во всех аспектах проектирования сетей. Учет фрактальных свойств сетевого трафика даст возможность более четко отобразить и воспроизвести мультимедийные видеоприложения, что, в свою очередь, даст возможность получения указанных параметров качества обслуживания. Измерение и управление трафиком важны во всех аспектах проектирования сетей. Этот тезис взял новую нерешенную проблему в этой области: как проанализировать фрактальные характеристики сетевого трафика. Поэтому, при

изучении свойств самоподобия мультимедийных видеоприложений, их влияние на характеристики инфокоммуникационных систем, является очень актуальным [5].

Литература

1. Сетевой трафик [Электронный ресурс]. – Режим доступа: <https://dic.academic.ru/dic.nsf/ruwiki/1866122> (дата обращения: 06.01.2019).
2. Самоподобный трафик [Электронный ресурс]. – Режим доступа: <http://studepedia.org/index.php?vol=1&post=93011> (дата обращения: 06.01.2019).
3. Применение фракталов [Электронный ресурс]. – Режим доступа: <http://www.michurin.net/fractals/app.html> (дата обращения: 06.01.2019).
4. Studfiles [Электронный ресурс]. – Режим доступа: <https://studfiles.net/preview/4198182/> (дата обращения: 06.01.2019).
5. Техносфера [Электронный ресурс]. – Режим доступа: <http://tekhnosfera.com/issledovanie-fraktalnyh-svoystv-potokov-trafika-realnogo-vremeni-i-otsenka-ih-vliyaniya-na-harakteristiki-obsluzhivaniya-> (дата обращения: 06.01.2019).



Пряничников Александр Александрович

Год рождения: 1994

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K41201

Направление подготовки: 11.04.02 – Программное обеспечение
в инфокоммуникациях

e-mail: pryaniq@gmail.com



Зудилова Татьяна Викторовна

Год рождения: 1946

Университет ИТМО, факультет инфокоммуникационных технологий,
к.т.н., доцент

e-mail: zudilova@ifmo.spb.ru

УДК 004.414.28

**ВОЗМОЖНОСТИ ВЗАИМОДЕЙСТВИЯ С WI-FI СЕТЯМИ В МОБИЛЬНОЙ
ОПЕРАЦИОННОЙ СИСТЕМЕ IOS**

Пряничников А.А.

Научный руководитель – к.т.н., доцент Зудилова Т.В.

В работе исследованы принципы работы имеющихся в мобильной операционной системе iOS инструментов, позволяющих организовать сетевой доступ. Проведен анализ возможностей реализации сервисных функций подключения к Wi-Fi сетям с использованием данных инструментов с целью создания концепции и прототипа приложения, позволяющего обеспечить доступ к беспроводной сети университета.

Ключевые слова: iOS, Wi-Fi, беспроводная сеть, мобильное устройство, Hotspots Helper, Network Extensions.

На сегодняшний день мобильные устройства закрепили свою позицию в жизни человека и стали ее неотъемлемым атрибутом. С каждым годом портативные устройства развиваются стремительным темпом, они становятся меньше, мощней, удобней и доступней. Мобильные гаджеты, все чаще, в использовании заменяют целые сложные технические устройства, такие как фотоаппарат, телевизор и даже компьютер. Большая часть используемых функций мобильных устройств – социальные сети, медиа ресурсы, поиск информации, мессенджеры и др., представленные на рис. 1 – так или иначе, связана с использованием интернета [1]. Одним из скоростных источников соединения к сети доступа с мобильного гаджета на сегодняшний день является беспроводная сеть доступа Wi-Fi.

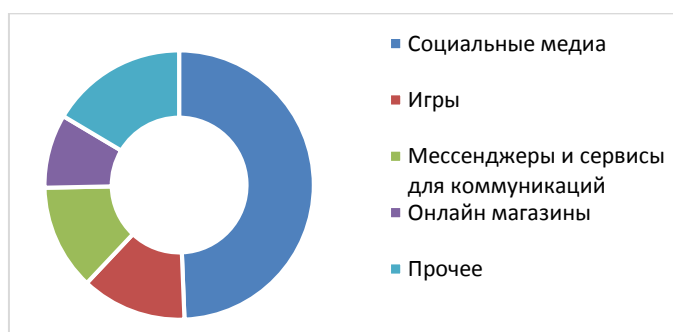


Рис. 1. Использование приложений на мобильных устройствах

Необходимость обеспечения должного уровня защищенности при использовании данных сетей влечет за собой уменьшение удобства подключения к сети для конечного пользователя. Для того чтобы улучшить пользовательский опыт, в данной работе была поставлена цель произвести анализ особенностей организации сетевого доступа в мобильной операционной системе (ОС) iOS для реализации функций Wi-Fi подключения на примере университетской сети.

В исследуемой ОС iOS основной библиотекой, позволяющей организовывать сетевой доступ, является Network Extension, а в частности, ее компоненты Hotspot Helper и Hotspot Configuration [2, 3].

Возможности, предоставляемые данными инструментами:

- возможность участия в авторизации пользователя в Wi-Fi сети на всех этапах – начиная с отображения списка доступных сетей и заканчивая автоматическим отключением от сети;
- обеспечение обработки команд при любом состоянии приложения, в том числе если пользователь завершил его использование;
- начиная с версии iOS 11, имеется возможность обеспечения соединения к Wi-Fi сетям прямо из приложения [4].

Известные ограничения:

- для разработки необходимы оплаченный аккаунт разработчика и разрешение от Apple, запрашиваемое на портале для разработчиков;
- ограничения программной реализации, отраженные в документации [5].

На основе изложенных возможностей и ограничений инструментов библиотеки Network Extension возможен следующий перечень примеров реализации сервисных функций подключения к Wi-Fi сетям университета в iOS-приложении:

- самодостаточное приложение или часть другого приложения, например, личного кабинета студента;
- функции автоматического и ручного режимов, последний включает в себя дополнительные конфигурации, такие как настройка времени существования сессии подключения;
- подключения к различным Wi-Fi сетям в зависимости от роли пользователя.

Результатом применения данных примеров реализации является концепция части функциональности мобильного приложения, обеспечивающего доступ к университетским Wi-Fi сетям, включающая несколько этапов (рис. 2):

1. аутентификация в личном кабинете студента информационной системы управления (ИСУ) Университета ИТМО;
2. выбор режима подключения – ручной, позволяющий изменить стандартные конфигурации подключения, которые были представлены выше, и автоматический;
3. осуществление подключения к Wi-Fi сети с прохождением всех необходимых этапов аутентификации без участия пользователя.

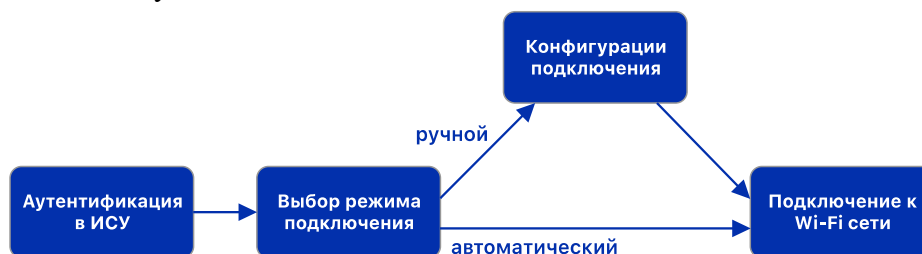


Рис. 2. Этапы подключения к Wi-Fi сети через приложение

Для представления графического интерфейса пользователя в мобильном приложении, на основе концепции, был разработан прототип приложения, включающий следующие экраны взаимодействия с пользователем, на этапе настройки подключения (рис. 3):

1. выбор режима подключения к сети;
2. пример использования ручного режима с выбором конкретной сети;

3. настройки подключения;
4. системное окно подтверждением запуска процесса подключения к Wi-Fi сети.

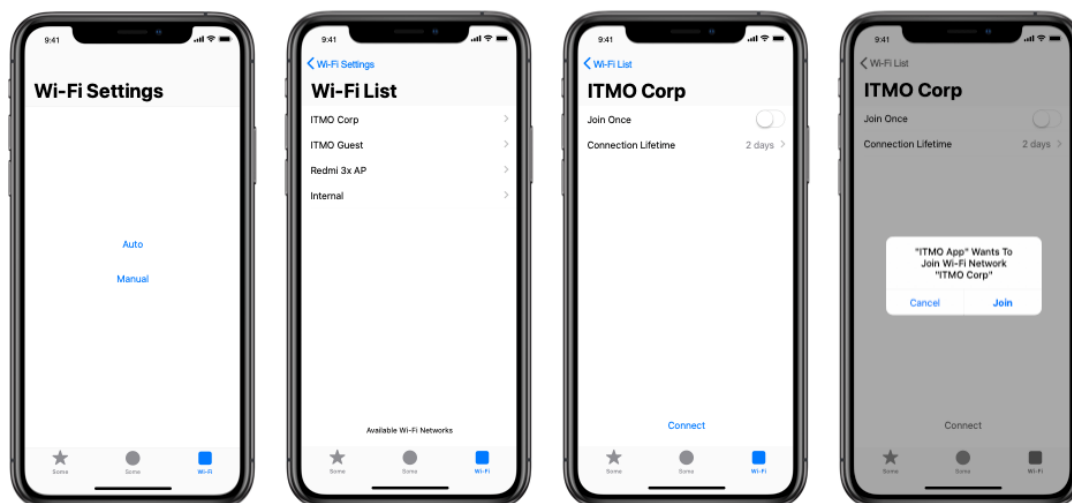


Рис. 3. Прототип пользовательского интерфейса приложения

В работе удалось произвести анализ особенностей организации сетевого доступа в iOS, выделить основные возможности практического применения, разработать концепцию и прототип мобильного приложения, реализующего сервисные функции Wi-Fi соединения. Данный результат показал реальность упрощения пользовательского опыта и повышения безопасности обеспечения доступа к беспроводным сетям за счет прохождения всех этапов авторизации и аутентификации внутри мобильного приложения. Дальнейшими этапами развития данной задачи являются анализ особенностей организации доступа к Wi-Fi сети университета и разработка приложения с минимальными функциями подключения к данной сети.

Литература

1. Mobile App Usage Statistics 2018 [Электронный ресурс]. – Режим доступа: <https://themanifest.com/app-development/mobile-app-usage-statistics-2018> (дата обращения: 27.01.2019).
2. Technical Q&A QA1942 iOS Wi-Fi Management APIs [Электронный ресурс]. – Режим доступа: https://developer.apple.com/library/archive/qa/qa1942/_index.html (дата обращения: 27.01.2019).
3. WWDC Session "Advances in Networking, Part 1" [Электронный ресурс]. – Режим доступа: <https://developer.apple.com/videos/play/wwdc2017/707/> (дата обращения: 06.01.2019).
4. NEHotspotHelper Official Documentation [Электронный ресурс]. – Режим доступа: <https://developer.apple.com/documentation/networkextension/nehotspothelper> (дата обращения: 06.01.2019).
5. WWDC Session "What's New in Network Extension and VPN" [Электронный ресурс]. – Режим доступа: <https://developer.apple.com/videos/play/wwdc2015/717/> (дата обращения: 11.01.2019).

**Синицына Анастасия Александровна**

Год рождения: 1995

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № К4232Направление подготовки: 27.04.03 – Системный анализ

и интеллектуальное управление в бизнесе

e-mail: siniz.tas@yandex.ru

Панова Анастасия Юрьевна

Год рождения: 1965

Университет ИТМО, факультет инфокоммуникационных технологий, к.э.н., доцент

УДК 33.338.3**АНАЛИЗ МЕТОДИК ОЦЕНКИ ЭКОНОМИЧЕСКОЙ БЕЗОПАСНОСТИ
ПРЕДПРИЯТИЯ****Синицына А.А.****Научный руководитель – к.э.н., доцент Панова А.Ю.**

В работе рассмотрены популярные методики оценки экономической безопасности предприятия, такие как индикаторная, ресурсно-функциональная, интегральная, мультипликативная, и рассмотрены детерминанты мезо-, микро- и макроэкономических уровней. Проведен анализ каждой из них, выявлены слабые и сильные стороны и сделаны соответствующие выводы.

Ключевые слова: экономическая безопасность предприятия, оценка, методики, индикаторы, алгоритм, показатели.

В настоящее время экономическая безопасность является актуальной проблемой как для государства в целом, так и для отдельных хозяйствующих субъектов. Однако следует понимать, что решение основных задач на макроуровне невозможно без должного внимания к проблемам многообразных организаций. Основанием для данного вывода являются такие факторы угроз экономической безопасности страны как: низкая конкурентоспособность продукции отечественных предприятий, приобретение иностранными фирмами российских фирм и др. Поэтому очень важно, чтобы хозяйствующие субъекты могли эффективно заниматься мониторингом экономической безопасности, чтобы вовремя устранять существующие угрозы и предпринимать меры по их нейтрализации, что, в свою очередь, повышало бы эффективность и функционирование предприятий.

Для того чтобы качественно проводить мероприятия по мониторингу экономической безопасности предприятий, необходимо адекватно оценивать факторы внешней и внутренней среды организации. Стоит отметить, что на сегодняшний день общей методики оценки экономической безопасности предприятия (ЭБП) не существует, тем не менее, есть множество подходов для оценки ЭБП, которые будут проанализированы ниже.

Первая рассмотренная методика, которая была предложена В.К. Сенчаговым – это индикаторная [1]. Суть этой методики заключается в том, чтобы использовать систему индикаторов, позволяющих эффективно идентифицировать опасности и угрозы. Под индикаторами понимаются нормативные показатели, характеризующие стратегические аспекты деятельности (финансовые, социальные, производственные индикаторы, индикаторы взаимоотношений с контрагентами). Для установления уровня ЭБП показатели проверяют на степень их соответствия пороговым значениям.

Положительные стороны данной методики: индикаторы рассматриваются в разрезе всех функциональных составляющих предприятия, также полученные результаты можно наглядно

представить в форме диаграмм, что облегчает работу сотрудникам, которые оценивают экономическую безопасность предприятия.

Отрицательные стороны: во-первых, на данный момент отсутствует методическая база для точного определения уровня индикаторов, во-вторых, так как показатели взаимосвязаны, а в деятельности предприятий постоянно происходят изменения, то уточнение величины индикаторов становится существенной проблемой, в-третьих, отсутствует сравнение с поведением и состоянием системы в будущем.

Таким образом, возникает необходимость в постоянной корректировке индикаторов, что, в свою очередь, приводит к большим трудозатратам. Также следует понимать, что от того, насколько точно будут определены индикаторы, зависят управленческие решения на предприятии. По мнению автора, данная методика хорошо подходит для малых предприятий, так как нет огромного количества данных, которые необходимо постоянно корректировать.

Следующая методика – ресурсно-функциональная. Методика основана на оценке показателей состояния использования корпоративных ресурсов по особым критериям. Положительной стороной данного подхода является то, что корпоративные ресурсы в данной методике рассматриваются как факторы бизнеса, используемые владельцем и менеджментом, т.е. рассматриваются со стороны целей предприятия. Однако чтобы оценить данную методику, предприятие должно обладать технологическим потенциалом, конкурентоспособностью, финансовой устойчивостью, обладать эффективным менеджментом и т.д. Также отрицательной стороной является, как и в индикаторной методике, не прогнозируемость результатов. Помимо прочего, ресурсно-функциональный подход не способен отобразить все составляющие ЭБП.

Интегральная методика оценки экономической безопасности, которая совокупно объединяет обычные финансово-экономические показатели, была предложена Р.М. Арзумановым [2]. При данном подходе можно использовать несколько уровней интеграции показателей, кластерный и многомерный анализ. Данный подход предполагает расчет пяти групп показателей: рентабельности, финансовой устойчивости, деловой активности, эффективного использования имущества и инвестиционной привлекательности.

Положительной стороной данного подхода является удобство использования большого количества алгоритмов на основе интегральных показателей.

К минусам можно отнести то, что данный подход больше подходит для математиков, а не для управленцев и менеджеров, поскольку эта методика очень сложна для оценки экономической безопасности, так как необходимо прибегнуть к математическому анализу.

Помимо выше перечисленных методик, разными авторами предлагается еще множество подходов к оцениванию ЭБП.

Так, А.В. Шохнех предлагает определять уровень экономической безопасности с помощью мультипликативной оперативной оценки [3]. Из его идеи можно сделать следующий вывод: оценка производится по вполне доступной информации, однако этого недостаточно, чтобы дать оценку ЭБП.

Также М.Т. Гильфанов предложил методику оценки, базирующуюся на характеристики детерминантов мезо-, микро-, макроэкономических уровней [4]. Достоинством данного подхода является возможность проанализировать и оценить влияние различных факторов на уровень ЭБП. Но такая методика требует большого объема экспертных оценок, а также не позволяет оценить прогнозируемое состояние предприятия.

Проанализировав лишь некоторые из многих методик оценки экономической безопасности предприятия, можно сделать следующую сравнительную таблицу.

Таблица. Сравнительная таблица методик оценки ЭБП

Название методики	Положительные стороны	Отрицательные стороны
Индикаторная	– предоставляются подробные данные; – наглядная форма	– нет методической базы; – сложность расчетов из-за изменений данных в режиме реального времени; – не отображает изменения в будущем
Ресурсно-функциональная	– рассматривается в разрезе корпоративных ресурсов и функционала предприятия	– используется только со стороны функционирования предприятия; – нет прогнозируемости результатов
Интегральная	– возможность использования множества алгоритмов	– сложность расчетов; – нет алгоритма для оценки всей ЭБП
Мультипликативная оперативная оценка	– легко достать необходимые данные	– невозможно определить ЭБП с помощью данной методики (недостаточно данных)
Детерминанты мезо-, микро- и макроэкономических уровней	– возможность проанализировать и оценить влияние различных факторов на ЭБП	– требуется большой объем экспертных оценок

Таким образом, рассмотрев несколько методик оценки экономической безопасности предприятия, автор предлагает использовать различные методики в зависимости от целей предприятия. Если необходимо дать полную оценку ЭБП, то автор советует использовать индикаторную и ресурсно-функциональную методику, так как они в полной мере позволяют провести анализ всех функциональных составляющих экономической безопасности. Если на предприятии необходимо рассчитать, к примеру, только финансовую составляющую экономической безопасности, то желательно использовать интегральный подход. Результаты оценок по вышеперечисленным методикам ЭБП должны отображать действующий уровень экономической безопасности.

Литература

1. Сенчагов В.К. Методология обеспечения экономической безопасности [Электронный ресурс]. – Режим доступа: <https://cyberleninka.ru/article/n/metodologiya-obespecheniya-ekonomicheskoy-bezopasnosti> (дата обращения: 06.01.2019).
2. Арзуманов Р.М. Методология организации системы оценки экономической безопасности предприятия [Электронный ресурс]. – Режим доступа: http://posnauka.org/assets/arzumantov_metodologiya_organizatsii_s.pdf (дата обращения: 06.01.2019).
3. Гильфанов М.Т. Инструментарий оценки и обеспечения экономической безопасности предприятия [Электронный ресурс]. – Режим доступа: https://new-dissert.ru/_avtoreferats/01006737252.pdf (дата обращения: 06.01.2019).
4. Шохнех А.В. Математические методы оценки экономической безопасности хозяйствующих субъектов [Электронный ресурс]. – Режим доступа: <http://uecs.ru/uecs> (дата обращения: 06.01.2019).



Смирнов Артем Евгеньевич

Год рождения: 1991

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № К41201

Направление подготовки: 11.04.02 – Программное обеспечение
в инфокоммуникациях

e-mail: artem1403@gmail.com



Зудилова Татьяна Викторовна

Год рождения: 1946

Университет ИТМО, факультет инфокоммуникационных технологий,
к.т.н., доцент

e-mail: zudilova@ifmo.spb.ru

УДК 004.738

АНАЛИЗ СОВРЕМЕННЫХ РЕШЕНИЙ ДЛЯ ПОСТРОЕНИЯ СЕТИ ИНТЕРНЕТА ВЕЩЕЙ

Смирнов А.Е.

Научный руководитель – к.т.н., доцент Зудилова Т.В.

Работа посвящена сравнению широко используемых технических решений для построения сети интернета вещей в масштабах квартиры и в масштабах города. Проведен сравнительный анализ существующих технических решений, и выбраны наиболее подходящие решения, в зависимости от необходимого масштаба сети интернета вещей.

Ключевые слова: интернет вещей, умный город, умный дом.

Для того чтобы внедрение интернета вещей было разумным и современным, прежде всего, необходимо уходить от проводных систем. Существующие около 100 лет системы проводной коммутации устройств постепенно исходят на нет. В условиях активно внедряемой цифровой экономики, а также при постепенном отказе от проводных устройств, актуально рассмотреть решения, на основе которых можно осуществлять подключение устройств к Интернету.

Однако бездумное и опрометчивое подключение устройств к Интернету, без изучения пропускной способности, а также без оценки коммерческой составляющей внедрения и эксплуатации сети, нельзя со стопроцентной уверенностью заявить, что та или иная технология подходит для определенных условий.

Цель работы – анализ существующих технических решений для построения сети интернета вещей и выбор адекватных методов и средств для оценки технологии с точки зрения функциональной применимости в сетях различного масштаба.

Существует множество беспроводных технологий для построения подобных сетей, ниже приведены наиболее широко применяемые:

- Wi-Fi;
- Bluetooth;
- ZigBee;
- GSM-900/GSM-1800;
- LTE;
- NarrowBand Internet of Things (NB-IoT);
- LoRaWAN.

Имеется множество протоколов, называемых Wi-Fi, благодаря которым осуществляется функционирование беспроводных сетей. Диапазон частот варьируется от 2,4 до 5 ГГц. Также меняется и скорость передачи данных от 11 Мбит/с до 320 Мбит/с. Максимальная дальность передачи сигнала будет изменяться от 100 м до 300–400 м в зависимости от внешних условий, таких как количество сооружений, находящихся в окружении работы источника сигнала [1].

- Wi-Fi используется для работы с серверами, хранящими базы данных или программные приложения, позволяет выйти в Интернет обычному пользователю. Развертывание сети происходит благодаря установке Wi-Fi роутера, который будет подключаться к глобальной сети, посредством различных проводных технологий.

Он широко используется в жизни и быту граждан в любой точке мира, владельцы кафе и различных других заведений часто предлагают бесплатный интернет, производя раздачу посетителям по сети Wi-Fi. Также он широко используется дома для подключения всевозможных гаджетов, так как обладает широкой пропускной способностью, его использование стало повсеместным, и в ближайшей перспективе на горизонте нет серьезных конкурентов. Однако такое соединение является сильно энергозатратным, что не позволяет дать большое время работы и энергонезависимость при проектировании сети интернета вещей [2].

- Bluetooth широко используется для обмена информацией между такими устройствами, как компьютеры, смартфоны, планшеты, принтеры, цифровые фотоаппараты, устройства ввода-вывода, аудиоустройства, в том числе гарнитуры и акустические системы. Bluetooth обладает маленьким радиусом действия, нормальное функционирование устройств появляется, когда радиус не превышает 240 м (протокол Bluetooth 5.0), и при этом очень зависит от помех и преград. Нарушение функционирования может наблюдаться даже в соседних комнатах кирпично-монолитного жилого дома.

Работает Bluetooth в диапазоне частот 2,402 ГГц–2,48 ГГц, что является не лицензируемыми частотами, и не требует регистраций и разрешений. С Bluetooth 5.0 устройства могут использовать скорость передачи данных до 2 Мбит/с [3]. Основным преимуществом является то, что в отличие от многих других типов сетей, сеть Bluetooth не нуждается в сервере, маршрутизаторе и другом оборудовании, необходимом для функционирования беспроводных сетей.

- ZigBee – беспроводная связь между устройствами с низким потреблением, с возможностью выстраивания ячеистой топологии сети и семейства IEEE 802. Стандарт дает возможность использовать каналы в нескольких частотных диапазонах: 868 МГц, 915 МГц и 2,4 ГГц. Максимальная скорость передачи данных и адекватная помехозащищенность достигается в диапазоне 2,4–2,48 ГГц.

ZigBee – сеть с низким потреблением, это достигается благодаря жесткому ограничению максимальной скорости. Средняя скорость передачи данных, в зависимости от загрузки сети, составляет от 5 до 40 кбит/с [4].

Данная сеть быстра и проста в развертывании, а благодаря высокой энергоэффективности, она допускает долгую автономную работу. Радиус действия рабочей станции составляет 10–20 м внутри помещения, а на открытом воздухе несколько сотен метров. Но для эффективного построения сети интернета вещей для города, необходимо чтобы маршрутизаторы, используемые для построения сети, имели большой радиус вещания. Данный факт является крайне важным с экономической точки зрения.

- GSM-900 – цифровой стандарт мобильной связи в диапазоне частот 890–915 МГц (от устройства к базовой станции) и 935–960 МГц (от базовой станции к устройству), и в случае с GSM-1800 в диапазоне частот 1710–1785 МГц (от устройства к базовой станции) и 1805–1880 МГц (от базовой станции к устройству) [5].

Это привычная и понятная сеть по работе мобильных телефонов и смартфонов. Лицензируемые радиочастоты, работы в данных частотах осуществляет ограниченное количество операторов. Далее рассмотрена GSM-1800 как более современная сеть.

Наблюдается высокая емкость сети, но основной нюанс, ограничивающий использование для интернета вещей, состоит в том, что зона охвата для каждой базовой станции значительно меньше, чем на более низких частотах. Тарифная сетка операторов большой тройки начинается от 100 руб. за устройство.

Такое ценообразование обусловлено безумно высокой стоимостью развертывания сети, сумасшедшей стоимостью базовых станций, а также высокими дополнительными расходами, такими как электроэнергия и обслуживание. Также не стоит забывать, что это лицензируемые частоты, а значит, установка оборудования требует разрешений и согласований. Срок сбора документации и разрешений для установки базовой станции для работы в сети GSM составляет от 6 месяцев до 1 года.

- LTE является стандартом беспроводной передачи данных и развитием стандартов GSM/UMTS. Главная цель появления LTE – это наращивание пропускной способности и скорости с использованием нового метода цифровой обработки сигналов и модуляции. Также LTE позволило реконструировать и упростить архитектуру сетей, основанных на IP, значительно уменьшив задержки при передаче данных по сравнению с архитектурой 3G-сетей.

Спецификация LTE обеспечивает скорость загрузки до 326,4 Мбит/с, скорость отдачи до 172,8 Мбит/с. На территории Российской Федерации операторы, предоставляющие возможность взаимодействия с сетью, работают на частоте 2600 МГц, это значительно более высокая частота, чем частоты на GSM900-1800, а значит, радиус охвата будет еще более низкий. Также от предшественника передалась все остальные свойства – высокая энергопотребляемость и лицензируемость частот.

Несмотря на это, данная сеть имеет просто потрясающие параметры в плане мобильности. Сеть LTE позволяет нормально функционировать устройствам на скорости до 500 км/ч. Данное обстоятельство позволило решить проблему с мобильной связью в скоростных поездах. Также у сети наблюдается крайне низкая задержка при передаче данных, что было бы полезно, не обладай она вышеперечисленными характеристиками.

- NB-IoT – стандарт сотовой связи для устройств телеметрии с низкими объемами обмена данными. Технология функционирует на базе сетевого оборудования поверх сети LTE на частотах 1800 МГц на базе принципов и подходов LTE для обеспечения подключений устройств IoT [6]. Это нюанс может быть актуален для маломощных модемов и встраиваемых модулей в рамках проекта «Умный город». Большим плюсом является то, NB-IoT может быть реализован на существующей сети LTE с минимальными дополнительными инвестициями.

Операторы большой тройки, начиная с 2018 года, активно используют и тестируют NB-IoT. NB-IoT специально создан для интернета вещей, поэтому учитывает специфические потребности, такие как улучшенная чувствительность к модуляции сигнала для подключения сотен тысяч устройств. Также для работы с NB-IoT не нужна SIM-карта, и достаточно небольшой мощности приемопередатчика, т.е. устройства IoT могут работать много лет от одной батарейки.

Единственное ограничение, которое накладывается на данную сеть – это то, что для функционирования NB-IoT необходимо наличие покрытия LTE. В дополнение к этому на устройства NB-IoT косвенно ложится стоимость развертывания LTE-сетей, в том числе работа в лицензируемой частоте, что тоже требует дополнительных вложений, а также не дает свободно развернуть собственную сеть для интернета вещей.

- LoRaWAN является технологией беспроводной передачи данных, работающей на физическом уровне модели OSI. Ключевыми особенностями технологии LoRa являются ее высокая энергоэффективность и дальность связи. Устройства, использующие LoRa, работают в нелицензируемых ISM (Industrial, Science, and Medical – промышленные, научные и медицинские) частотах, в России такой частотой является 868 МГц [7]. Работа в нелицензируемом диапазоне частот снимает с владельца сети правовые ограничения в план

развертывания. Благодаря этому у устройств, созданных для работы в сети LoRaWAN, наблюдается низкая стоимость. Данная сеть, как и NB-IoT, была специально создана для интернета вещей.

Обширное применение получили устройства, которые передают небольшой объем данных. LoRaWAN является сетью с небольшим трафиком. Объем одного пакета достигает лишь не более 255 Б.

Самое популярное применение – это автономные датчики, работающие годы от элемента питания, которые могут находиться на большом расстоянии от базовой станции. В условиях плотной городской застройки дальность сигнала составляет до 3 км, а на незастроенной территории до 10 км. В идеальных условиях дальность может достигать до 30 км, например при расположении базовой станции на воздушном шаре. Данная практика уже испытывается в Гонконге и Франции.

Развертываемость и настройка сети LoRaWAN будет занимать минимум времени ввиду того, что оборудование полностью готово для обычного пользователя. Для обширных территорий это будет также удобно, как в последние 10 лет стал удобен Wi-Fi для дома и квартиры.

Каждая исследуемая технология представляет собой фундамент для развития новых технологий в мире, и каждая имеет свои плюсы и минусы. В сводной таблице представлены основные характеристики, на которые стоит обращать внимание при построении собственной сети интернета вещей. Для построения сети, требующей обработку большого количества данных и на большом расстоянии, удобно использовать технологию LTE и GSM900-1800 в зависимости от технической возможности. Но данные технологии ограничены большим количеством лицензий и разрешений от государственных органов. Данные технологии невозможно использовать отдельно от операторов большой тройки, что создает невозможность построения полностью своей сети.

Таблица. Сравнение исследуемых технологий

	Wi-Fi 2,4 МГц	BT	ZigBee	GSM-1800	LTE	NB-IoT	LoRaWAN
Частота, ГГц	2,4	2,4	2,4	1,8	2,6	1,8	0,86
Скорость, ед/с	320 м	2 м	40 кБ	42 м	326 м	50 кБ	50 кБ
Объем пакета	бол.	мал.	мал.	бол.	бол.	мал.	мал.
Радиус охвата, м	30	240	20	3000	400	400	3000
Помехоуст.	нет	нет	да	да	да	да	да
Проприет.	нет	нет	нет	да	да	да	нет
Стоимость	\$	\$	\$	\$\$	\$\$\$	\$\$\$	\$

На рисунке представлены исследуемые решения в отношении пропускной способности к дальности передачи данных.

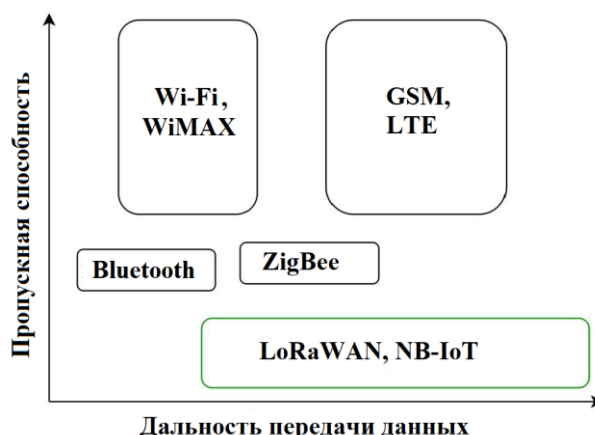


Рисунок. Сравнение исследуемых технологий

Исходя из этого, можно сделать вывод, что технологии Wi-Fi, а также ее аналог ZigBee отлично подходят для проектирования и построения сети интернета вещей в рамках «своего дома» или небольшого офиса. Но в перспективе развития и увеличения устройств, и расстояния между ними, покажет себя нефункционально и нерентабельно.

Bluetooth здесь представлена для дополнительного анализа в сравнении с технологией Wi-Fi. Данная технология отлично выполняет свои функции в плане подключения устройств к смартфону, но не пригодна для построения сети интернета вещей в разрезе целого города и даже дома.

Набирающая популярность технология NB-IoT была специально создана для интернета вещей, но она ограничена лицензией на частоту и дорогостоящими базовыми станциями операторов большой тройки, что также делает невозможность ее простого использования.

В заключение рассмотрена технология LoRaWAN для построения собственной сети. Нелицензируемость частот, а также дешевизна и простота в использовании, позволяют в короткие сроки спроектировать функционирующую сеть интернета вещей. Обширный радиус действия позволяет успешно внедрять устройства по диспетчеризации, управлению и сбору статистики в рамках целого города и даже региона.

Литература

1. Wi-Fi 6 [Электронный ресурс]. – Режим доступа: <https://www.wi-fi.org/discover-wi-fi/wi-fi-6> (дата обращения: 06.01.2019).
2. Беспроводная передача данных: типы, технология и устройства [Электронный ресурс]. – Режим доступа: <http://fb.ru/article/382356/besprovodnaya-peredacha-dannyih-tipyi-tehnologiya-i-ustroystva> (дата обращения: 06.01.2019).
3. Bluetooth 5.0: что нового и в чем отличие от предыдущих версий? [Электронный ресурс]. – Режим доступа: <https://guidepc.ru/articles/bluetooth-5-0-chto-novogo-i-v-chem-otlichie-ot-predydushhih-versij/> (дата обращения: 06.01.2019).
4. Сети ZigBee. Зачем и почему? [Электронный ресурс]. – Режим доступа: <https://habr.com/post/155037/> (дата обращения: 06.01.2019).
5. Сети GSM. Взгляд изнутри [Электронный ресурс]. – Режим доступа: <https://www.ixbt.com/mobile/gsm-nets.html> (дата обращения: 06.01.2019).
6. Технология NB-IoT: особенности развертывания в сотовых сетях! [Электронный ресурс]. – Режим доступа: <https://networkguru.ru/tekhnologiya-nb-iot-osobennosti-razvertyvaniia/> (дата обращения: 06.01.2019).
7. Технология LoRaWAN [Электронный ресурс]. – Режим доступа: <https://itechinfo.ru/node/25> (дата обращения: 06.01.2019).

**Соломонова Юлия Викторовна**

Год рождения: 1996

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № К4242Направление подготовки: 45.04.04 – Интеллектуальные системы

в гуманитарной среде

e-mail: solomonovajulia@gmail.com

**Хлопотов Максим Валерьевич**

Год рождения: 1980

Университет ИТМО, факультет инфокоммуникационных технологий,
к.т.н., доцент

e-mail: khlopotov@corp.ifmo.ru

УДК 004.912

**РАЗРАБОТКА РЕКОМЕНДАТЕЛЬНОЙ СИСТЕМЫ С УЧЕТОМ ТЕМАТИЧЕСКИХ
ПРОФИЛЕЙ ПОЛЬЗОВАТЕЛЕЙ СОЦИАЛЬНОЙ СЕТИ****Соломонова Ю.В., Хлопотов М.В.****Научный руководитель – к.т.н., доцент Хлопотов М.В.**

В работе представлен подход к векторизации текста на основе Государственного рубрикатора научно-технической информации. Описаны выбор основания для построения тематического профиля, последовательность отбора ключевых слов, а также алгоритм построения тематического профиля и сравнения тематических векторов. Представлены подходы к улучшению разработанного векторного пространства. Описано внедрение подхода к векторизации текста при разработке рекомендательной системы с учетом тематических профилей пользователей социальной сети.

Ключевые слова: классификация текстов, векторизация текста, тематическое моделирование, рекомендательная система, тематический профиль пользователя социальной сети.

Введение. Тематический профиль представляет собой вектор оценок, соответствующих некоторым темам [1]. Разработка рекомендательной системы с учетом тематических профилей объектов предполагает, что тематический профиль рассчитывается как для пользователя (на основе информации из связанного аккаунта в социальной сети), так и для поста на веб-сайте. Рассчитав эти значения, можно привести пользователя и пост в единое пространство, характеризуемое значениями тематических векторов. В таком случае, рекомендации для пользователя формируются на основании сравнения его тематического вектора и тематического вектора поста.

Промежуточные результаты исследования. В ходе работы был проведен обзор семи универсальных систем классификации, выбранных на основании ГОСТ 7.59-2003 [2], с целью выделить те системы, которые могут быть положены в основу тематического профиля. Была рассмотрена широта охвата классификаторами различных областей знаний, а также их избыточность. По итогам проведения обзора в качестве основы для построения тематического профиля был выбран Государственный рубрикатор научно-технической информации (ГРНТИ). Несмотря на то, что многие категории ГРНТИ посвящены разделам промышленности и технологий, данный рубрикатор, единственный из всех рассмотренных

универсальных систем классификации, наиболее полно и при этом избыточно описывает все области знаний.

Для описания категорий тематического профиля были выбраны ключевые слова. Для проведения автоматизированного поиска ключевых слов для каждой категории необходимо собрать достаточное количество текстов, однозначно принадлежащих к каждой из категорий классификатора. В качестве источника для сбора данных был выбран электронный каталог книг Государственной публичной научно-технической библиотеки Сибирского отделения РАН [3]. В связи с тем, что книги в библиотеке были распределены по категориям классификатора ГРНТИ экспертами, можно с уверенностью использовать собранные данные для отбора ключевых слов.

Для каждой из категорий ГРНТИ с помощью парсера на языке Python была собрана информация о соответствующих книгах (их названиях, аннотациях и предметных рубриках). Собранные данные были нормализованы с помощью морфологического анализатора rymorphy2.

Для каждой из категорий был произведен отбор 200 наиболее популярных ключевых слов (униграмм и биграмм). Полученные списки ключевых слов были вручную проверены на адекватность соответствия тематике категории; неподходящие слова были исключены. Пример результата отбора ключевых слов для категории «Литература. Литературоведение. Устное народное творчество» приведен в табл. 1.

Таблица 1. Ключевые слова категории «Литература. Литературоведение. Устное народное творчество»

Униграммы	Биграммы
проза, детективный, слово, рассказ, миф, приключенческий, язык, литература, художественный, отражение, фантастический, философия, собрание, история, творчество, текст, контекст, стиль, литературоведение, писатель, приключение, пушкин, чтение, фольклор, письмо, сказка, поэтика, достоевский, сочинение, поэт, книга, произведение, филология, литературный, поэзия, роман	литература зарубежный, зарубежный литература, драматический произведение, михаил булгаков, картина мир, серебряный век, литература xviii, сергей есенин, плата чехов, полка игорев, литературный процесс, полный собрание, левый толстой, ю лермонтов, малое собрание, русский зарубежье, русский словесность

Особенности некоторых ключевых слов («полка игорев», «левый толстой») обусловлены выбором морфологического анализатора rymorphy2. Однако, несмотря на их семантическую некорректность, подсчет ключевых слов в тексте все равно будет выполнен верно, поскольку исходный текст также будет приведен к нормальной форме с помощью того же морфологического анализатора.

По итогам работы был собран словарь, хранящий ключевые слова для каждой из 69 категорий ГРНТИ.

Формирование вектора текста состоит из следующих шагов:

1. для каждой категории в словаре подсчитывается количество ключевых слов, использованных в данном тексте;
2. значение соответствующего элемента вектора устанавливается равным отношению количества найденных ключевых слов данной категории к общему количеству слов в тексте;
3. значение вектора нормализуется таким образом, чтобы сумма элементов вектора равнялась единице.

Пример результата работы алгоритма построения тематического профиля текстов приведен в табл. 2.

Таблица 2. Результат работы алгоритма построения тематического профиля

Текст	Ключевые слова	Категории	Полученное значение вектора
Блюда из рыб фугу очень популярны в Юго-Восточной Азии. Не потому, что там живет много любителей пощекотать себе нервы, а потому что у мяса фугу приятный вкус. Китайские ученые попытались его воссоздать и определили 12 составляющих его веществ. Среди них оказались несколько аминокислот, в том числе глутаминовая, нуклеотид инозинмоно-фосфат и ионы калия, натрия и фосфорной кислоты	ученый (1)	00: Общественные науки в целом	[0.11111, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0,
	ученый (1)	12: Науковедение	0.11111, 0.0, 0.0, 0.0,
	число (1)	27: Математика	0.0, 0.0, 0.0, 0.0, 0.0,
	вещество (1)	31: Химия	0.0, 0.0, 0.0, 0.11111,
	рыба (1)	34: Биология	0.0, 0.0, 0.0, 0.11111,
	азия (1)	39: География	0.11111, 0.0, 0.0, 0.0,
	мясо (1)	65: Пищевая промышленность	0.11111, 0.0, 0.0, 0.0,
	рыба (1)	69: Рыбное хозяйство. Аквакультура	0.0, 0.0, 0.0, 0.0, 0.0,
ученый (1)	81: Общие и комплексные проблемы технических и прикладных наук и отраслей народного хозяйства	0.0, 0.0, 0.0, 0.11111, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.11111, 0.0, 0.0, 0.0, 0.0, 0.0]	

Полученные значения векторов могут использоваться для сравнения текстов. Так, можно рассчитать косинусное расстояние между векторами. В таком случае 1 будет означать, что рассматриваемые тексты максимально близки тематически, а 0 – что рассматриваемые тексты максимально далеки. Примеры расчета косинусного расстояния между тематическими векторами текстов приведены в табл. 3. Как можно увидеть, значение косинусного расстояния отражает тематическое сходство текстов.

Таблица 3. Примеры расчета косинусного расстояния между векторами текстов

Текст 1	Текст 2	Косинусное расстояние
В армии США есть необычный вид потерь – «переутомление в бою». К этой категории относят в первую очередь тех, кто оказался в плену. Так, при высадке в Нормандии в июне 1944-го года количество «переутомленных в бою» составило около 20% от общего числа выживших из боя. В целом, по итогам Второй мировой войны по причине «переутомления»	В последние месяцы войны немецкое командование драконовыми методами пыталось заставить войска сражаться, но тщетно. Особенно неблагоприятной была ситуация на Западном фронте. Там немецкие солдаты, зная о соблюдении Англией и США Женевской конвенции об обращении с военнопленными, сдавались гораздо охотнее, чем на Востоке	0,90453
	Музей Анны Ахматовой в Фонтанном доме напоминает: встреча уже послезавтра, 16 января! Во время экскурсии мы посмотрим на мемориальную экспозицию под углом страшных лет войны и блокады, расскажем о трагических страницах жизни обитателей квартиры №44. Начало в 18.00	0,62280
	Во вторник исполнилось 111 лет со дня рождения Льва Давидовича Ландау, нобелевского лауреата и гениального физика, который работал и читал лекции в МФТИ. Мы решили разобраться, как	0,07750

Текст 1	Текст 2	Косинусное расстояние
потери США составили 929307 человек.	ученый стал легендой Физтеха, и проанализировали типичные фольклорные сюжеты о нем	

В ходе работы также был разработан алгоритм, который для заданного пользователя социальной сети Вконтакте выгружает по 100 постов из каждой его группы и на основе всех собранных текстов строит единый тематический профиль. Примеры результатов работы данного алгоритма приведены в табл. 4. Как видно из таблицы, можно провести соответствие между сферой деятельности человека и пятеркой его наиболее популярных тематических категорий.

Таблица 4. Построение тематических профилей пользователей социальной сети

Пользователь Вконтакте	5 тематических категорий с максимальными значениями	Значения
id53083705 (Медведев Д.А., председатель Правительства РФ)	00: Общественные науки в целом	0,10528
	82: Организация и управление	0,05676
	26: Комплексные проблемы общественных наук	0,05580
	81: Общие и комплексные проблемы технических и прикладных наук и отраслей народного хозяйства	0,04859
	10: Государство и право. Юридические науки	0,04752
id4457105 (Шипулин А.В., российский биатлонист)	00: Общественные науки в целом	0,17190
	26: Комплексные проблемы общественных наук	0,13843
	23: Комплексное изучение отдельных стран и регионов	0,10909
	77: Физическая культура и спорт	0,10289
	18: Искусство. Искусствоведение	0,03306
id8530294 (Васильев А.А., историк моды, искусствовед)	00: Общественные науки в целом	0,09988
	03: История. Исторические науки	0,09406
	26: Комплексные проблемы общественных наук	0,07475
	64: Легкая промышленность	0,07169
	18: Искусство. Искусствоведение	0,06710

Разработанное пространство необходимо дорабатывать. Так, в словаре присутствуют схожие или слишком обобщенные категории, к примеру, «00: Общественные науки в целом» и «26: Комплексные проблемы общественных наук». Кроме того, некоторые ключевые слова встречаются во многих категориях, как, к примеру, ключевое слово «ученый» встречается в категориях «00: Общественные науки в целом», «12: Науковедение», «81: Общие и комплексные проблемы технических и прикладных наук и отраслей народного хозяйства».

В рамках работы были рассмотрены различные способы доработки разработанного векторного пространства.

Одним из способов сокращения разработанного пространства может являться слияние похожих категорий. На основе попарного расчета коэффициента Жаккара для всех категорий был сформирован список наиболее схожих пар категорий тематического профиля. Установив пороговое значение коэффициента Жаккара в 0,1, производилось последовательное слияние самых схожих пар категорий.

Другим способом может являться разделение категорий. Однако, в связи с тем, что более низкие уровни некоторых категорий ГРНТИ лишь усложняют существующую классификацию, данный подход подробнее не рассматривался.

Кроме непосредственной работы над категориями, также возможно внести изменения в сам алгоритм расчета тематического профиля. Так, возможно учитывать количество упоминаний ключевого слова во всем словаре при расчете его веса.

Каждый из рассмотренных подходов, включая исходный, был оценен на тестовом наборе данных, полученном из дополнительного каталога электронной библиотеки [3]. Для каждого из подходов оценивание проводилось тремя различными способами: результат алгоритма считался верным, если искомая категория находилась в Топ-1, Топ-3 или Топ-5 списка, сформированного алгоритмом построения тематического профиля. Результаты оценивания приведены в табл. 5.

Таблица 5. Оценивание подходов к улучшению векторного пространства

	Исходный словарь ключевых слов	Словарь ключевых слов с объединенными категориями
Исходный алгоритм расчета веса ключевого слова	Топ-1: 0.69118 Топ-3: 0.86765 Топ-5: 0.89706	Топ-1: 0.60294 Топ-3: 0.68015 Топ-5: 0.79412
Алгоритм с улучшенным расчетом веса ключевого слова	Топ-1: 0.72059 Топ-3: 0.88236 Топ-5: 0.91176	Топ-1: 0.64706 Топ-3: 0.77941 Топ-5: 0.80882

Как можно увидеть, наилучшие результаты показывает алгоритм с улучшенным расчетом веса ключевого слова. При этом объединение схожих категорий классификатора лишь ухудшает качество работы алгоритма.

Выводы и дальнейшая работа. Разработанный классификатор с улучшенным алгоритмом расчета веса ключевого слова был внедрен на веб-ресурс OpinionLine [4]. Тематический вектор поста автоматически пересчитывается при каждом обновлении его содержания – связанных с постом цитат и опросов. Тематический вектор пользователя рассчитывается на основе текстовой информации, собранной из связанного аккаунта социальной сети Вконтакте (в частности, его групп и подписок). На основе сравнения тематических векторов для пользователя на главной странице формируются индивидуально подобранные рекомендации постов веб-ресурса.

В дальнейшем планируется проведение эксперимента, в рамках которого различным пользователям будут назначаться различные стратегии формирования рекомендаций (в частности, будут рассмотрены контентная стратегия, при которой рекомендации основываются на 10 разработанных самостоятельно тематических категориях, а также стратегия с учетом тематических профилей, описанная в данной работе). На основе данных, собранных в ходе эксперимента, будут рассчитаны метрики, на основе которых можно будет делать вывод о том, как каждая из стратегий формирования рекомендаций влияет на поведение пользователя на веб-ресурсе.

Литература

1. Лексин В.А. Технология персонализации на основе выявления тематических профилей пользователей и ресурсов Интернет [Электронный ресурс]. – Режим доступа: <http://www.machinelearning.ru/wiki/images/c/c2/Lexin07master.pdf> (дата обращения: 06.01.2019).
2. ГОСТ 7.59-2003. СИБИБД. Индексирование документов. Общие требования к систематизации и предметизации. – Введен 01.01.2004. – М.: Стандартинформ, 2005. – 8 с.
3. Государственная публичная научно-техническая библиотека СО РАН [Электронный ресурс]. – Режим доступа: http://webirbis.spsl.nsc.ru/irbis64r_01/cgi/cgiirbis_64.exe?C21COM=F&I21DBN=HELP&S21FMT=web_rub_wn&S21ALL=%3C.%3E0%3C.%3E&P21DBN=CAT&S21CNR=20&Z21ID= (дата обращения: 06.01.2019).
4. OpinionLine [Электронный ресурс]. – Режим доступа: <https://www.opiniononline.ru/> (дата обращения: 06.01.2019).



Терентьев Александр Викторович

Год рождения: 1996

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K41211

Направление подготовки: 11.04.02 – Программное обеспечение
в инфокоммуникациях

e-mail: terentev.a@inbox.ru



Осипов Никита Алексеевич

Год рождения: 1972

Университет ИТМО, факультет инфокоммуникационных технологий,
к.т.н., доцент

e-mail: nikita@ifmo.spb.ru

УДК 004.93'14

АНАЛИЗ СУЩЕСТВУЮЩИХ ТЕХНОЛОГИЙ БЕСПИЛОТНЫХ АВТОМОБИЛЕЙ

Терентьев А.В.

Научный руководитель – к.т.н., доцент Осипов Н.А.

Автомобили с автоматизированными системами управления можно встретить на дорогах уже сегодня. И с каждым днем их будет становиться все больше.

Ключевые слова: автоматизированные системы, уровни беспилотности.

Введение. Благодаря нарастающей конкуренции производителей автомобилей и производителей технологий, таких как Apple и Google, ведется интенсивное развитие систем беспилотного управления. Такие системы наделены определенным уровнем беспилотности [1–3].

Уровни беспилотности. Сообщество автомобильных инженеров (SAE International) из Америки разделило системы беспилотности на уровни. Данная классификация называется шкалой SAE. Классификация включает шесть уровней – от нулевой до полной автоматизации:

1. нулевой уровень – отсутствие автоматизации. Системы нулевого уровня помогают водителю, предупреждая об опасных ситуациях. Они не выполняют никаких функций по управлению автомобилем. Примером систем нулевого уровня автоматизации являются:
 - система ночного видения предназначена для предоставления водителю информации об условиях движения в темное время суток. Данная система помогает водителю обнаружить различные препятствия на дороге. Она устанавливается как дополнительная опция в автомобиле премиум-класса;
 - система помощи при перестроении или система мониторинга «слепых» зон. При смене полосы движения система предупреждает водителя о возможности столкновения;
 - парковочная система или парктроник – второстепенная система активной безопасности автомобиля. Используя эту систему, водитель получит возможность припарковать автомобиль в условиях плохой видимости;
2. первый уровень – помощь водителю. Системы первого уровня выполняют некоторую часть функций управления автомобилем. Водитель должен постоянно следить за процессом и по необходимости взять полное управление на себя. Примером систем первого уровня автоматизации являются:

- адаптивный круиз-контроль. Данная система – это интеллектуальная форма круиз-контроля, которая автоматически замедляется и ускоряется, чтобы идти в ногу с автомобилем перед вами. Водитель задает максимальную скорость – просто как с круиз-контролем – тогда радарный датчик следит за движением;
 - активная система помощи движению по полосе. Эта система противоположна системе адаптивного круиз-контроля. Рулевое управление находится в руках системы, а за скоростью следит водитель. Система помогает удерживать автомобиль в выбранной полосе, предотвращая опасные ситуации;
 - система автоматической парковки. Система использует различные датчики для определения приблизительного размера пространства между двумя припаркованными транспортными средствами, а затем встроенный компьютер вычисляет необходимые углы поворота и скорости, чтобы безопасно перемещаться в парковочное место;
3. второй уровень – частичная автоматизация. Управление автомобилем полностью принадлежит системе второго. Руки можно убрать с рулевого колеса, ногу – с педали акселератора. Процесс работы системы требует непосредственного контроля со стороны водителя. Примером систем второго уровня автоматизации являются:
- автопилот Tesla. Автопилот, позволяющий автомобилю автономно двигаться на автомагистрали, используется на модели Tesla Model S с 2015 года. Тесла соответствует скорости движения условия, держит в пределах полосы движения, автоматически изменяет полосы движения без водителя, переходит от одного шоссе к другому, выходит на шоссе, когда пункт назначения находится рядом, самостоятельно паркуется, когда рядом есть место для парковки, и может быть вызван из гаража;
 - система автоматического движения в пробках (Traffic Jam Assistant) от компании Audi. Данная система помогает водителям прибыть в пункт назначения более расслабленными, даже в плотном транспортном потоке или в пробках. В качестве частично автоматизированной функции комфорта система берет на себя продольное и боковое наведение транспортного средства. Это означает, что автомобиль может автоматически отъезжать, ускоряться и тормозить, а также управлять транспортным средством в определенных ограничениях. Водитель должен постоянно контролировать систему и быть готовым взять на себя полный контроль над автомобилем в любой момент;
 - временный автопилот (Temporary Auto Pilot). Система позволяет автомобилю держать безопасное расстояние от автомобиля перед ним, соответствует скорости водителя выбора, замедляется по мере необходимости перед поворотами на дороге, и держит автомобиль в своей собственной полосе;
4. третий уровень – обусловленная автоматизация. В системах третьего уровня не требуется внимание водителя за контролем движения автомобиля. В это время он может читать, писать, смотреть видео. Система производит оповещение, когда требуется вмешательство водителя в процесс движения. К системам третьего уровня автоматизации относятся следующие разработки:
- система Super Cruise: система беспилотного управления Super Cruise от Cadillac обеспечивает движение автомобиля по автомагистрали. Маневрирование, торможение и движение по полосе осуществляется системой без участия водителя;
 - система SARTRE: система Safe Road Trains for the Environment (SARTRE) от компании Volvo, организывает движущуюся по дороге колонну из нескольких машин. В качестве головной машины выступает грузовой автомобиль с профессиональным водителем;
5. четвертый уровень – высокая автоматизация. Внимание водителя в системах четвертого уровня не требуется. Активировав систему, водитель может заниматься своими делами или вообще покинуть водительское сидение. В случае возникновения опасной ситуации система производит парковку в безопасном месте и доводит информацию до водителя. Систем, которые можно отнести к четвертому уровню автоматизации, пока немного:

- беспилотный автомобиль Google: в настоящее время система беспилотного управления установлена на автомобилях Toyota Prius, Lexus RX 450h и Audi TT, которые проехали в беспилотном режиме свыше двух с половиной миллионов километров. Работа системы осуществляется с помощью следующих входных устройств: лидар, радары, видеокамера, датчик оценки положения, инерционный датчик движения, GPS приемник;
 - система автономной парковки: система автономной парковки является дальнейшим развитием системы автоматической парковки. Разработкой этой системы занимаются компании: Volkswagen с Bosch, BMW с Continental;
6. пятый уровень – полная автоматизация. В системах пятого уровня автоматизации человек вообще не нужен. Поэтому органы управления (рулевое колесо, педали) могут быть убраны из салона автомобиля. Пассажир активирует (указывает пункт назначения) и деактивирует систему. Сегодня систем автоматизации пятого уровня нет.

Заключение. Систем с полной автоматизацией управления автомобиля или, коротко, беспилотных автомобилей еще не разработали. Существующим системам пока не позволят полностью доверить жизни пассажиров без их вмешательства в процесс управления. Каждая из систем не совершенна и требует доработок. Работы над исправлением ошибок и недочетов в этих системах ведутся до сих пор. Создание рабочих систем 5 уровня в ближайшее время не планируется, но выход полноценных систем 4 уровня в массы ожидается.

Литература

1. Беспилотный автомобиль [Электронный ресурс]. – Режим доступа: <http://systemsauto.ru/another/driverless-car.html> (дата обращения: 06.01.2019).
2. Уровни автоматизации автомобилей. Уровни беспилотных авто 0-5 [Электронный ресурс]. – Режим доступа: <https://bespilot.com/info/urovni-avtomatizatsii-avtopilota> (дата обращения: 06.01.2019).
3. Вчера и завтра: 6 уровней автоматизации автомобилей [Электронный ресурс]. – Режим доступа: <https://novate.ru/blogs/110617/41716/> (дата обращения: 06.01.2019).

**Тимоненков Денис Юрьевич**

Год рождения: 1997

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № К41201Направление подготовки: 11.04.02 – Программное обеспечение
в инфокоммуникациях

e-mail: mrdentimon@gmail.com

**Осипов Никита Алексеевич**

Год рождения: 1972

Университет ИТМО, факультет инфокоммуникационных технологий,
к.т.н., доцент

e-mail: nikita@ifmo.spb.ru

УДК 004.414.22**ОПЫТ КАСТОМИЗАЦИИ CRM-СИСТЕМЫ НА ПРИМЕРЕ РАЗРАБОТКИ
ВИДЖЕТА ИМПОРТА/ЭКСПОРТА ЦИФРОВОЙ ВОРОНКИ****Тимоненков Д.Ю.****Научный руководитель – к.т.н., доцент Осипов Н.А.**

В работе приведен способ решения проблемы переноса определенных настроек цифровой воронки CRM-системы. Проблема рассмотрена на примере компании, занимающейся интернет-маркетингом с использованием CRM-системы «amoCRM». Описаны настройки, которые возможно экспортировать/импортировать; составные части программной доработки, необходимой для решения проблемы; средства реализации.

Ключевые слова: CRM-система, amoCRM, виджет, цифровая воронка, экспорт настроек аккаунта, импорт настроек аккаунта.

CRM-системы на данный момент активно внедряются во все сферы бизнеса. Так как универсальных CRM не существует, то в случае с системами с открытым API кастомизация реализуется путем программных доработок. В качестве примера такой кастомизации в данной работе рассмотрена доработка CRM, связанная с маркетингом.

Заказчик занимается настройкой маркетинговых инструментов для своих клиентов, таких как автоматические электронные письма и SMS, движения сделок по воронке продаж, отправка данных в систему веб-аналитики, подписка контакта сделки на рекламные рассылки и цепочки писем. Для этого используется CRM-система «amoCRM». Нередко заказчику приходится работать с клиентами в одной сфере или направлении, где бизнес-процессы и маркетинговые инструменты либо очень близки, либо одинаковы. Чтобы сократить время на настройку для своих клиентов, была поставлена задача упрощения и автоматизации этого процесса. Необходимо переносить настройки одного аккаунта «amoCRM» в другой аккаунт.

«amoCRM» открыта для разработки и размещения в магазине интеграций новых виджетов [1], которые разрабатывают партнеры, поэтому проблем с доступом к API не возникло. Было принято решение разбить виджет на две части: одна для экспорта, вторая для импорта. Не все желаемые заказчиком настройки можно экспортировать из-за технических особенностей подключения виджетов внутри системы. Список настроек, доступных для экспорта, следующий:

1. воронки продаж;

2. дополнительные поля всей сущностей: сделок, контактов, компаний, покупателей;
3. инструменты «amoCRM» (рисунок):
 - создание задачи;
 - создание сделки;
 - создание покупателя;
 - отправка письма;
 - настройка Salesbot'a;
 - отправка webhook'a;
 - смена статуса сделки;
 - редактирование тегов;
 - завершение задач;
4. виджеты:
 - «MailChimp»;
 - «SMS-Центр»;
 - «SMS.RU».

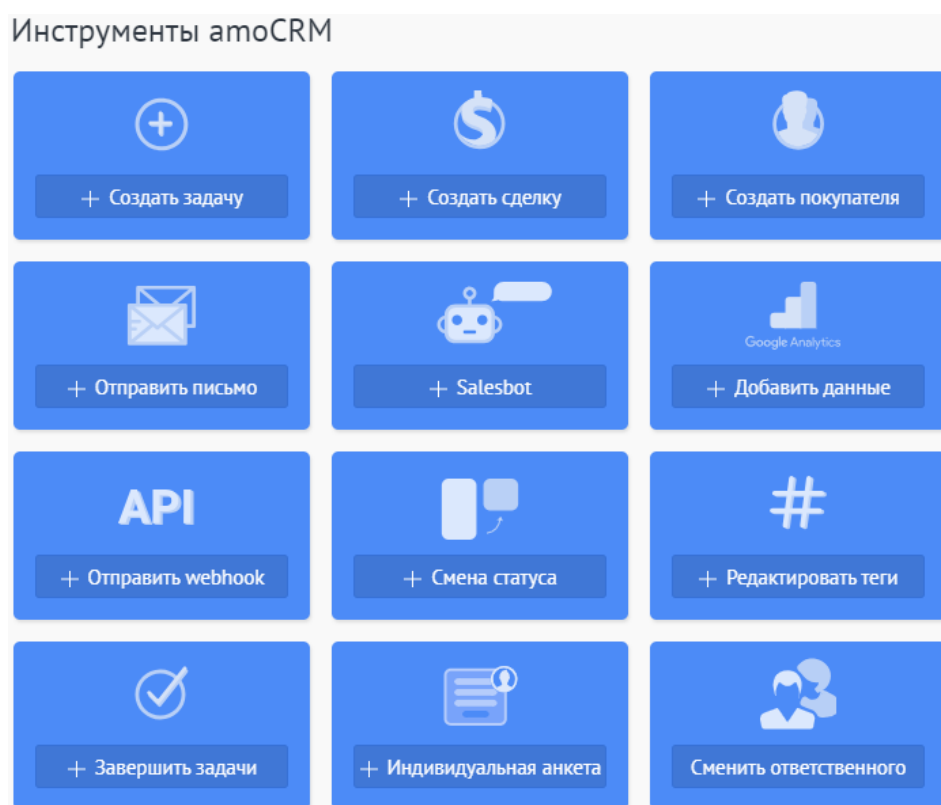


Рисунок. Инструменты «amoCRM»

Экспортная часть представляет собой виджет во вкладке интеграции в «amoCRM». В нем пользователь отмечает, какие настройки будут экспортированы из аккаунта. Можно экспортировать как одну, так и несколько воронок. Дополнительные поля экспортируются все. Получение настроек воронок, дополнительных полей и шаблонов электронных писем осуществляется через GET-запросы. Для экспорта виджетов и инструментов необходимо применить парсинг HTML-кода. Все полученные параметры используются для формирования JSON-файла, который сохраняется на устройстве пользователя.

Импортная часть также представляет собой виджет, предлагающий загрузить JSON-файл для импорта настроек в аккаунт. В импортной части необходимо предусмотреть предварительную подготовку аккаунта и ошибки, которые могут возникнуть при импорте.

В предварительную подготовку входят:

1. создание новых типов задач из файла;
2. добавление тегов в аккаунт;
3. добавление шаблонов электронных писем;
4. добавление шаблонов для сервисов «SMS.RU» и «SMS-Центр».

Ошибки могут возникнуть в следующих случаях:

1. в файле отсутствуют настройки для дополнительных полей, хотя эти поля используются, например, в условии для постановки задачи. В данном случае пользователю предлагается выполнить экспорт еще раз, на этот раз указав для экспорта нужные настройки;
2. в файле отсутствуют настройки воронок, необходимых для работы цифровой воронки. В данном случае пользователю предлагается выполнить экспорт еще раз, на этот раз указав для экспорта нужные настройки;
3. импортируемые виджеты неактивны. Проверка происходит GET-запросом. Если найдены неактивные импортируемые виджеты, то пользователю предлагается их активировать. Активация происходит пользователем самостоятельно, так как для активации необходимо вводить логин и пароль от сервиса.

Структурно виджет для «amoCRM» может состоять только из JS-, JSON- и PHP-файлов и изображений в форматах png, jpeg, jpg и gif [2], поэтому в качестве инструментов разработки будут использоваться JavaScript и PHP. На них будут написаны парсер и запросы. Интерфейс виджета будет реализован с помощью HTML и CSS.

Литература

1. Расширьте возможности с API [Электронный ресурс]. – Режим доступа: <https://www.amocrm.ru/developers/content/platform/abilities/> (дата обращения: 06.01.2019).
2. Структура [Электронный ресурс]. – Режим доступа: <https://www.amocrm.ru/developers/content/integrations/structure> (дата обращения: 06.01.2019).



Хмеленко Анастасия Юрьевна

Год рождения: 1995

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K4232

Направление подготовки: 27.04.03 – Системный анализ

и интеллектуальное управление в бизнесе

e-mail: n.khmelenko@yandex.ru



Максимова Татьяна Геннадьевна

Год рождения: 1967

Университет ИТМО, факультет технологического менеджмента
и инноваций, д.э.н., профессор

e-mail: maximovatg@gmail.com

УДК 65.01

СИСТЕМНЫЙ АНАЛИЗ И ОПТИМИЗАЦИЯ БИЗНЕС-ПРОЦЕССОВ ЧАСТНОЙ МЕДИЦИНСКОЙ КЛИНИКИ

Хмеленко А.Ю.

Научный руководитель – д.э.н., профессор Максимова Т.Г.

Моделирование бизнес-процессов и дальнейшая их оптимизация способствуют планомерному развитию частной организации. Дана общая характеристика деятельности медицинского учреждения, сформулированы существенные детали процесса, которые значительно влияют на реализацию бизнес-процесса. На основе разработанного бизнес-процесса версии AS IS был сформирован набор проблемных зон в деятельности организации, составлена версия процесса TO BE. Результатом оптимизации служит составленный список решенных проблем деятельности организации.

Ключевые слова: моделирование бизнес-процессов, методология BPMN, оптимизация бизнес-процессов, бизнес-процесс медицинской организации, здравоохранение.

В деятельности организации в сфере здравоохранения необходимо использовать процессный подход, поскольку ей присущи специфические черты, характеризующиеся мультидисциплинарностью. Административный персонал медицинской организации должен обладать знаниями и компетенциями в различных сферах, чтобы быть сведущими как в медицине, так и в информационных технологиях, юридических тонкостях и других смежных областях.

В настоящее время фокус на ключевых бизнес-процессах организации предполагает устойчивое положение среди конкурентов на рынке медицинских услуг. Использование таких моделей обеспечивает гибкость стратегии компании, что в дальнейшем, безусловно, положительно скажется на приспособлении к различным переменам во внешней среде.

В настоящее время такая проблема, как необходимость моделирования бизнес-процессов, затрагивается достаточно часто в научной литературе. Однако следует заметить, что на данный момент в полной мере широкий спектр методологий не заточен под отличительные характеристики системы здравоохранения, что является одним из значительных препятствий при моделировании бизнес-процессов медицинской организации.

Деятельность медицинской клиники представляет собой непрерывный процесс, начинающийся с обращения нового пациента в клинику до удовлетворения его потребности в получении медицинской помощи. В связи с тем, что данный процесс является одним из основных бизнес-процессов компаний этой сферы деятельности, оптимизация данного

процесса является целью большинства медицинских учреждений. Целью исследования является оптимизация бизнес-процесса частной медицинской клиники.

Основными методами при корректировке бизнес-процессов являются оптимизация бизнес-процессов и реинжиниринг бизнес-процессов. В качестве методологии в данном исследовании была выбрана оптимизация бизнес-процессов, потому что в качестве начальной точки лежит существующий бизнес-процесс, т.е. работа начинается не с чистого листа, и изменения носят нарастающий характер.

Поскольку бизнес-процесс организации должен быть понятным руководству бизнеса, исполнителям процессов, то графическое изображение модели бизнес-процесса не должно характеризоваться большим объемом деталей, она должна быть визуально простой и интуитивно понятной. Методология BPMN разработана так, чтобы была понятна как профессионалам бизнеса, так и IT-специалистам. Явный дизайн для нетехнических пользователей делает его перспективным кандидатом для моделирования процессов здравоохранения, где медицинский персонал должен понимать и обсуждать модели процессов [1].

С позиции предлагаемого подхода оценку качества оказания услуг медицинским учреждением целесообразно рассматривать как процесс исследования уровня удовлетворенности клиентов различными бизнес-процессами, с которыми они сталкиваются при получении услуги: лечебный процесс, поддерживающий процесс (инфраструктура), процесс сервиса (контактный персонал), процесс маркетинга (имидж компании) [2, 3].

Оптимизация бизнес-процессов проводится в медицинской частной клинике, которая предоставляет услуги медицинского профиля, в частности: прием врачей-клиницистов широкого спектра специализаций, прием врачей функциональной диагностики, услуги лабораторной службы. В клинике также работает административный персонал, отдел маркетинга, руководство, средний медицинский персонал и врачи.

Для моделирования бизнес-процессов удобно воспользоваться специальным программным обеспечением, позволяющим корректно описывать и оформлять бизнес-процессы в соответствии с той или иной нотацией. Процессы деятельности оформлены в нотации BPMN.

Оптимизация бизнес-процессов осуществляется в соответствии со следующими шагами:

1. моделирование текущих бизнес-процессов;
2. определение нерациональных действий;
3. обнаружение слабых мест;
4. моделирование будущих бизнес-процессов;
5. внедрение оптимизированных бизнес-процессов.

Отражение деятельности организации AS IS подразумевает описание процессов и процедур согласно текущей ситуации [4]. В клинике, для которой планируется оптимизация бизнес-процессов, внедрена медицинская информационная система (МИС), в которой ведется лишь график работы врачей и расписание приемов, данные пациентов. Врачи не пользуются этой МИС, поскольку МИС не обладает достаточным набором модулей для врачей. Врачи заполняют протоколы приемов на персональный компьютер на своих рабочих местах. Из этого следует, что История болезни пациента ведется лишь в бумажном варианте и не доступна врачам в электронном виде, в силу того, что она не хранится в единой базе данных. Все выдаваемые пациентам документы по приему, такие как индивидуальное информированное добровольное согласие (ИДС) на медицинское вмешательство, назначения и рекомендации, справки, заполняются врачом в рукописном варианте на заранее подготовленных бланках. По окончании приема врач отмечает на талоне (бумажный носитель), какие услуги были оказаны в рамках приема, и относит данный талон в регистратуру. Сотрудники регистратуры вносят данные пациентов в МИС, однако оплату принимают посредством другого программного обеспечения – 1С. При переносе данных из талона в 1С администраторы зачастую допускают ошибки, исправление которых негативно сказывается на времени пребывания пациента в клинике.

Общая схема оказания медицинской помощи в амбулаторных условиях версии AS IS представлена на диаграмме, оформленной в нотации BPMN, на рис. 1.

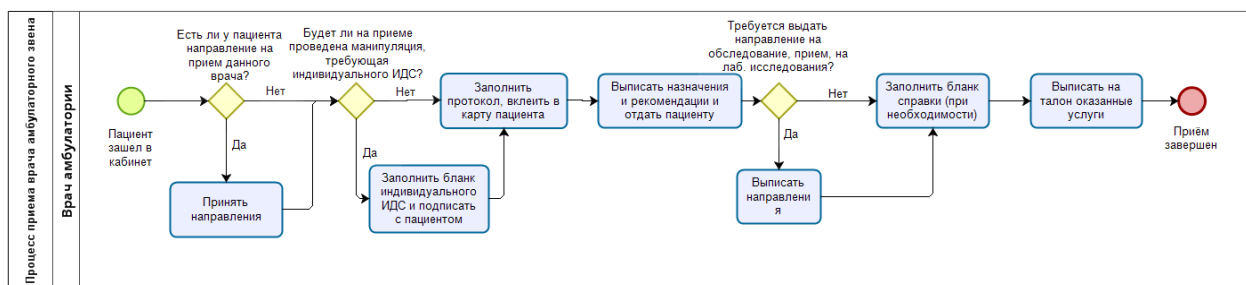


Рис. 1. Бизнес-процесс версии AS IS

Для простоты понимания незначительные детали не отражены в данной диаграмме бизнес-процесса.

Оптимизация является достаточно эффективным методом изменения работы организации. С введением нового порядка деятельности меняются многие процедуры, которые выполняются персоналом организации, меняется последовательность некоторых этапов работ, также меняются ответственные за ту или иную операцию. Переход от привычного распорядка работы к усовершенствованному должна быть аргументирована и подкреплена фактами. Таким образом, первоначальным обязательным этапом оптимизации является выявление слабых мест в деятельности организации.

Анализ данного бизнес-процесса позволяет составить причинно-следственную диаграмму Исикавы, которая отражает взаимосвязь узких мест в деятельности клиники; данная диаграмма изображена на рис. 2. При помощи данного инструмента были выявлены те проблемные зоны, в которых возникает большая часть ситуаций, которые негативно сказываются на качестве обслуживания пациента.



Рис. 2. Взаимосвязь причин, влияющих на качество обслуживания

В связи с наличием значительных недостатков в работе клиники руководством было принято решение оптимизации одного из основных процессов – процесс «Оказание медицинской помощи в амбулаторных условиях».

При переходе к новой МИС, обладающей достаточным функционалом для персонала, как медицинского профиля, так и административного, ряд проблем будет исключен. В результате внедрения новой МИС будут изменены некоторые действия, которые позволят врачу увеличить время общения с пациентом, администраторам – принимать оплату в той же ИС, где и ведут расписание и т.д.

Общая схема оказания медицинской помощи в амбулаторных условиях версии ТО ВЕ представлена на диаграмме, оформленной в нотации BPMN, на рис. 3.

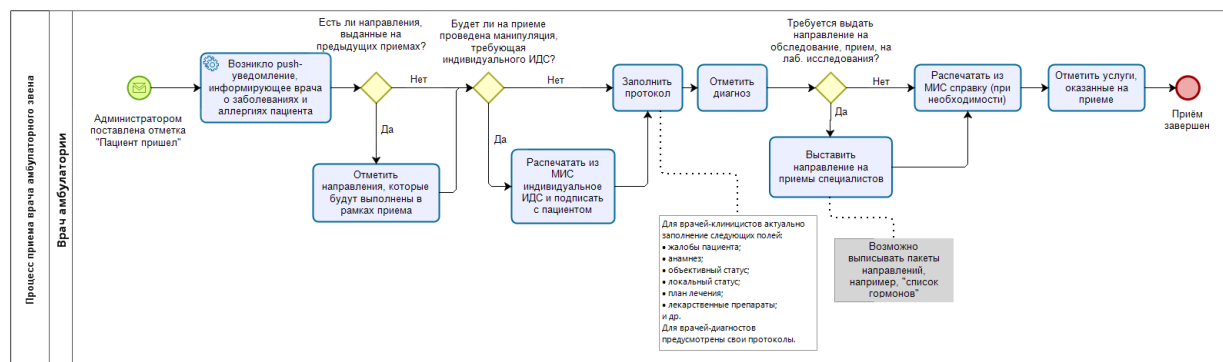


Рис. 3. Бизнес-процесс версии ТО ВЕ

Таким образом, изменение повлечет за собой повышение уровня удовлетворенности пациента. Безусловно, оптимизация одного процесса не является достаточным средством по повышению уровня удовлетворенности пациентом.

Список решенных проблем после оптимизации бизнес-процесса:

1. ведение единой Истории болезни пациента, доступной в любом месте при наличии интернета;
 - другие врачи клиники могут с легкостью ознакомиться с информацией по предыдущим приемам, исследованиям;
 - распечатать медицинские документы пациента можно в любое время, не только на приеме;
 - контроль главным врачом ведения протоколов для мотивации и поощрения врачей;
2. врач не тратит время приемов на встречу пациента у регистратуры, а готовится к следующему приему;
3. администратор не допускает ошибок при принятии оплаты – выставленные врачом услуги автоматически переносятся в талон;
4. запись на повторный прием в удобное ему время может осуществлять сам врач – время администратора не тратится на данный вид работы;
5. операторы call-центра фиксируют обращения пациентов с целью формирования статистики – наиболее часто запрашиваемые виды услуг фиксируются в отчетах, что позволит отделу маркетинга предложить новые услуги, в которых уже есть интерес со стороны пациентов;
6. минимизация ошибок врача при заполнении бланков документов – распечатываемые из МИС документы заполнены информацией о пациенте и враче.

Оптимизация бизнес-процесса положительно скажется на повышении эффективности работы клиники.

Таким образом, при внедрении новой МИС круг проблем, с которыми сталкивался персонал медицинской организации, существенно сократился. Все действия врачей и административного персонала при оформлении документов и ведении единой базы данных совершаются с минимальным количеством ошибок.

Выводы, направления дальнейших исследований. На основе полученных данных после проведения оптимизации бизнес-процесса руководство медицинской клиники способно четко сформулировать основные стратегии развития организации согласно обновленному бизнес-процессу, конкретизировать преимущества оказания услуг медицинского профиля.

К перспективным направлениям научного исследования можно отнести детальную проработку в части системного анализа остальных бизнес-процессов медицинской частной клиники, составление карты бизнес-процессов клиники, составление верхнеуровневого бизнес-процесса организации, проведение разработки новых бизнес-процессов, относящихся к основной группе.

Литература

1. Müller R., Rogge-Solti A. BPMN for Healthcare Processes [Электронный ресурс]. – Режим доступа: <http://ceur-ws.org/Vol-705/paper9.pdf> (дата обращения: 23.02.2019).
2. Уварина Ю.А., Шушкин М.А. Инновационные бизнес-модели медицинских центров: маркетинговый инструментальный анализ реализации бизнес-процессов // Право, менеджмент, маркетинг, инновации. – 2016. – № 1(207). – С. 99–108.
3. Рыжков Н.А., Верзилин Д.Н., Максимова Т.Г., Антохин Ю.Н. Управление экономической эффективностью программы информатизации многофункционального государственного учреждения // Научно-технические ведомости Санкт-Петербургского государственного политехнического университета. Экономические науки. – 2012. – № 3(149). – С. 73–76.
4. [Электронный ресурс]. – Режим доступа: http://www.uef.fi/documents/677096/736588/UBI_health_presentation_0.3.pdf/38913bba-4d57-4bf6-aa53-152f826159a1 (дата обращения: 06.01.2019).

**Шаркова Елизавета Игоревна**

Год рождения: 1996

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № K41211Направление подготовки: 11.04.02 – Программное обеспечение

в инфокоммуникациях

e-mail: sharkliza@gmail.com

**Ананченко Игорь Викторович**

Год рождения: 1968

Университет ИТМО, факультет инфокоммуникационных технологий,
к.т.н., доцент

e-mail: igor@anantchenko.ru

УДК 004.93'12

**ИССЛЕДОВАНИЕ МЕТОДОВ ОТОБРАЖЕНИЯ ВИРТУАЛЬНЫХ ГРАФИЧЕСКИХ
МОДЕЛЕЙ В СИСТЕМАХ ДОПОЛНЕННОЙ РЕАЛЬНОСТИ****Шаркова Е.И.****Научный руководитель – к.т.н., доцент Ананченко И.В.**

В работе рассмотрены особенности разработки систем дополненной реальности, а именно, виды и принцип действия специальных устройств, а также программные и алгоритмические решения для интеграции технологии в смартфоны.

Ключевые слова: дополненная реальность, видеопрозрачные устройства, безмаркерный алгоритм, Vuforia SDK, EasyAR SDK.

Дополненная реальность – это дословный перевод фразы augmented reality (AR). Термином augmented reality называют симбиоз нашей реальности, и виртуальной, компьютерной. Другими словами, дополненная реальность – это визуальное дополнение реального мира путем проецирования и введения каких-либо виртуальных, мнимых объектов на настоящее пространство (на экране компьютера, телефона и подобных устройств) [1].

Чем же должна обладать техника, поддерживающая данную технологию? В первую очередь, камерой для захвата объектов реального мира. Во-вторых, дисплеем для отображения виртуальных объектов.

На сегодняшний день для этого применяют устройства, закрепляющиеся на голове (Head Mounted Displays, HMD) или устройства с ручным дисплеем (Manual Displays, MD). Существует два вида HMD: видео и оптико-прозрачные. Видеопрозрачное устройство подразумевает, что пользователь носит две камеры на голове, изображение с них поступает в процессор устройства, где суммируется с дополнительными объектами. Далее полученное изображение выводится на два экрана перед глазами пользователя. В оптико-прозрачных устройствах применяется технология: «половины серебряного зеркала». Пользователь смотрит на мир с помощью специальной полупрозрачной линзы, через которую видно окружающий мир, и на которую накладывается дополнительная информация. Преимуществом видеопрозрачных устройств является более правдоподобная интеграция изображений, а оптико-прозрачные устройства выигрывают в лучшем восприятии окружающего мира [2]. Несмотря на это, данные технологии требуют больших затрат от производителей – на камеры или специальные линзы, поэтому и стоимость конечного продукта получается высокой – от

1000\$. Этот фактор тормозит выход технологии дополненной реальности в массовое использование. Однако ситуация в корне меняется при использовании MD. Ручной дисплей – это в первую очередь смартфоны и планшеты. Главное достоинство использования их для разработки дополненной реальности – это цена, а также массовая распространенность, меньшие габариты, большие объемы памяти. Недостатками применения их в данной области являются качество камеры и меньший, чем у HMD, эффект погружения.

Разберем подробнее алгоритм работы приложений AR на смартфонах. При запуске такого приложения включается камера (если требуется, то еще и датчики наклона, геолокация или микрофон). Происходит захват картинки окружающего мира. Как только выполняется одно из заложенных условий, начинается поиск соответствующего графического объекта для AR. Часто такие объекты хранятся на сервере, чтоб не занимать много места на устройстве пользователя. Когда нужный объект найден, вычисляются его координаты в пространстве реального мира относительно устройства. Далее картинка объекта накладывается на изображение с камеры, и полученное изображение выводится на экран смартфона. Разнообразие таких приложений зависит от графических объектов, смоделированных художниками, а также от условий их появления, заложенных разработчиками. Чаще всего встречаются графические, звуковые, геолокационные условия. Графические условия бывают маркерные и безмаркерные. Маркерные – это когда картинка с камеры сопоставляется с заранее заданным набором образов, маркерами. Безмаркерный алгоритм находит некие опорные точки и по ним определяет место, к которому должен быть «привязан» виртуальный объект. Преимущество такой технологии в том, что объекты реального мира служат маркерами сами по себе, и для них не нужно заранее создавать специальных идентификаторов [3]. Например, если требуется распознать собаку, то используя маркерную систему, придется сделать фотографии всех пород собак, всех окрасов и во всех возможных позах, а в безмаркерной системе надо смоделировать объект очертаний собаки, состоящий из некоторого числа соединенных точек, чем точек больше, тем точность определения выше.

Для создания вышеописанных приложений используются различные готовые библиотеки. Наиболее популярные бесплатные библиотеки для создания приложений дополненной реальности: Vuforia SDK, ARToolKit, EasyAR SDK. Данные библиотеки поддерживают кроссплатформенное программирование, т.е. можно создавать приложения для Android, iOS, и даже Web. У Vuforia SDK основным языком программирования – это C++, также она доступна на C#, Java, Objective-C [4]. У ARToolKit: Java, Objective-C, а у EasyAR SDK: C, C++, Java, Objective-C [5]. С 2018 года Vuforia SDK встроена в установочный пакет Unity – популярного игрового движка, также можно скачать и установить версию Vuforia SDK для Android Studio – IDE для создания Android приложений. У ARToolKit на сайте разработчика есть версии и для Unity, и для Android Studio, а вот настройка EasyAR SDK займет больше времени, но тоже возможна. Кроме того, для нее и документации на сайте разработчика существенно меньше. Исходя из описанного выше, Vuforia SDK лучше всего подойдет для создания приложения дополненной реальности на игровом движке Unity, а ARToolKit лучше всего подойдет, если приложение планируется программировать на Java в Android Studio.

В заключение стоит отметить, что технология дополненной реальности еще развивается, но уже используется во многих сферах жизни: в авиации – для отображения информации о скорости, наклоне самолета, на производстве – для отображения информации о сложных деталях, в маркетинге – для наглядной презентации товара, в электротехнике – для визуализации сложных схем, в обучении, медицине, геолокации, игровой индустрии и так далее. Это значит, что сейчас исследования данной технологии являются актуальными.

Литература

1. Применение технологии дополненной реальности с использованием QR-кодов в создании туристического маршрута г. Екатеринбурга [Электронный ресурс]. – Режим доступа: https://vuzlit.ru/379033/dopolnennaya_realnost (дата обращения: 06.01.2019).

2. Rolland J.P., Fuchs H. Optical Versus Video See-Through Head-Mounted Displays in Medical Visualization [Электронный ресурс]. – Режим доступа: https://www.researchgate.net/publication/220089776_Optical_Versus_Video_See-Through_Head-Mounted_Displays_in_Medical_Visualization (дата обращения: 06.01.2019).
3. Увлекательная реальность [Электронный ресурс]. – Режим доступа: https://funreality.ru/technology/augmented_reality/ (дата обращения: 06.01.2019).
4. Vuforia [Электронный ресурс]. – Режим доступа: <https://www.vuforia.com/content/vuforia/en/ar-augmented-reality-software-company.html> (дата обращения: 06.01.2019).
5. EasyAR [Электронный ресурс]. – Режим доступа: <https://www.easyar.com/view/sdk.html> (дата обращения: 06.01.2019).



Щаникова Каролина Евгеньевна

Год рождения: 1999

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № К3342

Направление подготовки: 45.03.04 – Интеллектуальные системы
в гуманитарной сфере

e-mail: karolina-99@mail.ru



Говоров Антон Игоревич

Год рождения: 1991

Университет ИТМО, факультет инфокоммуникационных технологий,
ассистент

e-mail: antongovorov@gmail.com

УДК 504.052, 004.622

**ОЦЕНКА НЕОБХОДИМОЙ ПЛОЩАДИ ЛЕСНЫХ НАСАЖДЕНИЙ
ДЛЯ ОБЕСПЕЧЕНИЯ ЗАДАННОГО РЕГИОНА КИСЛОРОДОМ С УЧЕТОМ
ПОТРЕБНОСТЕЙ НАСЕЛЕНИЯ**

Щаникова К.Е., Говоров А.И.

Научный руководитель – Говоров А.И.

В работе определена актуальность оптимизации вырубки лесов. Также в работе перечислены параметры, от которых зависит возможность вырубки лесов, и разработана математическая модель расчета одного из них: потребности в лесных насаждениях для конкретного региона с учетом затрат кислорода на дыхание населения и использование автотранспорта.

Ключевые слова: вырубка леса, дефорестация, потребность в кислороде, площадь лесных насаждений, продуцирование кислорода деревьями.

Так как леса – одна из важнейших составляющих биогеоценоза, при вырубке лесов большими массивами наносится вред экосистеме планеты. Для сохранения биологического равновесия в природе возникает потребность в оптимизации вырубки лесов. В первую очередь вырубка лесов должна осуществляться таким образом, чтобы регион, в котором она производится, оставался обеспеченным кислородом согласно затратам его населения в результате жизнедеятельности.

Актуальность рассматриваемой проблемы заключается в том, что определение мест вырубки лесов производится государством из экономических соображений, и влияние дефорестации непосредственно на биосферу при выборе зон обезлесения не учитывается. Более того, за последние десятилетия проблема вырубки деревьев превратилась в одну из глобальных проблем, это говорит о том, что на окружающую среду оказывается негативное влияние.

Способы оптимизации вырубки лесов недостаточно изучены предшественниками. Несмотря на то, что были созданы технологии, описывающие состояние лесов в мире, какие-либо указания по оптимальной с биологической точки зрения вырубке деревьев отсутствуют. Уже представленные технологии лишь отображают качественные свойства лесов и их целевое использование человеком, но не учитывают оптимальность происходящей на определенном регионе вырубки с точки зрения экологии.

К факторам, влияющим на определение возможности дефорестации в выбранном регионе, относятся:

- потребность в кислороде;
- степень обеспеченности лесными насаждениями представителей биологических видов;
- использование лесов в производстве;
- оборот леса (годовой прирост леса минус годовые потери);
- климат, рельеф, почва и прочие особенности местности [1];
- прочие факторы.

Целью данной работы являлось определение параметра, отражающего потребность в площади лесных насаждений для данного региона с учетом затрат его населения на процесс дыхания и расхода кислорода в результате пользования автотранспортом. Этот параметр необходим в дальнейших исследованиях для расчета общего коэффициента, определяющего допустимость вырубки лесов в данном регионе.

Чтобы найти площадь лесных насаждений, необходимую населению заданного региона для процесса дыхания в течение года, необходимо найти отношение массы кислорода, необходимой всему населению рассматриваемого региона для дыхания в течение года, к массе кислорода, выделяемого деревьями в данном регионе [2].

Масса кислорода в килограммах, необходимая населению заданного региона для дыхания в течение года, определяется следующими параметрами [2]: t_{m_y} – количество минут в году, $t_{m_y} = 525600$; V_l – средний объем легких человека, $V_l = 0,004 \text{ м}^3$; d – коэффициент обмена воздуха в легких человека; за один вдох поглощается 0,5 л воздуха [3], следовательно, это составляет 1/8 от объема легких, $d = 0,125$; F – среднее количество вдохов в минуту, $F = 14$ [3]; n_{O_2} – концентрация кислорода в воздухе, $n_{O_2} = 0,209$ [3]; $n_{O_{2ex}}$ – концентрация кислорода в выдыхаемом человеком воздухе, $n_{O_{2ex}} = 0,167$ [2]; p_{O_2} – плотность кислорода в нормальных условиях, $p_{O_2} = 1,42$ [2]; N_p – население региона.

Масса кислорода в килограммах, выделяемого деревьями в данном регионе, определяется суммой произведений следующих параметров: m_r – масса кислорода, выделяемая 1 га деревьев заданной породы в год; $P(m_r)$ – доля деревьев заданной породы по отношению ко всему лесному массиву рассматриваемого региона.

Тогда площадь лесных насаждений S_p в гектарах, необходимая населению заданного региона для дыхания в течение года, определяется следующей формулой:

$$S_p = \frac{t_{m_y} V_l d F (n_{O_2} - n_{O_{2ex}}) p_{O_2} N_p}{\sum m_r P(m_r)}. \quad (1)$$

В результате подстановки в формулу (1) известных величин получаем:

$$S_p = \frac{250,8 N_p}{\sum m_r P(m_r)}.$$

Чтобы найти площадь лесных насаждений, необходимую для воспроизводства кислорода, затраченного в результате сгорания топлива, необходимо найти отношение массы кислорода, необходимого для восполнения затрачиваемого автопарком региона кислорода в течение года, к массе кислорода, выделяемого деревьями в данном регионе [2].

Масса кислорода в килограммах, затрачиваемого в результате пользования населением автотранспортом, определяется следующими параметрами [2]: N_v – общее число транспортных средств в стране; N_p – население региона; N_{op} – население страны; S_{av} – средний годовой пробег автомобиля по региону; q – удельный расход топлива, $q = 0,00014 \text{ м}^3/\text{км}$ [2]; p_f – средняя плотность топлива, $p_f = 750 \text{ кг}/\text{м}^3$ [2]; Q_f – расход кислорода при сжигании 1 кг топлива, $Q_f = 1,338 \text{ кг}$ [2].

Тогда площадь лесных насаждений S_f в гектарах, необходимая населению заданного региона для дыхания в течение года, определяется следующей формулой:

$$S_f = \frac{N_v \frac{N_p}{N_{op}} S_{av} q p_f Q_f}{\sum m_r P(m_r)}. \quad (2)$$

В результате подстановки в формулу (2) известных величин получаем:

$$S_f = \frac{0,14 N_v \frac{N_p}{N_{op}} S_{av}}{\sum m_r P(m_r)}.$$

Таким образом, общая площадь лесных насаждений S_o в га, необходимая рассматриваемому региону для восполнения затраченного кислорода равна сумме площадей лесов, необходимых населению для дыхания и использования автотранспорта:

$$S_o = S_p + S_f = \frac{N_p (250,8 + 0,14 \frac{N_v}{N_{op}} S_{av})}{\sum m_r P(m_r)}.$$

Тогда для нахождения общего значения необходимой площади лесных массивов достаточно знать население региона, общее число транспортных средств в стране, население страны, средний годовой пробег автомобиля по региону, массу кислорода, выделяемую деревьями заданной породы, и долю деревьев заданной породы.

В данной работе была определена актуальность оптимизации вырубki лесов, перечислены параметры, от которых зависит возможность вырубki лесов как таковая, и разработана модель, отражающая потребность в лесных насаждениях для конкретного региона с учетом затрат кислорода на дыхание населения и использование автотранспорта.

В перспективе дальнейших исследований планируется расчет прочих перечисленных в данной работе параметров, определяющих возможность безопасной дефорестации в заданном регионе, с целью разработки системы, позволяющей построить карту, отражающую степень возможности вырубki лесных массивов в указанной области.

Литература

1. Руководство по организации и проведению рубок в лесах Республики Беларусь. – М.: МЛХ РБ, 2003. – 27 с.
2. Мельцаев И.Г., Сорокин А.Ф., Мурзин А.Ю. Методические указания для практических занятий по экологии. – Иваново: ИГЭУ, 2011. – С. 63–68.
3. Физиология человека. В 3-х т. – Т. 2. / Пер с англ. / Под ред. Р. Шмидта и Г. Тевса. – М.: Мир, 1996. – 313 с.

**Ямщиков Денис Викторович**

Год рождения: 1991

Университет ИТМО, факультет инфокоммуникационных технологий,
студент группы № К4220Направление подготовки: 11.04.02 – Инфокоммуникационные
технологии и системы связи

e-mail: lineonly@mail.ru

**Осипов Никита Алексеевич**

Год рождения: 1972

Университет ИТМО, факультет инфокоммуникационных технологий,
к.т.н., доцент

e-mail: nikita@ifmo.spb.ru

УДК 004.75**БЛОКЧЕЙН КАК ТЕХНОЛОГИЯ ТРАНЗАКЦИЙ: ВОЗМОЖНОСТИ РАЗВИТИЯ
В СФЕРЕ ТРАНСНАЦИОНАЛЬНЫХ ПЛАТЕЖНЫХ СИСТЕМ****Ямщиков Д.В.****Научный руководитель – к.т.н., доцент Осипов Н.А.**

Работа выполнена в рамках темы НИР «Анализ безопасности транзакций в финансовой сфере с применением технологии блокчейн».

В работе рассмотрены основные понятия технологии блокчейн, описан процесс ведения учета транзакций. Изучено внедрение рассматриваемой технологии международными компаниями Visa и MasterCard, а также раскрыты преимущества и недостатки блокчейн по сравнению с платежными системами в области проведения транзакций. Более того, в работе проанализирована возможность опосредованного конкурентирования технологии блокчейн с транснациональными платежными системами, а также сформулирован наиболее реалистичный вариант развития в данной сфере.

Ключевые слова: блокчейн, криптовалюта, банки, финансовая сфера, распределенная база данных, транзакции, платежные системы.

Технологию блокчейн обычно отождествляют с открытым и распределенным реестром. В качестве одной из экономических интерпретаций можно представить данную технологию в виде аналога бухгалтерской книги. Схожим с данным сравнением свойством обладает блокчейн, который использует структуру данных только для добавления записей. Это означает, что новые транзакции и данные могут быть добавлены, но прошлые записи не могут быть удалены. Таким образом, механизм технологии строится на постоянном добавлении проверенных данных и регистрации транзакций между двумя или более сторонами. Более того, немалая роль отводится децентрализации данной системы [1]. Совокупность описанных факторов способствует как росту транспарентности и подконтрольности на технологическом уровне, так и позитивному укреплению социальной и экономической системы в целом.

Блокчейн создается путем запуска программного обеспечения и связывания нескольких узлов. Рассматриваемая технология не является одной глобальной сущностью как, например, сеть Интернет – существует несколько блокчейнов. Механизм консенсуса и система вознаграждений необходимы для поддержания целостности и функциональности блокчейна. Так, например, в блокчейне биткойнов консенсус достигается путем «майнинга», а система

вознаграждений – это протокол, который присуждает майнеру некоторое количество биткойнов при успешном майнинге блока. Майнинг осуществляется мощными компьютерами, решающими сложные математические задачи. После того как транзакция проверена и принята как истинная для всей сети, майнеры начинают работать над следующим блоком. Таким образом, блокчейн продолжает расти (связывая каждый новый блок с предыдущим) [2].

Ведение учета транзакций является основной функцией всех предприятий. Эти записи предназначены для отслеживания прошлых результатов и являются основой прогнозирования и планирования на будущее. Создание записей большинства организаций занимает много времени и усилий, и часто процессы создания и хранения данных подвержены ошибкам. В настоящее время транзакции могут быть выполнены немедленно, но расчет может занять от нескольких часов до нескольких дней. Например, субъект, продающий акции в корпорации на бирже, может продать их немедленно, но расчет и урегулирование может занять несколько дней. Аналогичным образом, сделка по покупке дома или автомобиля может быть заключена и подписана быстро, но процесс регистрации (проверка и регистрация изменения в праве собственности на недвижимость) часто занимает несколько дней и может включать в себя адвокатов и государственных служащих. В каждом из этих примеров каждая сторона ведет свою собственную бухгалтерскую книгу и не может получить доступ к бухгалтерской книге других вовлеченных сторон.

С использованием технологии блокчейн, процесс проверки и записи транзакции является немедленным и постоянным. «Бухгалтерская книга» распределена по нескольким узлам, что означает, что данные реплицируются и мгновенно сохраняются на каждом узле в системе. Когда транзакция записывается в блокчейн, детали транзакции, такие как цена, актив и право владения, записываются, проверяются и рассчитываются в течение нескольких секунд для всех узлов. Подтвержденное изменение, зарегистрированное в одной бухгалтерской книге, также одновременно регистрируется во всех других копиях бухгалтерской книги. Поскольку каждая транзакция прозрачно и постоянно регистрируется во всех бухгалтерских книгах, открытых для всеобщего обозрения, нет необходимости в проверке третьей стороной [3].

Что же касается развития блокчейн в сфере таких транснациональных платежных систем, как Visa и Mastercard, то данные компании не остаются в стороне от столь стремительно развивающейся технологии.

Так, например, будучи уверенными в интеграции данной технологии в будущее коммерции, банковской сферы, работы предприятий торговли и сферу электронных платежей, компания Mastercard уже изучает ее способы применения и возможности. Директор по развитию цифровых технологий и бизнеса платежной системы Mastercard в России Михаил Федосеев обозначил основную цель развития в данном направлении: предложение удобных и простых решений для потребителей и бизнеса путем объединения знания локальных рынков, экспертизы партнеров и многолетнего опыта. По его словам, компания уже предприняла ряд важных шагов на пути к интеграции блокчейна в технологические и бизнес-процессы платежной системы, такие как [3]:

1. подача заявок на более чем 30 патентов, имеющих отношение к рассматриваемой технологии;
2. инвестирование в венчурный холдинг Digital Currency Group;
3. помощь в развитии блокчейн-стартапов в рамках программы Start Path Global;
4. наличие двух доступных API-интерфейсов (Blockchain Core API и Smart Contracts API) на блокчейн, которые уже используются разработчиками компании Mastercard для осуществления деятельности в сфере распределенного реестра и смарт-контрактов.

Что касается Visa, не менее крупной корпорации на рынке транснациональных платежных систем, то их специалисты отмечают, что инновационность и стремительное развитие технологии блокчейн может уже в ближайшем будущем значительно повлиять на рынок финансовых услуг. Более того, они не исключают возможность внедрения рассматриваемой

технологии в бизнес клиентов и сервисы Visa уже в ближайшие годы. Внимание на данную технологию Visa обратила еще на начальном этапе ее развития и осуществила ряд следующих шагов на пути к развитию блокчейн в рамках своей деятельности [4]:

1. инвестирование в 2015 году в стартап Chain Inc, который занимается разработкой и интегрированием данной технологии в финансовые учреждения;
2. формирование рабочей группы по анализу возможности имплементации инновационных решений в рассматриваемой области и возможности интеграции технологии в различные рынки.

Для того чтобы получить понимание того, сможет ли технология блокчейн опосредованно конкурировать с классическими платежными системами, необходимо сравнить их возможности на данный момент [4].

В классическом варианте, платежная система Visa ежесекундно обрабатывает порядка 1700 транзакций. За время проведения картой через терминал осуществляется запрос банковского баланса клиента, проведение авторизации, проверка счета организации-продавца и перевод на него средств. Таким образом, до момента перехода со счета на счет, деньги успевают пройти несколько этапов и побывать в нескольких «руках».

С точки зрения бизнес-аналитики, платежные системы являются посредниками, которые в течение процесса передачи права собственности на деньги играют роль арбитра. С целью проверки фактического владения и наличие денег у клиента, а также подтверждения требований по их передаче, MasterCard и Visa не могут не полагаться на централизованную IT-инфраструктуру. В отличие от них, аналогичные операции в блокчейн проводятся напрямую от одного криптовалютного кошелька к другому. Посредники в таком случае отсутствуют, что позитивно влияет на снижение рисков работы с третьими сторонами и понижает транзакционную стоимость. Таким образом, можно выделить следующие преимущества технологии блокчейн по сравнению с классическими платежными системами [5]:

1. низкая комиссия (в том числе и для крупных транзакций, стоимость проведения которой не связана с суммой, а напрямую зависит от требуемой вычислительной нагрузки);
2. автоматизация платежей на базе смарт-контрактов (как следствие, автоматизация сложных бизнес-моделей);
3. отсутствие посредников (например, отсутствие необходимости в эскроу-счетах при необходимости совершения денежных переводов после доставки товаров/услуг).

Однако стоит отметить, что наряду с очевидными преимуществами применения рассматриваемой технологии при совершении финансовых транзакций, существует также и ряд недостатков, которые в первую очередь связаны с децентрализацией:

1. масштабируемость – от 4 до 7 транзакций в секунду у блокчейна, против 1700–2000 в классических платежных системах;
2. энергозатратность консенсуса Proof-of-Work – из-за безопасности снижется скорость обработки информации;
3. обеспеченность пользователей ресурсами – необходим доступ в сеть Интернет, компьютер или смартфон, а также опыт работы с криптокошельками. Для использования Visa и Mastercard необходимо наличие пластиковой карты и банка.

Что же касается перспектив развития технологии блокчейн, то можно отметить, что росту пропускной способности уделяется немало внимания. Так, например, в ближайшем будущем значительно увеличить скорость обработки транзакций могут позволить такие технологии, как Plasma и Sharding. По оптимистичным прогнозам Виталия Бутерина (создатель Ethereum) скорость может увеличиться с 20 до 1 млн транзакций в секунду. Более того, полная имплементация гибридной версии протоколов Proof-of-Work и Proof-of-Stake, технология Casper, позволит достигнуть скорости в 800 транзакций/секунду [5].

Таким образом, по мнению автора, наиболее реалистичным вариантом развития технологии блокчейн в сфере транснациональных платежных систем можно считать не

полную замену одной технологии другой, а реструктуризацию существующих, которые на данный момент активно внедряют блокчейн в свои технологические процессы.

Стоит отметить, что на данный момент, технология Blockchain все еще находится на ранней стадии становления, а криптовалюты – это только ее первый основной пример использования. Помимо криптовалют, технология блокчейн стремится к модификации способа проведения, регистрации и проверки транзакций. Это, в свою очередь, революционизирует контракты и уменьшит трудности в процессе обмена активами. В течение следующих нескольких десятилетий технология блокчейн будет распространяться через организации и институты и формировать то, как они взаимодействуют друг с другом. Подобно тому, как Интернет продолжает поддерживать новые технологии, стоит ожидать появления новых вариантов использования технологии блокчейн во многих отраслях.

Литература

1. Nakamoto S. Bitcoin: a peer-to-peer electronic cash system [Электронный ресурс]. – Режим доступа: <https://bitcoin.org/bitcoin.pdf> (дата обращения: 06.01.2019).
2. Генкин А.С., Михеев А.А. Блокчейн. Как это работает и что ждет нас завтра. – М.: Альпина Паблишер, 2017. – 592 с.
3. Ray S. Blockchains: The Technology of Transactions [Электронный ресурс]. – Режим доступа: <https://towardsdatascience.com/blockchains-the-technology-of-transactions-9d40e8e41216> (дата обращения: 06.01.2019).
4. Шлыгин И. Mastercard и Visa про блокчейн и криптовалюту [Электронный ресурс]. – Режим доступа: <https://fomag.ru/news/mastercard-i-visa-pro-blokcheyn-i-kriptovalyutu> (дата обращения: 06.01.2019).
5. Лапшина А. Сможет ли блокчейн заменить Visa и MasterCard [Электронный ресурс]. – Режим доступа: <https://letknow.news/publications/smozhet-li-blokcheyn-zamenit-visa-i-mastercard-7457.html> (дата обращения: 06.01.2019).

**НАПРАВЛЕНИЕ
ДИЗАЙН И УРБАНИСТИКА**



Бочкарев Кирилл Олегович

Год рождения: 1996

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С41041

Направление подготовки: 09.04.03 – Прикладная информатика

e-mail: podockonnik@gmail.com



Митягин Сергей Александрович

Университет ИТМО, институт дизайна и урбанистики,
к.т.н.

e-mail: mityagin@corp.ifmo.ru

УДК 004.93'12

ОБНАРУЖЕНИЕ НЕЗАКОННОЙ РЕКЛАМЫ НА ФОТОГРАФИЯХ ФАСАДОВ ЗДАНИЙ

Бочкарев К.О.

Научный руководитель – к.т.н. Митягин С.А.

В работе рассмотрены существующие методы визуальной оценки городского пространства. Был проведен анализ применения методов компьютерного зрения для решения городских задач. Сформулирована проблема незаконной рекламы в городе. Предложен метод автоматизации слежения за соблюдением законодательства в области наружной рекламы.

Ключевые слова: умный город, городские пространства, наружная реклама, компьютерное зрение, машинное обучение, обнаружение.

Визуальная составляющая городских территорий является одной из важнейших характеристик города, ее можно назвать лицом города. Одним из элементов внешнего вида любого современного города является наружная реклама на фасадах зданий. Самой известной концепцией, объясняющей важность визуального состояния территории, является Теория разбитых окон [1]. Джордж Келлинг и Джеймс Вилсон в 1982 году пришли к выводу, что любые безнадзорные нарушения ведут к распаду общественного устройства.

Долгое время исследования в области рекламы в городе велись с позиции того, как город представляет себя в рекламе, как реклама может обрамлять образ города или как коммерческие объекты: торговые центры, рестораны и развлекательные предприятия, могут быть частью «бренда» города. Реклама может нести функцию средства коммуникации, быть неотделимой от места, стать частью архитектуры, как произошло в Лас-Вегасе. Можно даже говорить о целых зданиях в городе как о рекламе, например, об отелях. Реклама может стать частью семиотики городского пространства. Горожане часто не считают в прямом смысле рекламу вокруг них, но реклама имеет сильное влияние вне отдельной кампании. Реклама в городе получает свою собственную биографию, заводит диалог с биографией горожан и становится частью их жизненных ритмов [2].

Из этого можно сделать вывод, что реклама является важной частью городских пространств, и ее влияние на горожан может быть гораздо сильнее и глубже ожидаемого. В связи с этим актуальна задача регулирования рекламы, слежения за соблюдением установленного регламента и исследования влияния рекламы на городские территории и их жителей.

Визуальная составляющая города регулируется большим количеством законов, которые могут отличаться в разных городах и странах. Визуальная составляющая города регулируется большим количеством законов, которые могут отличаться в разных городах и странах. Например, в Сан-Паулу, Бразилия, законом от 2006 года о «Чистом городе» была запрещена вся наружная реклама. В то время как Таймс-сквер в Нью-Йорке невозможно представить без рекламы. Параметры законной рекламы, в частности, в Санкт-Петербурге, определяются постановлениями Правительства Санкт-Петербурга №1135 от 14 сентября 2006 года, №1002 от 20 сентября 2012 года, №576 от 13 июля 2015 года и №40 от 31 января 2017 года.

На данный момент для визуальной оценки городских территорий и контроля соблюдения законодательства есть лишь один метод: аудиты. Группа экспертов приезжает на место, требующее оценки, и проводит там необходимое исследование. Затем результаты обрабатываются, и выносится некоторое заключение. В роли экспертов могут выступать как уполномоченные властью лица, так и простые граждане, а аудит может представлять собой 5 потраченных минут горожанином на осмотр улицы по пути на работу.

Например, общественное движение «Центральный район за комфортную среду обитания» в марте 2018 года самостоятельно провело аудит Гончарной улицы и отправило жалобы почти на все рекламные вывески, присутствующие на улице. По словам инициативной группы 80% рекламных конструкций были установлены незаконно. В результате было демонтировано 22 вывески с двух домов, что заметно изменило внешность улицы.

С появлением панорамных карт ученые разных стран исследовали возможность их применения для проведения аудитов территорий без необходимости присутствовать на этой территории. Развитие областей компьютерного зрения и машинного обучения за последнее десятилетие позволяет говорить о возможности автоматизации всей процедуры аудитов, что позволит сократить затраты на оценку визуальной составляющей городских территорий и в целом ее упростить.

В 2011 году в Нью-Йорке было проведено исследование [3], в ходе которого осмотр 37 кварталов Нью-Йорка вживую сравнивался с исследованием Google Street View этих же районов. На каждый блок было потрачено 75 мин и еще 10 мин было потрачено на подсчет пешеходов. Подсчет пешеходов невозможно провести по панорамным картам, и это их очевидный минус. Тем не менее, они позволяют проводить оценку без траты времени на путь и сэкономили бы каждому аудитору по 33 ч. Согласованность между оценками, сделанными вживую и по фотографиям, для 60,2% оценок составила больше 80%, для 22,3% – от 60 до 80%. Таким образом, панорамные карты могут быть полезным инструментом для оценки городских пространств, если игнорировать параметры среды, которые нельзя запечатлеть на фото: запахи, активности, которые занимают непродолжительное время в жизни среды, сезонные особенности.

В 2014 году в Массачусетском технологическом институте (MIT) был представлен алгоритм StreetScore [4], по фотографии оценивающий безопасность улиц в человеческом восприятии. Обычно для создания карт восприятия большой территории проводится большое количество опросов граждан. В MIT предложили облегчить и ускорить этот процесс с помощью машинного обучения и панорамных карт. Предыдущие исследования [5] показали, что субъективная оценка безопасности территории в значительно большей мере зависит от внешнего вида территории, чем от возраста, пола или происхождения респондента. Таким образом, оценка исключительно по фотографии должна быть весомой.

Для корректной работы алгоритма [4] было необходимо выбрать признаки, которые будут хорошо характеризовать разницу между «безопасными» и «небезопасными» территориями. Было выбрано несколько общих признаков, таких как: Gist, Texton Histograms, Geometric Texton Histograms, Color Histogram, Geometric Color Histograms, HOG2x2, Dense SIFT, LBP, Sparse SIFT Histograms, SSIM. После обработки изображений данными методами с их помощью обучается SVM. В итоге точность предсказаний алгоритма достигает 93,49%.

Стоит отметить, что особые трудности у алгоритма вызывают неординарные места: с необычными архитектурными памятниками или граффити.

Эти исследования показывают, что с помощью GSV или любой другой панорамной карты города можно сократить время, необходимое на проведение аудитов территорий, а методы компьютерного зрения могут качественно заменить визуальную оценку человеком территории. Можно говорить о возможности автоматизации оценки визуального состояния территории: от общего вида зданий, до обнаружения отдельных проблем, например, наличия нелегальной наружной рекламы.

Для проверки гипотезы возможности обнаружения рекламы на панорамных картах (рисунок) были размечены 250 фотографий фасадов зданий Санкт-Петербурга, и по ним дообучена модель TensorFlow mobilenet, изначально обученная на наборе данных СОСО. Несмотря на маленький объем обучающей выборки, такая нейронная сеть способна обнаруживать вывески. В дальнейшем планируется использовать SVM, работающий с конкретными признаками изображения, как в работе [4].

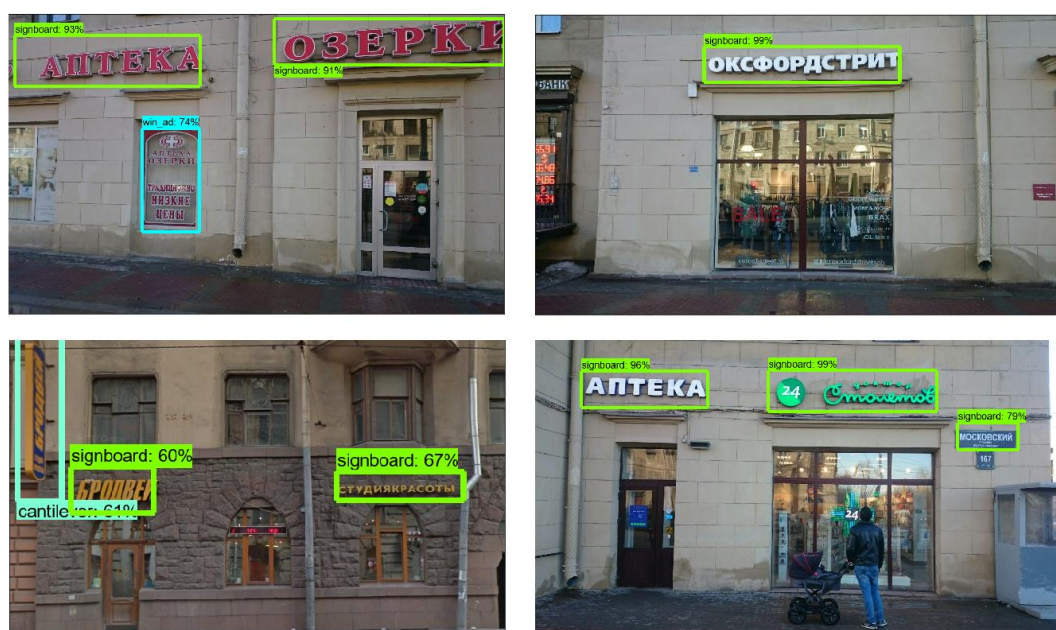


Рисунок. Примеры обнаружения

Литература

1. Wilson J.Q., Kelling G.L. Broken windows: The Police and Neighborhood Safety [Электронный ресурс]. – Режим доступа: https://media4.manhattan-institute.org/pdf/_atlantic_monthly-broken_windows.pdf (дата обращения: 06.01.2019).
2. Cronin A. Advertising and the metabolism of the city: Urban space, commodity rhythms // Environment and Planning D: Society and Space. – 2006. – V. 24. – P. 615–632. Wolf G., Venturi R., Brown D.S., Izenour S. Reviewed Work: Learning from Las Vegas by // Journal of the Society of Architectural Historians. – 1973. – V. 32. – № 3. – P. 258–260.
3. Rundle A.G. et al. Using google street view to audit neighborhood environments // Am. J. Prev. Med. – 2011. – V. 40(1). – P. 94–100.
4. Naik N. et al. Streetscore-predicting the perceived safety of one million streetscapes [Электронный ресурс]. – Режим доступа: http://streetscore.media.mit.edu/static/files/streetscore_paper.pdf (дата обращения: 06.01.2019).
5. Salesses P., Schechtner K., Hidalgo C.A. The Collaborative Image of The City: Mapping the Inequality of Urban Perception // [Электронный ресурс]. – Режим доступа: <https://dspace.mit.edu/bitstream/handle/1721.1/81230/Philip%20Salesses-2013-The%20Collaborative%20Im.pdf?sequence=1&isAllowed=y> (дата обращения: 06.01.2019).

**Вишератин Александр Александрович**

Год рождения: 1991

Университет ИТМО, национальный центр когнитивных разработок

e-mail: alexvish91@gmail.com

**Мухина Ксения Дмитриевна**

Год рождения: 1992

Университет ИТМО, факультет информационных технологий
и программирования, аспирантНаправление подготовки: 09.06.01 – Информатика и вычислительная
техника

e-mail: mukhinaks@gmail.com

**Струков Алексей Максимович**

Год рождения: 1993

Университет ИТМО, институт дизайна и урбанистики,
аспирантНаправление подготовки: 09.06.01 – Информатика и вычислительная
техника

e-mail: as5423.ru@gmail.com

**Насонов Денис Александрович**

Год рождения: 1985

Университет ИТМО, институт дизайна и урбанистики,
к.т.н.

e-mail: denis.nasonov@gmail.com

УДК 004.042**РАЗРАБОТКА ИНТЕЛЛЕКТУАЛЬНОЙ ПЛАТФОРМЫ МОНИТОРИНГА
СОСТОЯНИЯ ГОРОДСКОЙ СРЕДЫ НА ОСНОВЕ РАЗНОРОДНЫХ ДАННЫХ****Вишератин А.А., Мухина К.Д., Струков А.М.****Научный руководитель – к.т.н. Насонов Д.А.**

Работа выполнена в рамках темы НИР № 618274 «Разработка методов мониторинга и прогнозирования состояния городской среды на основе алгоритмов машинного обучения».

Растущая популярность социальных сетей сделала их отличным источником данных для приложений интеллектуального анализа данных, и обнаружение событий не является исключением. Исследователи стремятся не только находить сетевые события, но, что наиболее важно, выявлять и определять местонахождение событий, происходящих в реальном мире. В работе представлена расширенная версия дерева квадрантов – сверточное дерево квадрантов (ConvTree) – и продемонстрировано его преимущество по сравнению со стандартным деревом. Также был представлен алгоритм поиска событий разных масштабов с использованием геопространственных данных, полученных из

социальных сетей. Алгоритм основан на статистическом анализе исторических данных, генерации ConvTree, представляющих нормальное состояние города, и оценке аномалий для обнаружения событий. Экспериментальное исследование, проведенное на наборе данных из 60 миллионов постов в Instagram с геотегами в районе Нью-Йорка, демонстрирует, что предлагаемый подход способен охватывать широкий спектр событий от очень локальных (концерт или вечеринка инди-группы) до городских (бейсбольный матч или праздничный марш), и даже до событий масштаба страны (политические акции протеста или мероприятия в рождественские праздники). Это открывает возможности для создания простой и быстрой, но мощной системы мониторинга событий в реальном времени.

Ключевые слова: дерево квадрантов, поиск событий, социальные сети, мониторинг, свертка, геопространственные данные.

Введение. Социальные сети играют важную роль в современном обществе – люди делятся своими мнениями, эмоциями и опытом, связанными со всеми областями жизни. Информация во всем своем многообразии – текст, фото, видео, геолокация может быть использована для предоставления пользователям полезной информации о происходящем вокруг. Широкий спектр событий различного масштаба, например, концерты под открытым небом, политические демонстрации, автокатастрофы, отражаются в социальных сетях. Вот почему анализ социальных сетей для обнаружения событий набирает обороты в последнее десятилетие. Было продемонстрировано, что во многих случаях, информация о событиях может быть найдена в социальных сетях быстрее, чем в СМИ [1–5].

Тем не менее, обнаружение событий на основе данных из социальных сетей является непростой задачей из-за большого количества шума. Посты, сделанные в близких местах, могут быть не связаны с каким-либо значимым событием. Извлечение информации, которая может быть полезна для идентификации действительно важных событий, требует тщательного анализа пространственно-временных особенностей сообщений, их содержания и дополнительной информации. Еще одной проблемой обнаружения событий в социальных сетях является правильное определение и анализ аномалий.

Обычно эти проблемы решаются с использованием сложных методов кластеризации, которые находят близкие посты с точки зрения пространственно-временного распределения и/или контекста. Однако способ разделения целевой области на части предопределен и не зависит от самих данных, что приводит к чрезмерному или недостаточно подробному разделению и, следовательно, к неэффективному обнаружению событий.

Для решения данной проблемы был предложен метод обнаружения событий, при котором пространственное и временное распределение данных тщательно учитывается. Разработанный подход основан на построении нормальных состояний целевой области для всех комбинаций месяца, типа дня (рабочий день или выходные) и часа, чтобы исключить возможное временное смещение во время пространственного анализа. Нормальное состояние представлено адаптивной географической сеткой – деревом квадрантов, где каждый лист определяет часть целевой области, в которой соблюдаются определенные критерии, такие как среднее число постов или размер самой части. Тем не менее, стандартные деревья квадрантов не учитывают распределение данных, и поэтому была разработана сверточная версия дерева квадрантов – ConvTree, которая использует расположение точек в пространстве при разделении целевой области. Использование ConvTree позволяет точно различать области высокой и низкой активности публикации и, таким образом, повысить чувствительность алгоритма обнаружения событий. Алгоритм обнаружения событий состоит из поиска аномалий, в котором текущая пост-активность в ячейке геосетки сравнивается с базовой линией ячейки, и идентификации события, которая основана на построении связанных компонентов ключевых слов в аномальной ячейке.

Для тщательной оценки разработанного подхода был собран большой набор данных постов Instagram из города Нью-Йорк с 2010 по 2018 год. Полученный набор данных содержит более 62 миллионов постов в 114 000 местоположений. Это позволило создать базисные

адаптивные геосетки на целый год и исследовать эффективность данного решения в двух экспериментах. Экспериментальные результаты показывают, что данный метод достигает точности 73%, что близко или даже выше, чем у самых передовых методов в данной области.

Сверточное дерево квадрантов. Представлено описание алгоритма построения сверточного квадрирования. Было решено использовать дерево квадрантов в качестве структуры данных по двум причинам. Во-первых, дерево квадрантов покрывает всю целевую область за все время. Это важно, потому что для обнаружения событий с использованием статистического анализа ретроспективных данных нужна информация о любой части целевой области для обнаружения аномалий. Если используется, например, R-дерево, которое не покрывает всю целевую зону, в какой-то момент придется расширить одну из ячеек исторической сетки, чтобы покрыть новую область, таким образом, разрушая любой смысл статистических метрик этой ячейки. Во-вторых, иерархическая структура дерева квадрантов, обеспечивает функции увеличения и уменьшения масштаба для последующей локализации событий. Недавние исследования показывают, что использование деревьев квадрантов помогает эффективно обрабатывать большие объемы данных, а также визуализировать события на разных масштабах. Основным недостатком дерева квадрантов является то, что он не учитывает пространственное распределение данных во время разделения. Именно поэтому было принято решение усовершенствовать дерево квадрантов, добавив функциональность поиска точек разделения на основе пространственного распределения данных.

Свертка. Основной принцип построения сверточного дерева квадрантов основан на концепции сверточных нейронных сетей (CNN), которые широко используются в области глубокого обучения. CNN способны извлекать важные признаки из входных изображений посредством набора сверток, активаций и пулов.

В данной работе использовались несколько понятий CNN – сверточное ядро, рецептивное поле, шаг и отступы. Ядро свертки представляет собой небольшую матрицу, значения которой определяют влияние значений рецептивного поля на свернутую величину. Рецептивное поле является частью входной матрицы того же размера, что и ядро свертки. Ядро проходит через все рецептивные поля входной матрицы и выполняет расстановку точек между каждым из них и ядром. Полученная матрица называется картой объектов.

То, как ядро «сканирует» входную матрицу, определяется с помощью шага и заполнения. Шаг контролирует количество элементов, на которое ядро сдвигается. Заполнение округляет входную матрицу с нулями, чтобы сохранить размеры после свертки.

Многомасштабное обнаружение событий. Далее описан алгоритм обнаружения событий с использованием адаптивных геосеток. Адаптивные геосетки, по сути, являются ConvTree, в которых точки являются постами из социальных сетей с привязанными геолокациями. ConvTree, как и любая древовидная структура данных, может легко адаптироваться к содержащимся в них данным. Это означает, что для большого количества точек ConvTree будет разбиваться на множество ветвей и листьев, но при небольшом количестве точек оно будет объединяться в меньшее количество ветвей. Это свойство используется для анализа многомасштабных аномалий и обнаружения событий.

Процесс обнаружения событий состоит из четырех этапов – сбор данных, создание исторических сеток, поиск аномалий и обнаружение событий. На этапе сбора данных собирается информация из социальных сетей за период не менее одного года. Это делается для того, чтобы иметь возможность статистически анализировать все месяцы в году. На этапе создания исторических сеток собранные данные разбиваются на подмножества на основе отметки времени сообщений. Каждое подмножество содержит данные за конкретный месяц, тип дня (рабочий или выходной) и время суток. Для каждого подмножества удаляются выбросы, вычисляется среднее распределение данных и строится ConvTree, который

представляет нормальное состояние города в конкретное время, например, 14:00 рабочего дня в марте. На шаге поиска аномалий анализируемые данные помещаются в соответствующие исторические сетки, и производится поиск ячеек с количеством сообщений, которые значительно отличаются от их стандартного значения. Ячейки с аномальным поведением обрабатываются на этапе обнаружения событий, где события обнаруживаются путем анализа использования хэштегов.

Экспериментальная оценка

1. Сбор данных. Для экспериментальной оценки предложенного подхода был выбран Нью-Йорк в качестве целевой области, потому что в этом городе наибольшая плотность постов в Instagram, и почти каждый его день полон различных событий от грандиозных и всемирно известных, таких как New York Fashion Week, до довольно маленьких и ориентированных на конкретную аудиторию, например, местные кулинарные фестивали и тематические выставки. Оба эти факта означают, что возможно собирать огромные объемы данных, связанных с событиями любых типов и размеров, что является наиболее благоприятными условиями тестирования предложенного подхода к обнаружению событий.

Используя разработанный сканер, удалось собрать данные из более чем 114 872 мест за период до 7 лет. Собранный период времени варьируется от одного местоположения к другому, так как сбор происходил несколько раз. Исходные данные, собранные из Instagram, имели объем 118 ГБ. Общее количество сообщений, извлеченных из района Нью-Йорка, составляет 67 486 683. Для данного исследования не загружались изображения из постов, но в будущем планируется использовать их в качестве дополнительного источника информации о содержании поста и его географическом расположении. Для каждого поста было извлечено время его создания, координаты (широта и долгота) и описание.

2. Экспериментальная оценка. В первом эксперименте было исследовано качество обнаружения указанного набора событий различных типов – конвенций, медиа-событий и т.д. Для этих целей был использован список основных событий в Нью-Йорке за 2017 год, созданный популярным веб-сайтом BizBash, который включает в себя 16 типов событий. Хотя этот набор создается редакторами сайта и может быть произвольным, это наиболее полный источник ретроспективных коллекций событий за целый год. Категории, соответствующие профессиональной деятельности и предназначенные для конкретной узкой аудитории, были исключены из исследования. В силу своей специфики такие типы как реклама, Public relations (PR), бизнес, гостиничный бизнес и политика, не могут получить достаточного освещения в социальной сети и, следовательно, такие события не могут быть обнаружены через Instagram. В результате удалось определить 54 события из 70, находящихся в списке, что составляет 77,1%.

Для второго эксперимента было выбрано две недели 2018 года и вручную проверено, какая часть событий, обнаруженных предложенным методом, была фактической. Этот эксперимент был разработан, чтобы проверить способность системы найти значимые события в большом количестве шума посторонних сообщений. Первый временной диапазон охватывает период с 12 марта 2018 года 00:00 до 18 марта 2018 года в 23:59. В течение этой недели произошло несколько знаменательных событий, таких как День Святого Патрика (17 марта), который является большим праздником с массовым празднованием в Нью-Йорке. В течение недели 291 аномалия была идентифицирована как событие, и 238 из которых (приблизительно 82%) были связаны с фактическими событиями в Нью-Йорке в этот период. Эхо (отзывы и впечатления) Armory Show, ежегодной выставки современного искусства, проходившей 8–11 марта, было обнаружено в начале целевой недели. Данный метод успешно идентифицировал события различных типов и масштабов от изменений погоды (неожиданный снегопад 7, 13 марта), до новой работы Бэнкси (15 марта). Важно отметить, что, несмотря на низкое присутствие политических тем в Instagram, огромные социальные события, такие как

акции протеста, шествия или выборы привлекают большое внимание, и эти события могут быть успешно обнаружены с самого начала.

Вторая неделя (19–25 февраля) была выбрана потому, что не была наполнена событиями. Это было сделано для того, чтобы определить, как алгоритм находит события во время спокойного течения жизни. Единственным крупным событием за этот промежуток времени стала Североамериканская международная ярмарка игрушек, которая прошла 17–20 февраля. В результате было найдено 571 событие, а подтверждено оказалось 417, что составляет 73%.

Заключение. В работе представлена сверточная версия дерева квадрантов ConvTree, которая расширяет стандартную функциональность дерева квадрантов с учетом пространственного распределения данных во время генерации дерева. Также был предложен метод обнаружения событий на основе исторических данных и ConvTree.

Результаты экспериментальной оценки предлагаемого подхода показывают, что адаптивные геосетки достигают высоких показателей как до, так и после отзыва, что открывает перспективы их использования для разработки систем мониторинга в реальном времени.

Литература

1. Atefeh F., Khreich W. A survey of techniques for event detection in Twitter [Электронный ресурс]. – Режим доступа: <http://doi.wiley.com/10.1111/coin.12017> (дата обращения: 06.01.2019).
2. Boettcher A., Lee D. EventRadar: A real-time local event detection scheme using Twitter stream [Электронный ресурс]. – Режим доступа: <http://dx.doi.org/10.1109/GreenCom.2012.59> (дата обращения: 06.01.2019).
3. Chen L.C., Papandreou G., Kokkinos I., Murphy K., Yuille A.L. DeepLab: Semantic Image Segmentation with Deep Convolutional Nets, Atrous Convolution, and Fully Connected CRFs // IEEE Transactions on Pattern Analysis and Machine Intelligence. – 2018. – V. 40. – № 4. – P. 834–848.
4. Zhang C., Lei D., Yuan Q., Zhuang H., Kaplan L., Wang S., Han J. GeoBurst+: Effective and Real-Time Local Event Detection in Geo-Tagged Tweet Streams // ACM Transactions on Intelligent Systems and Technology. – 2018. – P. 1–24.
5. Gao Y. and Zhao L. Incomplete Label Multi-Task Ordinal Regression for Spatial Event Scale Forecasting // AAAI. – 2018. – P. 2999–3006.



Вычужанин Павел Витальевич

Год рождения: 1995

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С4217с

Направление подготовки: 01.04.02 – Прикладная математика
и информатика

e-mail: pavel.vychuzhanin@gmail.com



Калюжная Анна Владимировна

Год рождения: 1990

Университет ИТМО, институт дизайна и урбанистики,
к.т.н., доцент

e-mail: kalyuzhnaya.ann@gmail.com

УДК 004.021

РОБАСТНАЯ КАЛИБРОВКА ПАРАМЕТРОВ ЧИСЛЕННОЙ МОДЕЛИ ВЕТРОВОГО ВОЛНЕНИЯ SWAN

Вычужанин П.В., Калюжная А.В.

Научный руководитель – к.т.н., доцент Калюжная А.В.

Работа выполнена в рамках темы НИР № 618265 «Интеллектуальные методы и технологии поддержки жизненного цикла цифровых образов природных и социальных систем».

Задача адаптации численных моделей ветрового волнения под различные условия моделирования может быть решена различными методами калибровки. Однако в условиях низкого пространственно-временного охвата имеющихся исторических измерений волн полученный набор параметров модели может быть сильно приспособлен к каким-то специфичным сценариям и на практике значительно снижать качество долгосрочных прогнозов. В работе предложен алгоритм робастной калибровки, основанный на алгоритме SPEA2, позволяющий находить более устойчивые для различных сценариев параметры модели. Алгоритм был протестирован на задаче калибровки модели SWAN для региона Карского моря.

Ключевые слова: калибровка моделей, модель ветрового волнения SWAN, робастная оптимизация, эволюционный алгоритм.

Для различных задач в области строительства в открытом море, а также прибрежного судоходства необходимо применять региональные численные модели ветрового волнения для воспроизведения исторических экстремальных явлений и прогнозирования потенциальных опасностей. Для получения прогнозов желаемого качества необходимо определять подходящие физические параметры моделей для конкретных условий моделирования.

Калибровка числовой модели ветрового волнения океана включает согласование результатов моделирования с локальными измерениями, а также спутниковыми данными. Целью калибровки является идентификация набора физических параметров, которые позволяют минимизировать расхождение между моделью и наблюдениями.

Однако даже для эксперта в области гидрометеорологии, ручная калибровка модели представляется достаточно трудоемкой задачей. Современные модели ветрового волнения требуют значительных вычислительных ресурсов, и каждый запуск модели может занять несколько часов. Кроме того, значительно низкий пространственно-временной охват имеющихся исторических измерений волн и низкое качество атмосферных реанализов в

некоторых регионах (например, в Арктических морях, в частности в регионе Карского моря) затрудняет надежную проверку набора параметров. Полученные параметры с минимальным расхождением с наблюдениями могут быть, в значительной степени, сильно приспособлены к каким-то специфичным сценариям, и в случае чрезмерного соответствия небольшому количеству наблюдаемых точек данных на практике могут фактически снижать качество результатов долгосрочного моделирования в ненаблюдаемых местоположениях или временных диапазонах [1].

В настоящее время современные атмосферные реанализы все еще имеют проблемы с качеством в Арктическом регионе [2]. В работе предложен алгоритм, который устанавливает искусственное разнообразие в поля скорости ветра, которые далее используются для генерации вероятностного ансамбля входных полей ветра для учета влияния неопределенности в поверхностном давлении. Далее, для достижения компромисса между надежностью и производительностью оптимизированной модели, используется многоцелевая функция пригодности, на основе которой производится отбор наилучших индивидов в ходе эволюционной оптимизации.

Настройка гидродинамической модели путем подбора параметров модели (или калибровки модели) может быть сформулирована как задача оптимизации. Для этой цели целесообразно представить процесс моделирования в общих математических обозначениях:

$$Y = \{Y_1, Y_2, \dots, Y_k\} = M(\xi | \theta), \quad (1)$$

где Y – многомерные выходные данные (моделируемые поля, например, высоты волн); M – оператор модели, который преобразует граничные начальные условия при установленном наборе параметров θ .

При этом настройка параметров модели (или калибровка модели) может быть формализована с точки зрения многоцелевой оптимизации в пространстве параметров модели и записана в виде:

$$\begin{aligned} \theta_{opt} &= \operatorname{argmin}_{\theta} F(\theta), \\ F(\theta) &= G(f_i(\theta, Y, \{x, y\})), \end{aligned} \quad (2)$$

где G – оператор, преобразующий значения целевых функций f_i в пространственных координатах $\{x, y\}$ в F . В случае ретроспективы характеристик ветровых волн, плохой временной и пространственный охват наблюдений значительно усложняют оптимизацию модели.

Избыточное соответствие решения конкретным случаям, представленным в небольших выборках данных, может привести к неоптимальной конфигурации модели для случаев с другими внешними условиями. Одним из способов повышения надежности результатов оптимизации является расширение набора обучающих данных новыми экземплярами, внося в них относительно небольшие искусственные возмущения. Это позволяет учитывать факторы неопределенности моделирования.

Неопределенность в моделях ветрового волнения может быть представлена не только возмущениями в расчетных переменных [3]. Существуют отклонения в переменных среды, которые могут быть представлены посредством разнообразия входных наборов данных (для модели Simulating WAves Nearshore (SWAN) наиболее важным является атмосферный форсинг, полученный в результате атмосферного реанализа). В этом случае входные данные ξ преобразуются в реализацию ансамбля $\{\xi\}_n = \{\xi_1, \dots, \xi_n\}$ путем добавления искусственных возмущений (или шума) в уравнение (2) так, что оно преобразуется в:

$$\begin{aligned} \theta_{rob} &= \operatorname{argmin}_{\theta} \tilde{F}(\theta | \{\xi\}_n), \\ \tilde{F}(\theta | \{\xi\}_n) &= \mathcal{G}(\tilde{f}_i(\theta | \{\xi\}_n, Y, \{x, y\})). \end{aligned} \quad (3)$$

Целевая функция ансамбля $\tilde{f}_i$ определяет поверхности целевой функции в пространстве параметров с учетом ансамбля входных состояний $\{\xi\}_n$.

В качестве примера гидродинамической модели для экспериментальных исследований была выбрана волно-ветровая модель третьего поколения SWAN. Ветровые волны – это поверхностные волны в океанах и морях, вызванные взаимодействием водных масс и ветра на уровне моря. Модели ветровых волн третьего поколения (например, SWAN) позволяют моделировать волновые спектры и восстанавливать характеристики волн (например, высоты, периоды, направления).

Для модели SWAN наиболее значимыми параметрами [4] являются: энергия ветра, которая характеризуется функцией сопротивления (DRG), диссипация волн (функция разбивания волн (STMP) и нижнее трение (CFW). Поток энергии от нелинейных взаимодействий относительно невелик и не учитывался в данной работе.

В качестве базового решения задачи калибровки модели был выбран широко известный алгоритм многокритериальной оптимизации SPEA2 [5]. С точки зрения эволюционных алгоритмов, в нашем случае каждый индивид соответствует генотипу, представленному в качестве набора значений параметров модели, а фенотипом (значения целевой функции) являются ошибки прогнозирования модели, соответствующие этим параметрам. На каждой итерации эволюции оптимальное по Парето множество индивидов выбирается в соответствии со значениями функции приспособленности, а все не доминирующие решения сохраняются в архиве. Затем формируется набор для скрещивания путем бинарного турнирного отбора. Полученное множество индивидов подвергается рекомбинации и мутации. Итоговое множество становится новой популяцией на следующей итерации алгоритма.

Основным недостатком базового алгоритма является то, что переменные модели являются оптимальными именно для конкретных условий, которые использовались при оценке фитнес-функции. Это позволяет максимизировать производительность для наблюдаемого случая, однако найденное решение может быть нестабильным даже после небольших изменений внешних условий.

Более надежный подход к оптимизации параметров модели может быть реализован с использованием ансамбля волновых моделей, сконфигурированных с использованием различных входных наборов данных. Мы можем сформировать стохастический ансамбль волновых моделей с возмущениями в полях форсингов ветра и искать более надежные значения параметров модели, используя этот ансамбль вместо одной модели с определенным форсингом.

Для этого можно адаптировать базовый алгоритм SPEA2 (который был представлен выше), изменив стратегию вычисления функции приспособленности: для данного генотипа оценивается набор фенотипов, соответствующих элементам ансамбля и на основе его значений вычисляется устойчивая метрика. Схема работы предлагаемого алгоритма представлена на рис. 1.

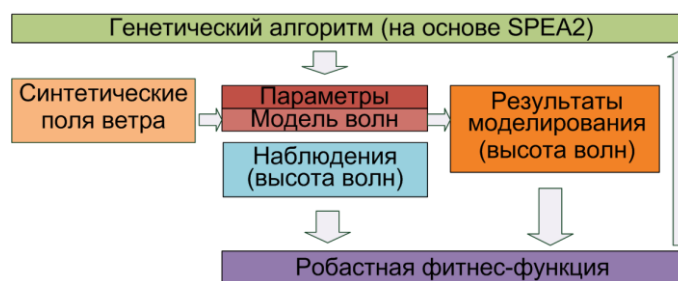


Рис. 1. Схема работы робастного алгоритма калибровки

В качестве примера была взята задача калибровки конфигурации модели SWAN для региона Карского моря. В качестве целевой переменной было взято значение высоты волны (H_{sig}). Кроме того, были проанализированы результаты в 9 репрезентативных точках (P1–P9, представленные на рис. 2), чтобы учесть возможную пространственную изменчивость оптимального решения.

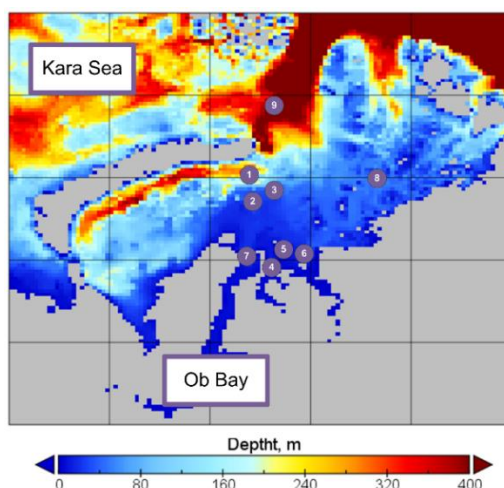


Рис. 2. Часть батиметрии исследуемой области моделирования: Карское море и Обская губа

Был проведен ряд экспериментов для сравнения результатов экспериментов по оптимизации. Начальная популяция для обоих подходов калибровки была произведена с использованием латинского отбора из гиперкуба (LHS) в пространстве параметров. Параметры обоих алгоритмов калибровки были выбраны следующим образом: размер популяции составляет 20 особей, количество поколений – 60, размер архива – 5 особей, вероятность мутации и кроссовера – 0,2.

Для оценки качества результатов модели были выбраны две целевые функции: средняя абсолютная ошибка (MAE) и среднеквадратичная ошибка (RMSE). Калибровка для каждого сценария повторялась 100 раз, чтобы получить распределение относительного улучшения RMSE и MAE по сравнению с конфигурацией модели со значениями параметров по умолчанию (DRF=1,0, CFW=0,015, STPM=0,00302). Также было учтено среднее относительное стандартное отклонение откалиброванного набора параметров. Результаты представлены на рис. 3.

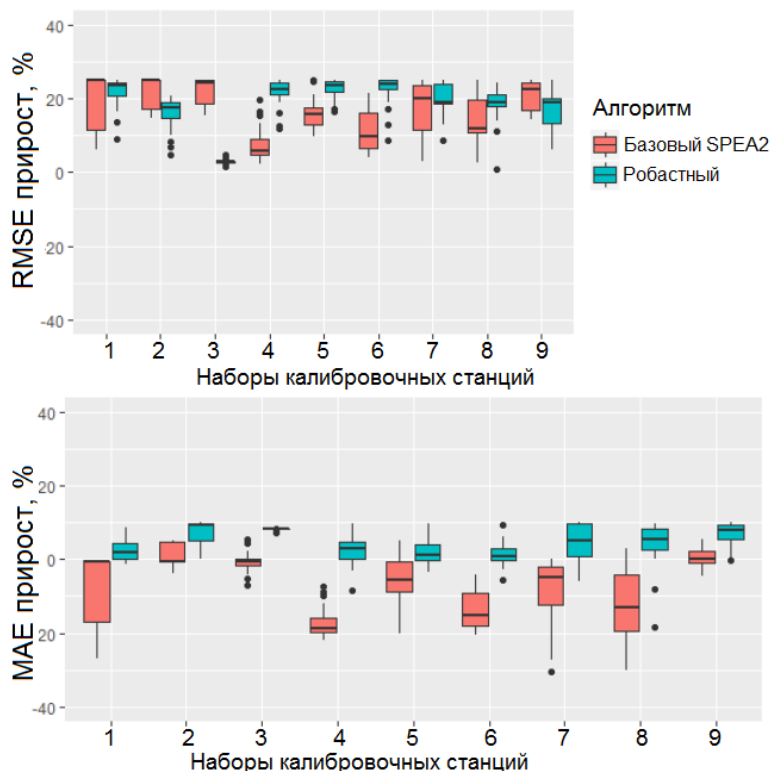


Рис. 3. Сравнение производительности базового и робастного алгоритмов на проверочном наборе станций

Можно заметить, что робастный подход обеспечивает лучшее или эквивалентное улучшение производительности модели на проверочных точках во всех группах сценариев. Среднеквадратичное отклонение, как для параметров модели, так и относительных значений улучшения, также ниже, чем у базового алгоритма. В среднем прирост составляет около 1–2% RMSE и 10% MAE по отношению к значениям базового алгоритма. Среднеквадратичное отклонение (СКО) полученных метрик меньше для всех сценариев, также как и для СКО параметров модели. Мы можем утверждать, что робастный подход эффективен для случая с несколькими пространственно-распределенными точками, которые можно применять для калибровки. Важно отметить, что качество точек калибровки не ухудшают результат.

Литература

1. Brynjarsdóttir J., O'Hagan A. Learning about physical parameters: The importance of model discrepancy // *Inverse Problems*. – 2014. – V. 30. – № 11. – P. 114007.
2. Fredriksen L.E. An evaluation of the reanalyses ERA-Interim and ERA5 in the Arctic using N-ICE2015 data [Электронный ресурс]. – Режим доступа: <https://munin.uit.no/bitstream/handle/10037/13509/thesis.pdf?sequence=2&isAllowed=y> (дата обращения: 06.01.2019).
3. Paenke I., Branke J., Jin Y. Efficient search for robust solutions by means of evolutionary algorithms and fitness approximation // *IEEE Transactions on Evolutionary Computation*. – 2006. – V. 10. – № 4. – P. 405–420.
4. Nikishova A. et al. Uncertainty quantification and sensitivity analysis applied to the wind wave model SWAN // *Environmental modelling & software*. – 2017. – V. 95. – P. 344–357.
5. Zitzler E., Laumanns M., Thiele L. SPEA2: Improving the strength Pareto evolutionary algorithm [Электронный ресурс]. – Режим доступа: <https://pdfs.semanticscholar.org/6672/8d01f9ebd0446ab346a855a44d2b138fd82d.pdf> (дата обращения: 06.01.2019).

**Деева Ирина Юрьевна**

Год рождения: 1994

Университет ИТМО, институт дизайна и урбанистики,
аспирантНаправление подготовки: 09.06.01 – Информатика и вычислительная
техника

e-mail: iriny.deeva@gmail.com

**Калюжная Анна Владимировна**

Год рождения: 1990

Университет ИТМО, институт дизайна и урбанистики,
к.т.н., доцент

e-mail: kalyuzhnaya.ann@gmail.com

УДК 004.942**МНОГОМАСШТАБНОЕ МОДЕЛИРОВАНИЕ ЦИФРОВОГО ОБРАЗА ЧЕЛОВЕКА
В КИБЕРПРОСТРАНСТВЕ****Деева И.Ю.****Научный руководитель – к.т.н., доцент Калюжная А.В.**

В настоящее время предложено множество подходов для автоматического определения личности пользователей по их пользовательскому контенту в социальных сетях. Подходы различаются с точки зрения алгоритмов машинного обучения и используемых наборов характеристик, структуры используемого цифрового следа и среды социальных сетей, используемых для сбора данных. В работе проведен сравнительный анализ современных вычислительных методов распознавания личности на разнообразном наборе данных в социальных сетях из Facebook и ВКонтакте. Даны ответы на три вопроса: следует ли рассматривать прогнозирование психометрик как задачу прогнозирования с несколькими откликами (т.е. все черты личности данного пользователя прогнозируются одновременно), или каждый признак должен идентифицироваться отдельно? какие признаки хорошо работают в разных онлайн-средах? и каково снижение точности при переносе моделей, обученных на данных одной социальной сети, в другую?

Ключевые слова: цифровой образ, цифровой след, цифровой двойник, машинное обучение, модель, психометрика.

Исследования в области психологии показали, что поведение и предпочтения людей могут быть в значительной степени объяснены базовыми психологическими конструкциями: личностными чертами. Знание личности человека позволяет нам делать прогнозы относительно предпочтений в разных контекстах и средах и совершенствовать системы рекомендаций [1]. Личность может влиять на процесс принятия решений, и было показано, что она влияет на предпочтения для веб-сайтов [2], продуктов, брендов и услуг [1], а также для контента, такого как фильмы, телешоу и книги [3].

Наиболее широко принятая модель личности, модель «Большой пятерки», включает в себя пять черт: открытость опыту, добросовестность (сознательность), экстраверсия, доброжелательность и эмоциональная стабильность (часто наоборот, называемая нейротизм). Дальнейшие объяснения каждой черты приведены в табл. 1.

Традиционный подход к измерению личности требует, чтобы участники отвечали на ряд вопросов (обычно от 20 до 360), оценивая их поведение и предпочтения [4]. Этот подход отнимает много времени и нецелесообразен, особенно в контексте онлайн-услуг. Онлайн-пользователи могут не захотеть тратить значительное количество времени на заполнение анкеты, чтобы персонализировать результаты поиска или рекомендации по продукту.

Таблица 1. Обзор модели личности «Большой пятерки»

Черта	Описание
Открытость опыту	Открытость связана с воображением, творчеством, любопытством, терпимостью, политическим либерализмом и ценностью культуры. Люди, которые высоко ценят открытость, любят перемены, ценят новые и необычные идеи и имеют развитое чувство эстетики
Добросовестность (сознательность)	Сознательность измеряет предпочтение организованному подходу к жизни, а не стихийному. Люди с высокими показателями добросовестности имеют больше шансов быть хорошо организованными, надежными и последовательными. Им нравится планировать, стремиться к достижениям и преследовать долгосрочные цели. Недобросовестные люди, как правило, более беспечные, спонтанные и творческие. Они склонны быть более терпимыми и менее связанными правилами и планами
Экстраверсия	Экстраверсия измеряет тенденцию искать стимуляции во внешнем мире, в обществе других и выражать положительные эмоции. Люди, получающие высокие оценки по экстраверсии, обычно более общительны, дружелюбны и социально активны. Они обычно энергичны и разговорчивы; они не против быть в центре внимания и заводят новых друзей легче. Интроверты, более вероятно, будут одиночками и сдержанными, они ищут среду, характеризующуюся более низким уровнем внешней стимуляции
Доброжелательность	Доброжелательность связана с акцентом на поддержание позитивных социальных отношений, дружелюбия, сострадания и сотрудничества. Люди, получившие высокую оценку «доброжелательности», как правило, доверяют другим и приспосабливаются к их потребностям. Менее доброжелательные люди больше сосредоточены на себе, менее склонны к компромиссу и могут быть менее легковверными. Они также имеют тенденцию быть менее связанными социальными ожиданиями и условностями и более напористыми
Эмоциональная стабильность	Эмоциональная стабильность, также называемая невротизмом, измеряет склонность испытывать перепады настроения и такие эмоции как вина, гнев, беспокойство и депрессия. Люди, получившие низкую оценку эмоциональной стабильности (высокий невротизм), чаще испытывают стресс и нервозность, тогда как люди, набравшие высокие оценки эмоциональной стабильности (низкий невротизм), обычно спокойнее и увереннее в себе

Однако недавно было показано, что цифровой след пользователей может использоваться для автоматического определения их личности. Например, Kosinski et al. (2013) и Youyou et al. (2015) показали, что автоматические суждения о личности, основанные на лайках в Facebook, являются более точными, чем оценки, сделанные друзьями пользователей или даже их супругами. Также Shuotian Bai et al. (2015) показали, что подобные прогнозы могут основываться на языке, используемом в социальных сетях. Было предложено множество других подходов с использованием различных механизмов прогнозирования, пространств признаков и сосредоточения внимания на различных онлайн-средах [5].

В данном исследовании проведен сравнительный анализ современных компьютерных методов распознавания личности на разнообразном наборе данных социальных сетей, собранных в Facebook и ВКонтакте. Цель – ответить на три следующих вопроса:

1. следует ли рассматривать прогнозирование личности как задачу прогнозирования с несколькими откликами (т.е. все черты личности данного пользователя прогнозируются

одновременно), или каждый признак должен идентифицироваться отдельно? Учитывая пользовательский контент каждого пользователя, цель состоит в том, чтобы получить набор из пяти оценок (действительных чисел), представляющих измерения «Большой пятерки». Данная задача может быть рассмотрена как задача регрессии, включающая в себя различные методы одномерной и многомерной регрессии. В этой работе авторы сравнивали методы многомерной регрессии, например, комбинирование моделей множественной регрессии, ансамбль цепочек регрессоров и множественные случайные леса, с одномерными подходами, такими как машины опорных векторов и деревья решений. В качестве baseline предсказания для каждой точки данных было рассмотрено среднее значение по обучающим данным;

2. какие признаки хорошо работают в разных онлайн-средах? Был извлечен широкий спектр лингвистических и эмоциональных особенностей из обновлений статуса ВКонтакте. Основное обоснование включения лингвистических и эмоциональных особенностей заключается в том, что люди с разными чертами личности будут выражать себя по-разному и, следовательно, будут использовать разные слова (фразы) и эмоции (гнев, радость). Была произведена оценка силы взаимосвязи между различными прогностическими признаками и личностными чертами, определяя их взаимосвязи. Проведено сравнение результатов корреляции по различным наборам данных. Однако, насколько известно авторам, нет работы, которая сравнивала бы результаты с различными наборами эталонных данных. Цель состояла в том, чтобы определить, какие отношения между особенностями и чертами личности являются общими для различных платформ социальных сетей;
3. каково снижение точности при переносе моделей, обученных на данных одной социальной сети, в другую? Предсказание личностных характеристик является сложной задачей; в отличие от демографических данных, результаты социальных опросов являются относительно скудной информацией и измеряются со значительной ошибкой. Farnadi et al. (2013) предложили перекрестное обучение или разработку моделей прогнозирования личности с использованием различных цифровых сред. Преимущество перекрестного обучения заключается в том, что примеры обучения из разных социальных сетей могут быть объединены для повышения точности других тестовых данных. В работе были изучены возможности перекрестного обучения для прогнозирования личности с использованием наборов эталонных данных из двух разных сред (Facebook и ВКонтакте).

Исследования, представленные в данной работе, используют два набора данных, собранных с самых популярных социальных сетей (Facebook и ВКонтакте). Набор данных Facebook находится в открытом доступе. Набор данных социальной сети ВКонтакте собирался самостоятельно посредством добровольного участия в опросе. Структура обоих наборов данных представляет собой матрицу из шести полей – пять психометрик и количество друзей пользователя в социальной сети. Для подтверждения статистической значимости взаимосвязи указанных полей, была составлена матрица корреляции, приведенная в табл. 2 и 3 для каждого набора данных.

Таблица 2. Матрица корреляции Facebook

sEXT	sAGR	sCON	sNEU	sOPN	friends
1,0	0,2	0,2	-0,4	0,1	0,3
	1,0	0,1	-0,4	0,3	0,0
		1,0	-0,3	0,1	0,2
			1,0	-0,2	-0,2
				1,0	0,0
					1,0

Таблица 3. Матрица корреляции ВКонтакте

sEXT	sAGR	sCON	sNEU	sOPN	friends
1,0	0,2	0,2	-0,2	0,7	0,3
	1,0	0,4	-0,1	0,3	0,1
		1,0	-0,2	0,2	-0,1
			1,0	-0,1	0,1
				1,0	0,1
					1,0

Так как количество друзей показало статистически значимую зависимость от метрики нейротизма (обратную), авторы попробовали предсказать степень выраженности именно этой психометрики по количеству друзей. Для этого была проведена дискретизация психометрики sNEU, что позволило оценивать, насколько точное значение психометрики, сколько степень ее выраженности.

Для создания модели были использованы основные методы машинного обучения, в том числе метод наименьших квадратов, метод построения случайных лесов, метод K-ближайших соседей, логистическую регрессию, а также метод опорных векторов с опорным ядром. В качестве оценки степени влияния каждого предиктора использовалась такая метрика, как коэффициент детерминации (R-квадрат). Результаты представлены в табл. 4.

Таблица 4. Результаты работы моделей машинного обучения

	LinearRegression	RandomForest	KNeighbors	LogisticRegression
R-квадрат	0,061	0,969	0,962	-0,503
MSE	0,482	0,088	0,097	0,609

Таким образом, в этом исследовании выполнен сравнительный анализ современных компьютерных методов распознавания личности на разнообразном наборе достоверных данных социальных сетей из Facebook и ВКонтакте.

Авторы попытались ответить на три вопроса исследования следующим образом. Использовались различные одномерные методы и методы многомерной регрессии. При использовании этих учащихся в двух разных наборах данных модели дерева решений в основном превосходили модели опорных векторов, в то время как учащиеся с многомерной регрессией с деревом решений в качестве базового обучаемого часто опережали одномерные регрессионные.

Литература

1. Lambiotte R., Kosinski M. Tracking the digital footprints of personality // Proceedings of the Institute of Electrical and Electronics Engineers (PIEEE). – 2014. – P. 35–39.
2. Kosinski M., Stillwell, D.J., Graepel T. Private traits and attributes are predictable from digital records of human behavior // Proc. Natl. Acad. Sci. (PNAS). – 2013. – V. 110. – P. 5802–5805.
3. Cantador I., Fernández-Tobías I., Bellogín A., Kosinski M., Stillwell D. Relating personality types with user preferences in multiple entertainment domains [Электронный ресурс]. – Режим доступа: https://www.researchgate.net/publication/283270671_Relating_Personality_Types_with_User_Preferences_Multiple_Entertainment_Domains (дата обращения: 06.01.2019).
4. John O.P., Srivastava S. The Big Five trait taxonomy: history, measurement, and theoretical perspectives // Handb. Pers. Theory Res. – 2008. – P. 132–138.
5. Farnadi G., Sitaraman G., Rohani M., Kosinski M., Stillwell D., Moens M., Davalos S., De Cock M. How are you doing? Emotions and personality in Facebook // Proceedings of the EMPIRE. – 2014. – P. 45–56.

**Деревицкий Илья Владиславович**

Год рождения: 1992

Университет ИТМО, институт дизайна и урбанистики,
аспирантНаправление подготовки: 09.06.01 – Информатика и вычислительная техникаe-mail: ilyaderevitskiy@gmail.com**Ковальчук Сергей Валерьевич**

Год рождения: 1984

Университет ИТМО, институт дизайна и урбанистики,
к.т.н.e-mail: sergey.v.kovalchuk@gmail.com

УДК 004.852

**ПРЕДСКАЗАТЕЛЬНОЕ МОДЕЛИРОВАНИЕ ХРОНИЧЕСКОГО
ИНСУЛИНЕЗАВИСИМОГО ДИАБЕТА****Деревицкий И.В.****Научный руководитель – к.т.н. Ковальчук С.В.**

Работа выполнена в рамках темы НИР № 618267 «Интеллектуальные технологии гибридного предсказательного моделирования в задачах Р4-медицины».

Для здравоохранения сахарный диабет 2 типа представляет одну из наиболее приоритетных проблем, так как данное заболевание связано с большим числом сопутствующих заболеваний, приводящих к ранней инвалидизации и повышенному сердечно-сосудистому риску. Основной задачей исследования являлось построение модели раннего прогнозирования возникновения преддиабета, включающего нарушение толерантности к глюкозе и диабета 2 степени. Данная модель должна оценивать вероятность наличия преддиабета в определенный момент времени.

Ключевые слова: диабет, нарушение толерантности к глюкозе, диабет инсулинезависимый, машинное обучение, глюкоза плазмы в крови.

Исследование выполнялось на данных из 5791 медицинских карт, предоставленных ФГБУ «НМИЦ им. В.А. Алмазова» Минздрава России (Центр Алмазова). Ни у одного пациента на момент начала наблюдения нет преддиабета и сахарного диабета (СД) любого типа. Посмотрим на некоторые характеристики пациентов, на данных о которых будет строиться модель (рис. 1).

Сопутствующие хронические заболевания для диабетиков и недиабетиков

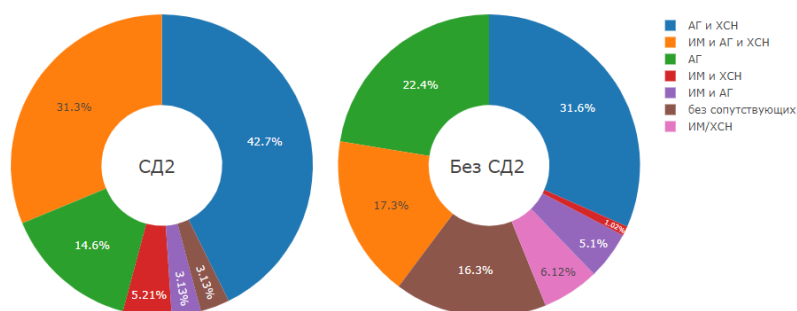


Рис. 1. Сопутствующие хронические заболевания в разрезе наличия сахарного диабета 2 типа

Стоит проверить гипотезу: диабет статистически значимо увеличивает риски хронической сердечной недостаточности (ХСН), инфаркта миокарда (ИМ) и артериальной гипертензии (АГ) (рис. 2).

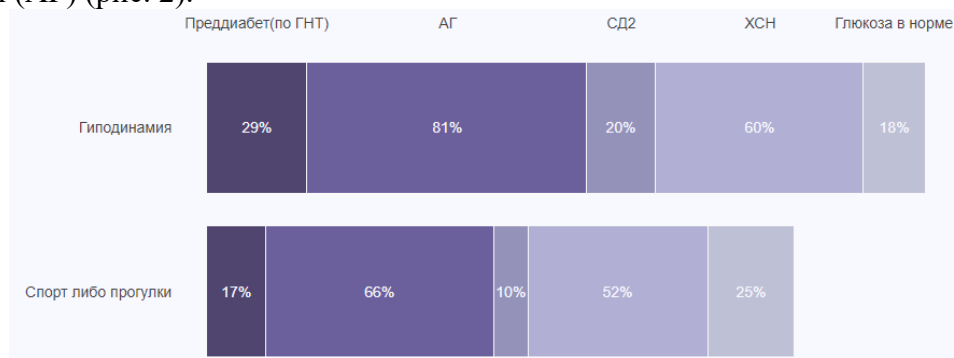


Рис. 2. Исследуемые в разрезе образа жизни – неспортивный (верх) и спортивный

Исследуемые – малоподвижные люди, с множеством сопутствующих сердечных заболеваний. Для таких пациентов стандартные распространенные модели скоринга на наличие диабета либо преддиабета, например FINDRISK [1], будут работать плохо, так как такие модели обучались на среднестатистической выборке исследуемых.

Для таких пациентов необходима более надежная прогнозная модель, предсказывающая вероятность наличия преддиабетических состояний, учитывая сердечные осложнения и неактивный образ жизни. Построим такую модель с помощью методов машинного обучения.

У каждого пациента были взяты два измерения глюкозы в плазме крови натощак, в начале наблюдаемого в Центре Алмазова периода, и в конце.

Задача предсказания наличия диабетического уровня глюкозы у пациента решалась построением бинарного классификатора, определяющего, превысит ли анализ, сделанный в конце наблюдения, преддиабетическую отметку 6,1 ммоль/л для глюкозы плазмы крови натощак, согласно рекомендациям [2]. Модель классификатора выбиралась из моделей RandomForestClassifier, KNeighborsClassifier, LogisticRegression, SVC. Признаки для модели выбирались с помощью одномерного отбора признаков (univariate feature selection) и метода рекурсивного исключения признаков (recursive feature elimination, RFE). Значимыми по score оказались возраст, индекс массы тела, максимальная глюкоза (в начале наблюдения, не превышает 7) и объем талии, RFE подтвердил данный вывод. Также был проведен анализ важности признаков с помощью метода Feature importances.

Модель выбиралась по показателям доли правильных ответов модели (accuracy) показателю точности определения преддиабетиков (precision) и показателю полноты (recall), который показывала ненастроенная методом перебора гиперпараметров по сетке модель на кросс-валидации (средняя оценка каждого показателя) и на отложенной выборке, по которой не обучалась модель.

Лучшая модель – LogisticRegression, параметры модели настроены с помощью метода grid search и кросс-валидации.

Среднее accuracy по кросс-валидации 0,8, полнота recall более 70 для отложенной и обучающей выборки, обе выборки сбалансированы (отношение классов 1:1) и стратифицированы (рис. 3).

Наиболее влияющие на показатель признаки – возраст пациента на момент начала наблюдения, максимальная глюкоза в плазме крови натощак (в момент начала наблюдения делалось несколько измерений, максимальное из них) время между измерениями и объем талии.

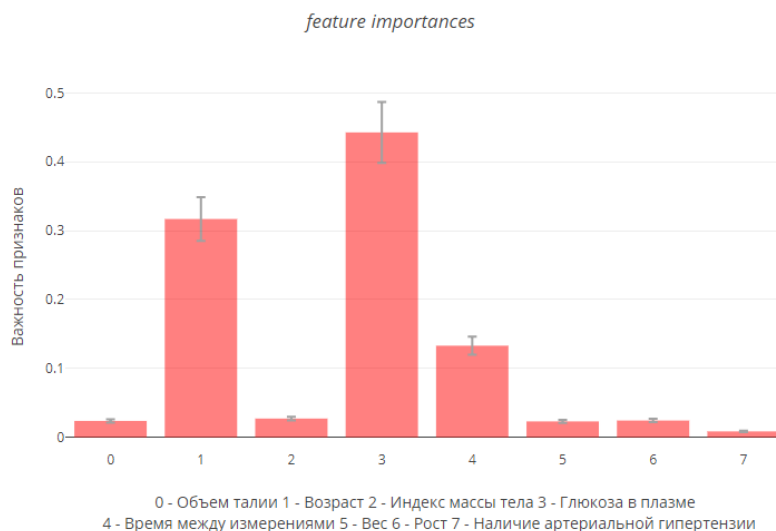


Рис. 3. Анализ важности признаков с помощью Feature Importances

Построенная модель осуществляет прогноз на наличие диабета у пациентов с большим количеством сопутствующих болезней (наиболее частые АГ и ИБС) и делает оценку риска возникновения диабета точнее, чем FINDRISK для группы пациентов, на которой была протестирована модель (рис. 4). FINDRISK явно сильно недооценивает риск развития инсулиннезависимого диабета у данной группы пациентов, построенная модель дает более реальный прогноз для превышения глюкозы выше преддиабетического уровня, который можно использовать для принятия ранних превентивных мер против развития преддиабета и диабета [3].

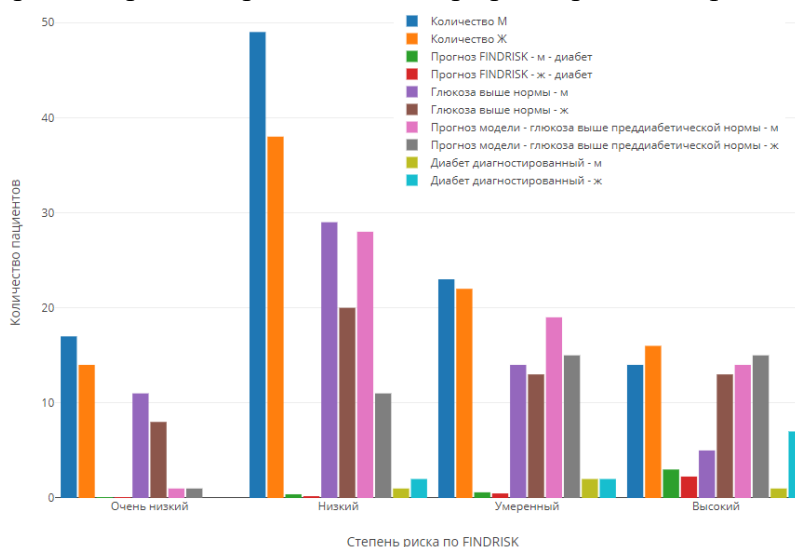


Рис. 4. Наличие диабета, гендерный состав и прогноз шкалы FINDRISK для пациентов, попавших в тестовую выборку

Литература

1. Lindström J., Tuomilehto J. The diabetes risk score: A practical tool to predict type 2 diabetes risk // Diabetes Care. – 2003. – V. 26(3). – P. 725–731.
2. Рекомендации по диабету, предиабету и сердечно-сосудистым заболеваниям. EASD/ESC [Электронный ресурс]. – Режим доступа: https://scardio.ru/content/Guidelines/Diabet_esc_2014.pdf (дата обращения: 06.01.2019).
3. Heianza Y., Hara S., Arase Y., Saito K., Fujiwara K., Tsuji H., Kodama S., Hsieh S.D., Mori Y., Shimano H., Yamada N., Kosaka K., Sone H. HbA1c 5,7–6,4% and impaired fasting plasma glucose for diagnosis of prediabetes and risk of progression to diabetes in Japan (TOPICS 3): a longitudinal cohort study // Lancet. – 2011. – V. 378. – P. 147–155.



Дунаенко Сергей Сергеевич

Год рождения: 1995

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С41011

Направление подготовки: 27.04.07 – Наукоемкие технологии
и экономика инноваций

e-mail: s3.dunaencko@gmail.com

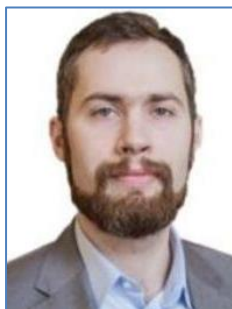


Смирнов Егор Владимирович

Год рождения: 1988

Университет ИТМО, институт дизайна и урбанистики

e-mail: smirnov.egor.v@gmail.com



Митягин Сергей Александрович

Университет ИТМО, институт дизайна и урбанистики,
к.т.н.

e-mail: mityagin@corp.ifmo.ru

УДК 004.942

**ПРИМЕНЕНИЕ МЕТОДОВ АГЕНТНОГО МОДЕЛИРОВАНИЯ ДЛЯ ПОВЫШЕНИЯ
БЕЗОПАСНОСТИ ДВИЖЕНИЯ ПЕШЕХОДОВ**

Дунаенко С.С., Смирнов Е.В.

Научный руководитель – к.т.н. Митягин С.А.

Работа выполнена в рамках темы НИР № 618266 «Интеллектуальные технологии в организации процесса создания умного города».

В работе рассмотрены проблемы пересечения пешеходами дороги в неполюженном месте и оптимальной расстановки пешеходных переходов в условиях современного города. Изучены методы исследования данных проблем, сделан вывод о целесообразности применения метода компьютерного моделирования. Предложен метод определения опасных для пешехода мест и оптимальной расстановки пешеходных переходов, на основе разрабатываемой в Институте дизайна и урбанистики Университета ИТМО симуляции пешеходного движения AntRoadPlanner.

Ключевые слова: мультиагентное моделирование, компьютерная симуляция, пешеходные переходы, пешеходная безопасность, урбанистика, умный город.

Постоянно растущие, меняющиеся и развивающиеся города встречаются множество проблем, связанных с комфортом и безопасностью их жителей. Если рассматривать такую группу пользователей города, как пешеходы, то одна из проблем, не теряющих для них актуальность, это небезопасное пересечение дороги вследствие неоптимальной расстановки пешеходных переходов. Согласно статистике, в России ежегодно около 20% от всего количества дорожно-транспортных происшествий составляют наезды на пешехода вне пешеходного перехода. В США ежегодно около 10% от всех смертей в дорожно-транспортных

происшествиях происходят вследствие наезда на пешехода вне перехода на перекрестке. Данная проблема часто освещается средствами массовой информации [1]. При этом согласно исследованиям поведения пешеходов стремление сократить путь и, возможно, пересечь дорогу даже при отсутствии перехода является частью психологии пешеходов, склонностью к нахождению оптимального пути и не решается запретительными мерами вроде установки заграждений.

Для исследования безопасности пешеходов и повышения их безопасности существует ряд методов. Одним из самых простых методов является учет статистики дорожно-транспортных происшествий. Согласно государственному стандарту в области пешеходных переходов, переход необходимо ставить в местах, в которых произошло три и более аварии с участием пешехода за последние 12 месяцев [2]. Другим способом является измерение данных пешеходной и дорожной инфраструктуры (например, ширина проезжей части, количество полос) и данных о движении пешеходов и автомобилей (интенсивность потока). Этот метод также находит отражение в вышеупомянутом государственном стандарте в виде выделения четырех типов дорог, для каждого из которых ставятся особые требования относительно организации пешеходного перехода. Третьим способом является экспертная оценка. Этот способ упоминается и часто используется в методических рекомендациях по организации пешеходных пространств, изданных Министерством транспорта [3]. Также для повышения безопасности пешеходов и оптимальной установки пешеходных переходов можно пользоваться и другими методическими рекомендациями в данной области, а также указаниями в государственных стандартах. В качестве примера зарубежных методических рекомендаций по установке переходов можно привести рекомендации, принятые в США [4]. Однако такие рекомендации и документы стандартизации чаще всего ограничиваются достаточно общими требованиями, касающимися рекомендуемого среднего расстояния между переходами и необходимостью учитывать устоявшиеся пешеходные пути и точки притяжения пешеходов.

В качестве отдельного метода исследования безопасности пешеходов можно выделить математическое моделирование. Развитием данного метода является компьютерное моделирование и симуляция пешеходного движения. Достоинствами данного метода являются скорость исследования, ограниченная только вычислительной мощностью, и уникальность применения, выражающаяся в возможности адаптации моделей под любую территорию. На данный момент существует множество подходов к моделированию движения пешеходов, но почти все подходы объединяет ряд особенностей и ограничений. Основной целью такого моделирования чаще всего выступает исследование взаимодействия пешеходов с автомобилями на отдельных участках дороги или пешеходных переходах, следствием чего является использование моделей на локально выбранных участках и необходимость точной калибровки при смене такого участка, а также отсутствие учета точек притяжения пешеходов при симуляции.

Для данного исследования было решено использовать программу симуляции пешеходного движения AntRoadPlanner, разрабатываемую в Институте дизайна и урбанистики Университета ИТМО [5]. Программа использует метод мультиагентного моделирования на основе алгоритма поиска пути A*. В качестве исходных данных симуляция использует карту с размеченными препятствиями и элементами пешеходной инфраструктуры, на которой также обозначаются точки притяжения пешеходов (входы в здания, остановки транспорта и пр.). Назначением программы является построение прогнозируемых маршрутов пешеходов на основе движения агентов симуляции между точками притяжения и выделение вытопанных троп, образующихся в результате неоптимальной организации сети пешеходных дорог. Программа позволяет обрабатывать данные о пешеходной сети, как для существующих территорий, так и для территорий на стадии проектирования. Таким образом, данная симуляция представляет собой инструмент, который может помочь выявить недостатки проектируемой пешеходной инфраструктуры еще до ее реализации. Достоинствами

симуляции для данного исследования являются учет точек притяжения пешеходов и возможность моделирования достаточно крупных территорий, включающих в себя части дорог между перекрестками и пешеходные переходы через эти участки. Пример работы симуляции можно увидеть на рис. 1, красным цветом обозначены вытоптанные в симуляции тропы.

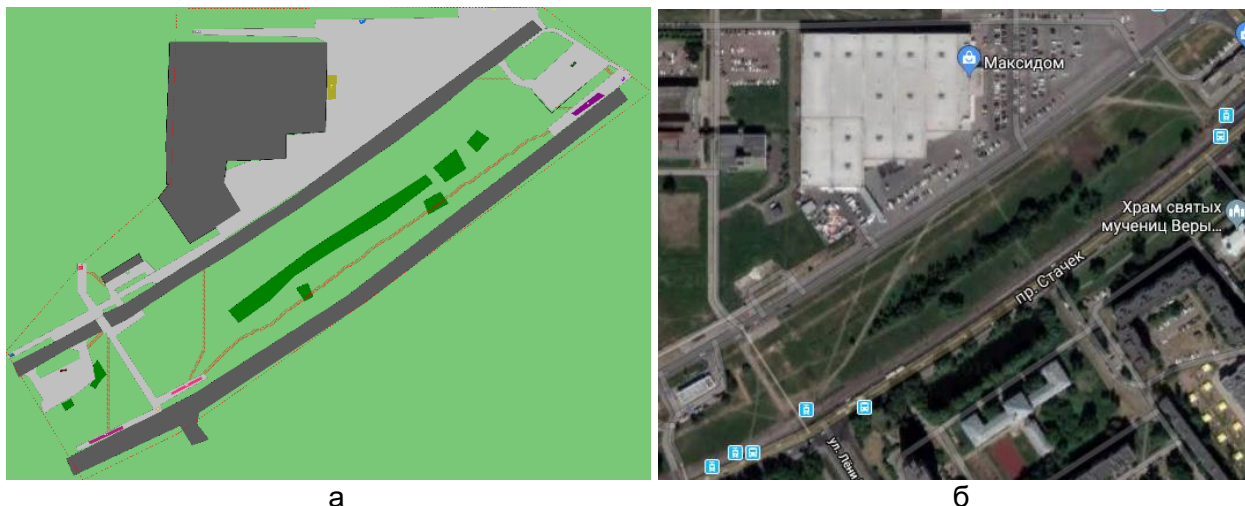


Рис. 1. Изображение вытоптанных троп, полученное симуляцией (а); спутниковый снимок территории, используемой для симуляции (б)

Идея метода, описываемого в данной работе, состояла в развитии данной симуляции с помощью программных расширений, позволяющих использовать полученные симуляцией данные для вывода пересечения пешеходами-агентами симуляции дороги в непопавшем месте, и обозначения на основе данного вывода рекомендуемых пешеходных переходов, позволяющих учесть потребности пешеходов.

Для реализации данного метода были созданы два программных расширения. Первое расширение заносит полученные симуляцией данные в отдельно создаваемую схему в базе данных. Внутри схемы создаются две таблицы. Первая таблица содержит информацию о координатах всех вершин навигационного графа, выстраиваемого симуляцией, и количество пешеходов, прошедших по каждой вершине. Вторая таблица содержит информацию о расположении и типе каждого элемента размеченной карты.

Второе расширение позволяет на основе запроса к базе данных, передаваемого в качестве отдельного параметра, выводить на изображение места, соответствующие условиям запроса. В данном методе используется вывод точек, в которых побывал хотя бы один пешеход, лежащих внутри элементов, соответствующих автомобильной дороге. В качестве результата получается изображение всех пересечений дороги вне пешеходного перехода. Изображение таких пересечений можно увидеть на рис. 2, а, пересечения обозначены красным цветом.

Следует отдельно упомянуть о возможности использования этих программных расширений для вывода и другой информации о результатах работы симуляции, что открывает потенциал к дальнейшему развитию программы.

Следующим шагом является объединение различных пересечений дороги в кластеры. В один кластер объединяются пересечения, находящиеся на расстоянии не более 50 м друг от друга и не разделенные существующим пешеходным переходом. Максимальным размером кластера устанавливается 200 м, при превышении данного размера кластер разбивается на два равных кластера. Кластер включает в себя минимум два пересечения, одиночные пересечения, находящиеся на расстоянии более 50 м либо отделенные существующими пешеходными переходами от других пересечений, считаются выбросами. Для каждого кластера определяется центроид – среднее арифметическое от положений пересечений дороги внутри кластера.

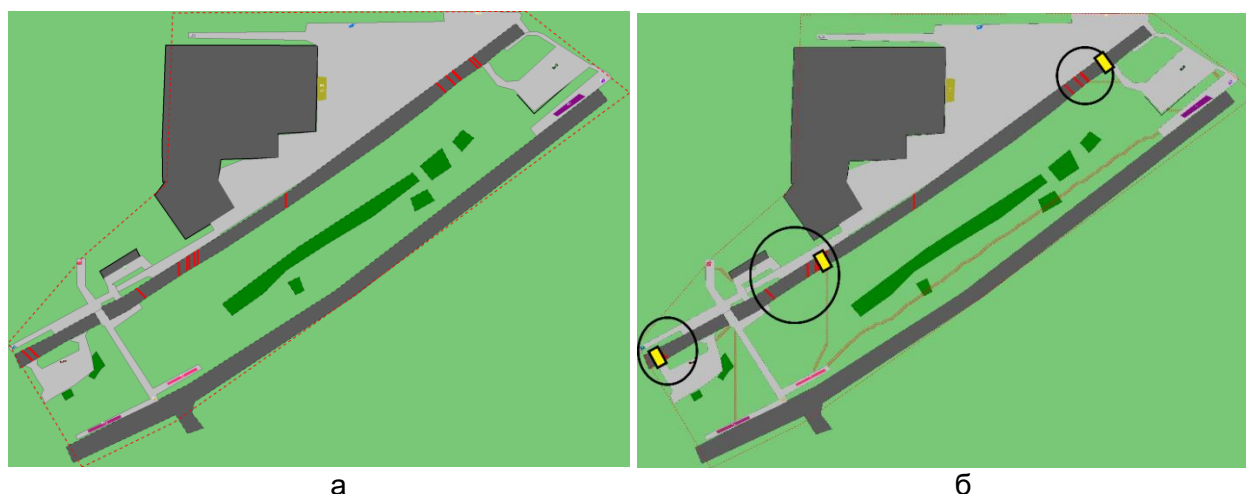


Рис. 2. Изображение пересечений дороги в процессе симуляции (а); обозначение кластеров и рекомендуемых переходов (б)

Далее происходит определение положений рекомендуемых пешеходных переходов внутри каждого кластера. Переход ставится как можно ближе к центроиду кластера. При наличии внутри кластера примыкающих к дороге пешеходных дорожек или вытоптанных троп переход ставится в местах таких примыканий, при этом пешеходные дорожки имеют больший вес по сравнению с вытоптанymi тропами. Обязательным условием для постановки пешеходного перехода является минимальное расстояние от него до существующего перехода в 50 м. При нарушении этого условия пешеходный переход внутри кластера не ставится.

Пример кластеризации и обозначения рекомендуемых пешеходных переходов можно увидеть на рис. 2, б. Кластеры обозначены эллипсами, рекомендуемые переходы обозначены желтыми прямоугольниками.

Для оценки работы данного метода территория, представленная на изображениях, была размечена заново с учетом обозначенных переходов. Результат повторной симуляции с учетом этих переходов можно увидеть на рис. 3. Как видно на изображении, наличие дополнительных переходов полностью исключило пересечения дороги в неполюженном месте (отсутствие красных линий, пересекающих дорогу).

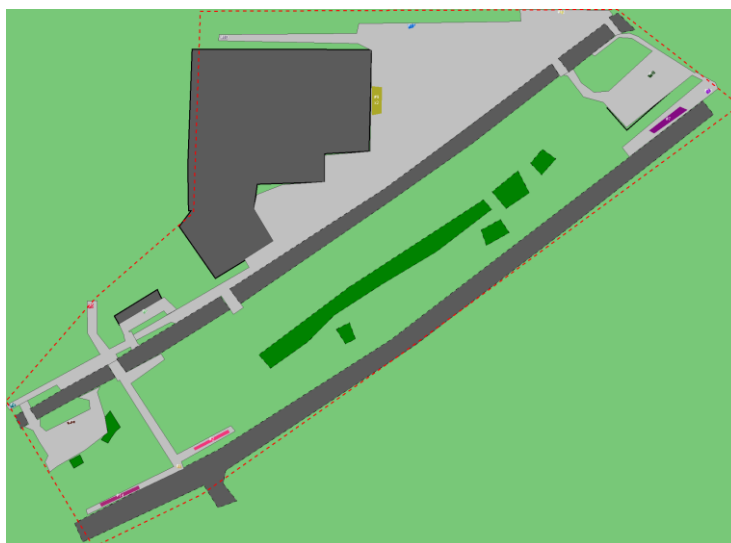


Рис. 3. Результат повторной симуляции

В результате данного исследования разработаны основы метода, позволяющего с помощью данных симуляции выстраивать оптимальную расстановку пешеходных переходов, учитывающих потребности пешеходов и их устоявшиеся маршруты движения. В качестве направлений дальнейшей работы планируется уточнение и формализация алгоритма

определения положения рекомендуемых пешеходных переходов и автоматизация данного алгоритма в виде программных расширений к симуляции. Также планируется проведение натурных исследований для оценки работы данного метода и последующего уточнения алгоритма. Система обозначения рекомендуемых пешеходных переходов и пересечений пешеходами дороги позволит упростить проектирование оптимальной пешеходной инфраструктуры и избежать появления опасных для пешеходов маршрутов.

Литература

1. Почему нарушаем? Переход дороги в неположенном месте [Электронный ресурс]. – Режим доступа: <https://primamedia.ru/news/706894/> (дата обращения: 07.03.2019).
2. ГОСТ 32944-2014. Дороги автомобильные общего пользования. Пешеходные переходы. Классификация. Общие требования. – Введен 08.09.2016. – М.: Стандартинформ, 2016. – 14 с.
3. Методические рекомендации по разработке и реализации мероприятий по организации дорожного движения. Развитие пешеходных пространств поселений, городских округов в Российской Федерации. – М.: Министерство транспорта Российской Федерации, 2018. – 7 с.
4. Broek N.V. The When, Where and How of Mid-Block Crosswalks [Электронный ресурс]. – Режим доступа: <http://www2.ku.edu/~kutc/pdf/LTAPFS11-Mid-Block.pdf> (дата обращения: 06.01.2019).
5. Kudinov S., Smirnov E., Malyshev G., Khodnenko I. Planning Optimal Path Networks Using Dynamic Behavioral Modeling // Springer Nature. – 2018. – P. 129–141.

**Елховская Любовь Олеговна**

Год рождения: 1996

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С4113сНаправление подготовки: 01.04.02 – Вычислительная биомедицина

e-mail: lyubov.elkhovskaya@gmail.com

**Ковальчук Сергей Валерьевич**

Год рождения: 1984

Университет ИТМО, институт дизайна и урбанистики,
к.т.н.

e-mail: sergey.v.kovalchuk@gmail.com

УДК 004.942, 004.896**ПЕРСОНАЛИЗИРОВАННОЕ ПРЕДСКАЗАТЕЛЬНОЕ МОДЕЛИРОВАНИЕ
ЖИЗНЕННЫХ ПОКАЗАТЕЛЕЙ ПАЦИЕНТОВ С ХРОНИЧЕСКИМИ
ЗАБОЛЕВАНИЯМИ В УСЛОВИЯХ ОГРАНИЧЕННОСТИ НАБЛЮДЕНИЙ****Елховская Л.О.****Научный руководитель – к.т.н. Ковальчук С.В.**

Работа выполнена в рамках темы НИР № 618267 «Интеллектуальные технологии гибридного предсказательного моделирования в задачах Р4-медицины».

Объектом исследования данной работы являлись технологии гибридного предсказательного моделирования состояния пациента и клинического течения хронических заболеваний с целью повышения эффективности медицинских учреждений в лечении и контроле заболеваний и обеспечения оптимального соотношения между индивидуальным подходом к пациенту (персонализацией) и обоснованностью предлагаемых решений, и как следствие, улучшения оказания медицинской помощи в целом. В работе рассмотрены подходы к решению проблем в области медицины и здравоохранения в рамках концепции Р4-медицины, и приведены результаты промежуточных экспериментов.

Ключевые слова: Р4-медицина, персонифицированные рекомендации, предсказательное моделирование, интеллектуальный анализ данных, машинное обучение.

На сегодняшний день одной из важных проблем в области медицины и здравоохранения является поддержка и оценка состояния здоровья пациентов с неинфекционными хроническими заболеваниями (ХНИЗ), являющимися одними из ведущих причин смертности в мире по данным Всемирной организации здравоохранения (ВОЗ) [1]. В то время как популяция пациентов с ХНИЗ растет с каждым годом, чему способствует урбанизация, старение населения и распространение нездорового образа жизни, осуществление достаточного медицинского контроля становится затруднительным в условиях ограниченности наблюдений: амбулаторные пациенты редко посещают специализированные медицинские центры и не всегда соблюдают рекомендации врача. Снизить уровень неопределенности в условиях, в которых медицинский специалист должен принимать решения о дальнейшем курсе лечения, а также помочь пациентам самостоятельно следить за их здоровьем можно посредством анализа данных и предсказательного моделирования

жизненных процессов, используя данные, полученные дистанционно с носимых устройств (телемониторинг), и сведения об амбулаторных приемах.

В работе рассматриваются подходы к решению обозначенной проблемы в рамках концепции P4-медицины (*personalized, predictive, preventive, participatory*). Основной целью работы являлась разработка комплекса технологий, обеспечивающих выработку персонифицированных рекомендаций для пациентов с хроническими заболеваниями по результатам непрерывного или периодического мониторинга ключевых показателей жизнедеятельности, с совместным использованием средств интеллектуальной обработки больших данных и предсказательного моделирования. На текущем этапе работы первоочередной задачей являлся анализ состояния исследований по разным направлениям в области P4-медицины с целью выявления подходов и методов для решения задач непрерывной диагностики и оценки рисков развития заболевания.

В рамках первого этапа работы были выявлены средства и методы для решения проблем в разных направлениях задач P4-медицины:

1. для диагностики и прогнозирования риска развития диабета [2–4] обычно используются такие модели, как логистическая регрессия, наивный байесовский классификатор, решающие деревья и случайный лес;
2. логистическая регрессия и регрессионная модель пропорциональных рисков Кокса широко применяются в прогнозном моделировании времени выживания и повторной госпитализации [5–7] пациентов с сердечной недостаточностью;
3. многие научно-исследовательские работы посвящены прогнозному моделированию неблагоприятных явлений у пациентов с сердечной недостаточностью и артериальной гипертензией [8–10] на основе мобильных клинических данных, где также преобладают регрессионные модели наравне с наивной байесовской моделью и случайным лесом. Стоит отметить, что, несмотря на доминирование случайного леса в показателях эффективности и обширности применения в различных задачах, концепции, выработанные данным методом, могут быть не интерпретируемыми с точки зрения медицинского эксперта.

В контексте последнего пункта в работе приведены результаты эксперимента применения предложенного подхода к решению задачи прогнозирования состояния контролируемости артериальной гипертензии (АГ): на основе одной из проанализированных научных статей [11] был проведен эксперимент на реальных данных программы дистанционного мониторинга, в которой участвовали более 7 тысяч пациентов, страдающих повышенным артериальным давлением. Показатели артериального давления собирались телемедицинской системой в течение 2018 года и всего составили более 3 миллионов записей о домашних измерениях давления. Следует отметить, что были доступны данные только по демографическим (возраст, пол, место жительства) и основным жизненным показателям (вес, систолическое и диастолическое давление, частота сердечных сокращений).

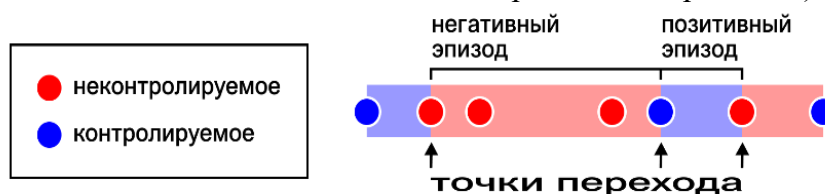


Рис. 1. Оценка состояния контроля артериальной гипертензии для одного пациента с течением времени в виде последовательности событий

Была разработана модель логистической регрессии для предсказания статуса состояния пациента (контролируемое/неконтролируемое, рис. 1), а также новое признаковое пространство согласно методу, предложенному авторами рассматриваемой статьи [11]. В обучающую и тестовую выборку вошли только те пациенты, чьи временные рамки программы полностью покрывали выбранное окно наблюдения – 2018 год, а также те, кто с постоянной периодичностью делал замеры давления (в каждом месяце есть данные и на протяжении не

менее трети года) с учетом того, что один сеанс измерения давления утром и вечером состоял из 1–4 контрольных замеров. Таким образом, полученная выборка состояла из 223 пациентов с 261 015 записями (рис. 2).

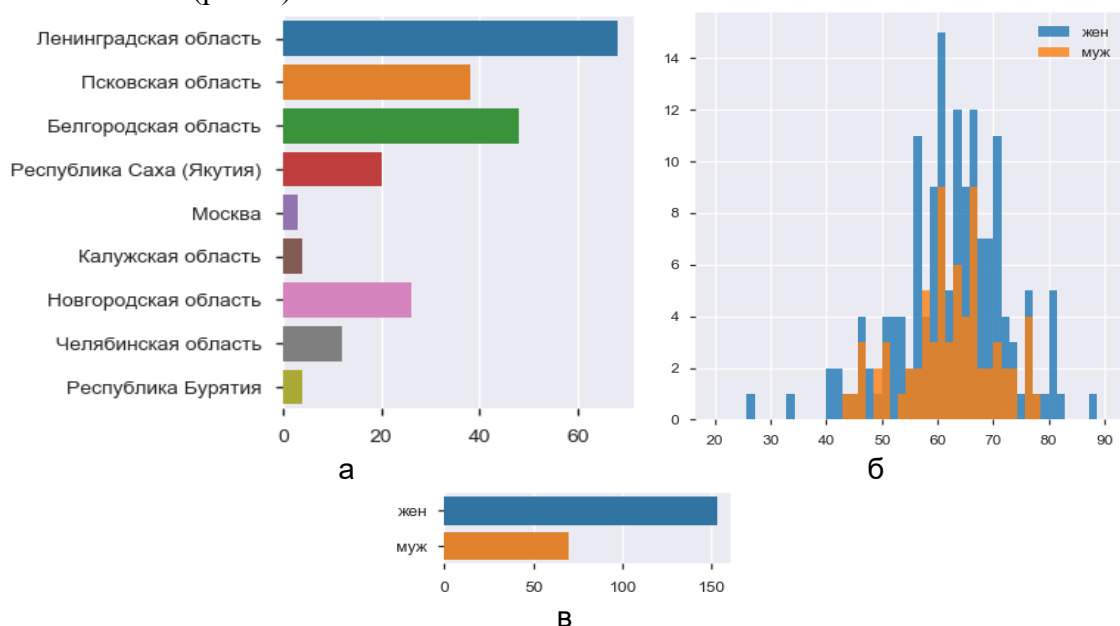


Рис. 2. Некоторые статистические данные по полученной таблице: число пациентов по регионам (а); возрастное распределение пациентов по полу (б); число пациентов мужского и женского пола (в)

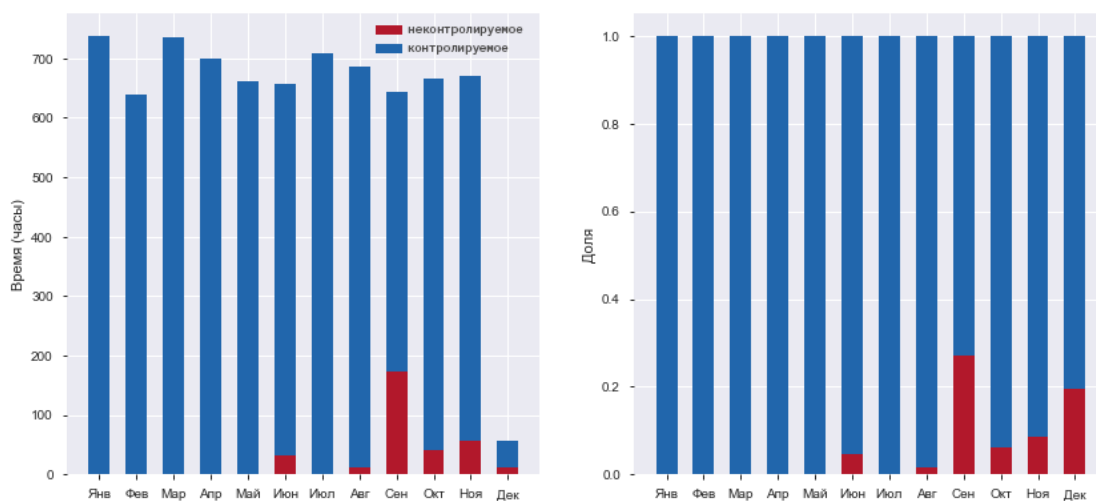


Рис. 3. Общая продолжительность одного из двух состояний (суммарная длительность соответствующих эпизодов) конкретного пациента за 2018 год

Затем были сделаны оценки состояния пациента: поскольку измерения проводились в домашних условиях, давления выше 135/85 мм рт. ст. определялись как индикатор проявления АГ. Последовательность оценок показателей артериального давления пациента объединялись в негативные и положительные эпизоды в плане контролируемости заболевания соответственно (рис. 1 или рис. 3). На их основе были сконструированы новые предикторные переменные, такие как число точек перехода C_p из одного состояния пациента в другое, взаимоисключающие индикаторные переменные, соответствующие нахождению пациента p в основном контролируемом I_p , в основном неконтролируемом O_p или в смешанном состоянии M_p . Таким образом, $I_p + O_p + M_p = 1$, и I_p означает, что суммарная длительность периодов нахождения пациента в контролируемом состоянии превышает суммарную длительность неконтролируемого в 1,5 раза, аналогично для O_p (M_p в случае неопределенности).

Таблица. Основные показатели, используемые в рамках исследования

Концепция признаков	Название переменной	Тип переменной	Тип агрегации
Демографические характеристики	Возраст	Числовой (целый)	Последнее значение
	Пол	Бинарный	Последнее значение
	Регион	Категориальный	Последнее значение
Жизненно важные показатели	Систолическое АД	Числовой (целый)	Среднее
	Диастолическое АД	Числовой (целый)	Среднее
	Частота сердечных сокращений	Числовой (целый)	Среднее
Контроль состояния АГ	C_p	Числовой (целый)	Сумма
	I_p	Бинарный	Последнее значение
	O_p	Бинарный	Последнее значение
	M_p	Бинарный	Последнее значение

В итоге модель включала переменные, представленные в таблице. Обработка пропущенных значений в полученном наборе данных не требовалась в виду их отсутствия. Выборка разбивалась на обучающую и тестовую в соотношении 70% и 30% соответственно, затем были применены методы конструирования признаков (масштабирование числовых и кодирование категориальных). Оптимизация гиперпараметров модели производилась с помощью прохода по сетке. Предложенная модель логистической регрессии показала лучшие результаты по сравнению с той же моделью без нововведенных признаков: значение F-меры повысилось с 0,67 до 0,82. Тест хи-квадрат выявил высокую корреляционную зависимость предложенных новых переменных с целевой, в особенности число транзитных переходов из одного состояния пациента в другое. Концепция основных жизненных показателей (значения систолического и диастолического давления) и новая концепция контроля состояния АГ внесли наибольший вклад в предсказательную способность модели.

Дальнейшая научная работа будет сосредоточена на разработке подходов и методов по выработке персонализированных рекомендаций для пациентов с ХНИЗ с учетом данных непрерывного или периодического мониторинга показателей жизнедеятельности. Причина такого решения заключается в том, что исследование касается случая амбулаторного наблюдения. Более того, измерения жизненно важных показателей, собираемые дистанционно, обеспечивают персонифицированный подход к предсказательному моделированию, что делает его более точным и подходящим для раннего выявления негативных клинических случаев. Текущие результаты уместны для использования в будущих исследованиях, которые могут не только предоставить врачам инструмент, повышающий качество оказания медицинской помощи, но и помочь пациентам лучше справляться с заболеваниями.

Литература

1. WHO. World Health Organisation, 2018 [Электронный ресурс]. – Режим доступа: <https://apps.who.int/iris/bitstream/handle/10665/272596/9789241565585-eng.pdf?ua=1> (дата обращения: 06.01.2019).
2. Arjun C., Anto S. Diagnosis of Diabetes Using Support Vector Machine and Ensemble Learning Approach // IJEAS. – 2015. – V. 2. – № 11. – P. 68–72.
3. Meng X., Huang Y., Rao D., Zhang Q., Liu Q. Comparison of three data mining models for predicting diabetes or prediabetes by risk factors // Kaohsiung Journal of Medical Sciences. – 2013. – V. 29. – № 2. – P. 93–99.
4. Kandhasamy J.P., Balamurali S. Performance Analysis of Classifier Models to Predict Diabetes Mellitus // Procedia Computer Science. – 2015. – V. 47. – P. 45–51.

5. Eapen Z.J., Liang L., Fonarow G.C., Heidenreich P.A., Curtis L.H., Peterson E.D., Hernandez A.F. Validated, Electronic Health Record Deployable Prediction Models for Assessing Patient Risk of 30-Day Rehospitalization and Mortality in Older Heart Failure Patients // JACC: Heart Failure. – 2013. – V. 1. – № 3. – P. 245–251.
6. Lin Y., Chen H., Brown R.A., Li S., Yang H. Predictive Analytics for Chronic Care: A Time-to-Event Modeling Framework Using Electronic Health Records // SSRN Electronic Journal. – 2014. – P. 1–54.
7. Njagi E.N., Rizopoulos D., Molenberghs G., Dendale P., Willekens K. A joint survival-longitudinal modelling approach for the dynamic prediction of rehospitalization in telemonitored chronic heart failure patients // Statistical Modelling. – 2013. – V. 13. – № 3. – P. 179–198.
8. Rothman M.J., Rothman S.I., Beals J. IV Development and validation of a continuous measure of patient condition using the Electronic Medical Record // Journal of biomedical informatics. – 2013. – V. 46. – № 5. – P. 837–848.
9. Larburu N., Artetxe A., Escolar V., Lozano A., Kerexeta J. Artificial Intelligence to Prevent Mobile Heart Failure Patients Decompensation in Real Time: Monitoring-Based Predictive Model // Hindawi. – 2018. – V. 2018. – 11 p.
10. Henriques J., Carvalho P., Paredes S., Rocha T., Habetha J., Antunes M., Morais J. Prediction of heart failure decompensation events by trend analysis of telemonitoring data // IEEE Journal of Biomedical and Health Informatics. – 2014. – V. 19. – № 5. – P. 1757–1769.
11. Sun J., Mcnaughton C.D., Zhang P., Perer A., Gkoulalas-divanis A., Denny J.C., Kirby J., Lasko T., Saip A., Malin B.A. Predicting changes in hypertension control using electronic health records from a chronic disease management program // JAMIA. – 2014. – V. 21. – № 2. – P. 337–344.



Киселева Валерия Олеговна

Год рождения: 1997

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С41551

Направление подготовки: 09.04.03 – Прикладная информатика

e-mail: kiseleva.v97@mail.ru



Видясова Людмила Александровна

Университет ИТМО, Центр технологий электронного правительства,
к.социол.н.

e-mail: bershadskaya.lyudmila@gmail.com

УДК 004.056.5

ОПРЕДЕЛЕНИЕ ДОВЕРИЯ В КОНТЕКСТЕ ИСПОЛЬЗОВАНИЯ ИНФОРМАЦИОННЫХ ТЕХНОЛОГИЙ: АНАЛИТИЧЕСКИЙ ОБЗОР

Киселева В.О., Видясова Л.А.

Научный руководитель – к.социол.н. Видясова Л.А.

Работа выполнена в рамках проекта «Исследование киберсоциального доверия в контексте использования и отказа от информационных технологий» при поддержке Российского фонда фундаментальных исследований (грант РФФИ № 18-311-20001).

В работе представлен аналитический обзор нормативно-правовой базы, научных статей и статистических данных с целью выявления факторов доверия в контексте использования информационных технологий.

Ключевые слова: доверие, информационные технологии, информационная безопасность, государственные услуги, электронное правительство.

Несмотря на масштабную распространенность информационных технологий, существует категория граждан, которая отказывается от их использования. Согласно данным Министерства связи и массовых коммуникаций Российской Федерации (РФ), доля граждан, использующих механизм получения государственных и муниципальных услуг в электронной форме, составляет 51,3% [1]. А согласно рейтингу [2] по регионам доля граждан в субъектах Российской Федерации, зарегистрированных в ЕСИА по состоянию на 22 января 2018 г с учетом населения старше 14 лет, г. Москва находится на 74 месте с показателем в 43,3%, г. Санкт-Петербург на 65 месте с показателем 47,4%.

В работе представлены результаты исследования, целью которого являлось выявление факторов доверия в контексте использования информационных технологий с помощью аналитического обзора нормативно-правовой базы, научных статей и статистических данных.

В ходе анализа нормативно-правовой базы, организующей деятельность по разработке и внедрению информационных технологий в РФ, было выявлено, что с обозначенными рисками реализации программы не проводят работу по предотвращению угроз. Такими рисками являются [3]:

1. пассивное сопротивление использованию органами государственной власти инфраструктуры электронного правительства и распространению современных информационных технологий;

2. пассивное сопротивление отдельных граждан и общественных организаций проведению мероприятий программы по созданию информационных баз, реестров, классификаторов и единого идентификатора граждан по этическим, моральным, культурным и религиозным причинам.

С перечисленными выше пунктами должна проводиться работа по направлению вовлечения граждан и органов в информационное общество. Программа отвечает на вопросы о техническом доступе информационных технологий, ориентируясь на распространение информации, но не учитывает причины пассивного сопротивления, в том числе возможности различных социальных групп.

Во время работы с данными рисками необходимо получить обратную связь с целью выявления признаков, отвечающих за доверия пользователя, таким образом, чтобы на этапе разработке учесть данные особенности.

По итогам анализа научных статей были выявлены факторы доверия в контексте информационных технологий. Были изучены статьи, ключевыми словами которых было «доверие» и «информационные технологии». В отобранных научных статьях авторами проводился анализ информационных технологий различных отраслей, а также в некоторых из них дана оценка полностью сфере информационных технологий в Российской Федерации.

В таблице отображен результат анализа научных статей, исходя из анализа, был выявлен объект – информационная технология, которую описывал автор, а также выделены возможные меры предотвращения.

Таблица. Факторы, влияющие на доверие пользователя к информационным технологиям

№	Фактор и автор статьи	Информационные технологии	Категория	Меры предотвращения
1	Неактуальность информации (Пономарев С.В., 2014)	Электронное правительство	Контроль информации	Постоянное обновление и актуализация информации
2	Конфиденциальность информации (Овчинников С.А. Гришин С.Е., 2012)	Социальные сети	Контроль информации	Разработка закона об авторском праве с учетом распространения текстов в сети Интернет
3	Несоответствие морально-этическим ценностям (Дедюлина М.А., 2016)	Социальные сети	Контроль информации	Фильтрация контента в сети Интернет с учетом потребностей пользователя
4	Недоверие к оператору информационной системы (Астахова Л.В., 2015)	Социотехнические системы, используемые на предприятиях	Кадровое обеспечение	Дополнение оценочных уровней доверия к безопасности информационных системы компонентами безопасности кадров
5	Киберпреступления (Карпова Д.Н., 2014)	Социальные сети	Цифровая грамотность	Повышение уровня компетенции пользователей
6	Информационная безопасность (Шапиро Л.А., 2018)	Информационные технологии, используемые на предприятиях	Цифровая грамотность	Распространение информации о существующих способах информационной безопасности
7	Недоверие к пользователю	Электронное правительство	Субъект-субъект	Аутентификация пользователей

№	Фактор и автор статьи	Информационные технологии	Категория	Меры предотвращения
	информационной системы (Дураковский А.П., 2015)			
8	Низкий уровень цифровизации (Digital Planet, 2017)	Информационные технологии	Внешняя среда	Повышение уровня цифровизации, что в дальнейшем приведет к инвестициям в экономику
9	Отсутствие «человечности» в информационных технологиях (John F. Tripp., 2016)	Информационные технологии	Субъект-система	Внедрение эффекта социального присутствия в систему

Отказ от использования информационных технологий связан с понятием доверия, и если причина отказа – недоверие к технологии, то на этапе разработки и внедрения на доверие пользователя можно повлиять.

Анализ статистических данных. Для того чтобы оценить степень отказа от использования информационных систем, следует оценить масштаб использования сети Интернет, поскольку Интернет является основным ресурсом, обеспечивающим возможность доступа к множествам информационных систем. Оценка степени проникновения сети Интернет в России осуществлена с помощью статистических данных Фонда общественного мнения [4].

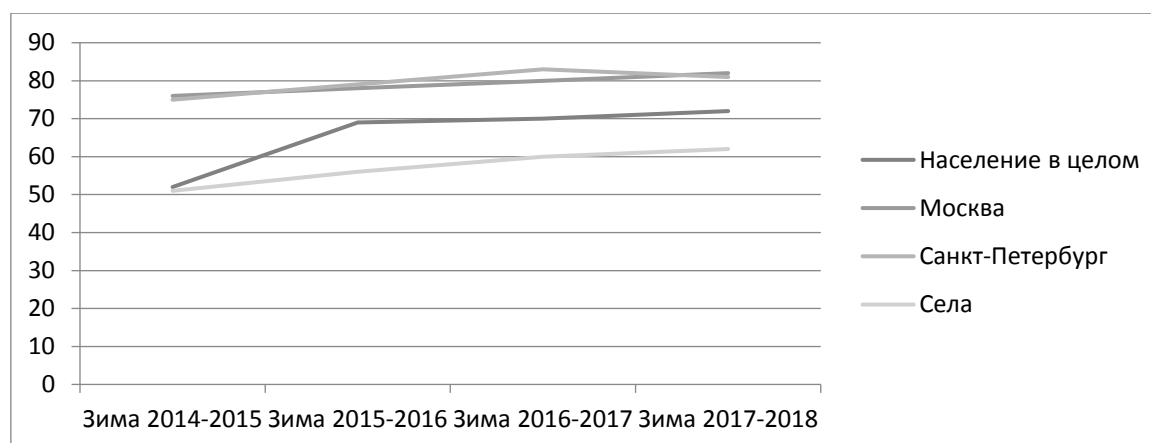


Рис. 1. График динамики проникновения Интернета по населению в целом, в Москве, в Санкт-Петербурге и в селах

Согласно информации, отображенной на рис. 1, можно определить положительную динамику и рост проникновения Интернета как в населении в целом, так и в селах, но разница между городами федерального значения и селами составляет 20%.

Далее был проведен анализ использования информационных технологий для получения государственных и муниципальных услуг, а именно информация о доли населения, использовавшего сеть Интернет для получения государственных и муниципальных услуг, по типам поселения и полу (от общей численности населения, получившего государственные и муниципальные услуги, соответствующих типов поселения) за 2014–2016 годы [5].

На основе данных была построена диаграмма (рис. 2).

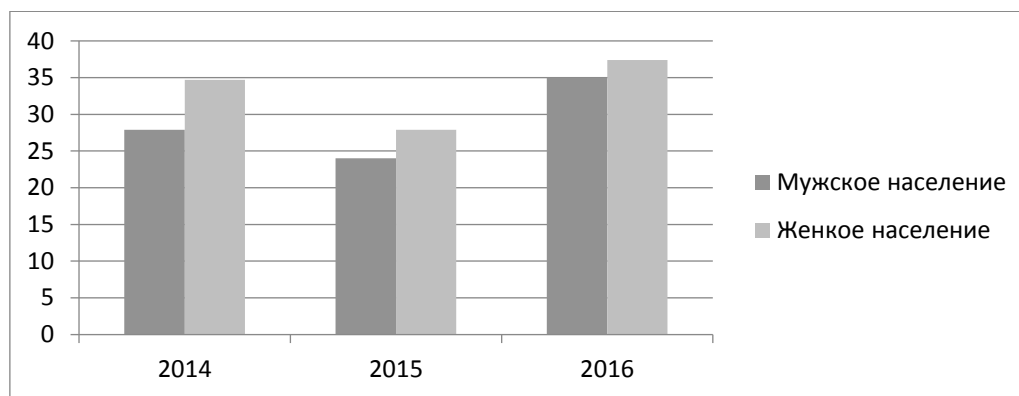


Рис. 2. Диаграмма доли населения, использовавшего сеть Интернет для получения государственных и муниципальных услуг

На рис. 2 можно отметить, что в 2015 году показатели упали по сравнению с 2014, затем снова восстановили показатели в 2016, но все еще с отставанием. За 2015 и 2016 год женское население незначительно опережает мужское по использованию сети Интернет для получения государственных и муниципальных услуг. Согласно данным чуть меньше половины населения получили государственные и муниципальные услуги, пользуясь услугами интернет-ресурсов, но больше половины все еще отказываются от использования информационных технологий и предпочитают пользоваться multifunctional centers.

В результате исследования было выявлено, что на данном этапе термин доверия в контексте информационных технологий не проработан, но является актуальным как со стороны внедрения и разработки информационных технологий, так и со стороны вовлеченности в процесс пользователей. Выявленные факторы доверия пока что применимы только к информационным технологиям, по которым проводился анализ в научных статьях, поэтому необходимо определение доверия, которое будет содержать в себе универсальные факторы. Данные факторы будут определять ключевые аспекты безопасности информационных технологий, которые должны быть учтены при разработке и внедрении.

Литература

1. Постановление Правительства РФ от 15.04.2014 № 313 (ред. от 25.09.2018) «Об утверждении государственной программы Российской Федерации «информационное общество (2011–2020 годы)» [Электронный ресурс]. – Режим доступа: <http://www.consultant.ru/cons/cgi/online.cgi?req=doc&base=LAW&n=307801&fld=134&dst=8957,0&rnd=0.21844842765069017#09938618441574505> (дата обращения: 23.12.2018).
2. Доля граждан в субъектах Российской Федерации, зарегистрированных в ЕСИА по состоянию на 22 января 2018 [Электронный ресурс]. – Режим доступа: <http://er35.ru/upload/Приложение%201%20к%20протоколу%20Рейтинг%20ЕСИА.pdf> (дата обращения: 23.12.2018).
3. Распоряжение Правительства РФ от 28.07.2017 № 1632-р «Об утверждении программы «Цифровая экономика Российской Федерации» [Электронный ресурс]. – Режим доступа: http://www.consultant.ru/document/cons_doc_LAW_221756/ (дата обращения: 23.12.2018).
4. Фонд общественного мнения [Электронный ресурс]. – Режим доступа: <https://fom.ru/> (дата обращения: 20.12.2018).
5. Федеральная служба государственной статистики [Электронный ресурс]. – Режим доступа: <http://www.gks.ru> (дата обращения: 15.12.2018).



Петрушина Татьяна Дмитриевна

Год рождения: 1995

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С4255

Направление подготовки: 09.04.03 – Прикладная информатика

e-mail: petrushina2117@gmail.com



Клечиков Александр Владимирович

Год рождения: 1977

Университет ИТМО, институт дизайна и урбанистики,
к.т.н., доцент

e-mail: a.klechikov@corp.ifmo.ru

УДК 004.9

АВТОМАТИЗАЦИЯ КОНТРОЛЬНО-НАДЗОРНОЙ ДЕЯТЕЛЬНОСТИ В ОБЛАСТИ ГОСУДАРСТВЕННОГО ТАРИФНОГО РЕГУЛИРОВАНИЯ В САНКТ-ПЕТЕРБУРГЕ

Петрушина Т.Д.

Научный руководитель – к.т.н. Клечиков А.В.

Работа посвящена разработке унифицированного (типового) бизнес-процесса контрольно-надзорной деятельности в области государственного тарифного регулирования в Санкт-Петербурге, особенностям автоматизации разработанного бизнес-процесса.

Ключевые слова: контрольно-надзорная деятельность, бизнес-процесс, моделирование, автоматизация.

Реформирование системы государственного контроля (надзора) является одним из основных направлений стратегического развития Российской Федерации. Эти направления реализуются под эгидой Совета при Президенте Российской Федерации по стратегическому развитию и приоритетным проектам [1].

На заседании президиума Совета при Президенте РФ по стратегическому развитию и приоритетным проектам 21 декабря 2016 года утвержден паспорт приоритетной программы «Реформа контрольной и надзорной деятельности». Срок ее реализации – до 2025 года. Реформа объединяет 12 контрольно-надзорных органов, включает более 25 видов контроля. Сопровождает деятельность 4 федеральных органа исполнительной власти – Министерство экономического развития, Министерство юстиции, Министерство цифрового развития, связи и массовых коммуникаций и Министерство труда и социальной защиты – и регионов [2].

Программой определено 8 самостоятельных проектов. Предметом рассмотрения данной работы является автоматизация контрольно-надзорной деятельности. Однако проект носит сквозной характер, и, в некоторой степени, касается всех направлений реформы.

Автоматизация бизнес-процессов – это внедрение программного или аппаратного обеспечения, либо их комплекса, который совместно с новыми правилами выполнения типовых процедур обеспечивает качественное выполнение процесса на оптимальном уровне [3].

Одной из целей автоматизации бизнес-процессов является повышение качества исполнения самого процесса. Автоматизированный процесс обладает более стабильными характеристиками, чем процесс, выполняемый вручную сотрудниками организации. Во

многих случаях целями автоматизация бизнес-процессов могут являться повышение производительности, сокращение времени выполнения процесса, уменьшение риска появления ошибок, увеличение точности и стабильности выполняемых операций [4].

Перед тем как приступать к автоматизации бизнес-процессов необходимо провести их предварительный анализ. Объектами исследования в данной работе являлись бизнес-процессы осуществления регионального государственного контроля (надзора) в области государственного регулирования цен (тарифов). Государственную функцию исполняет Комитет по тарифам Санкт-Петербурга.

При исполнении государственной функции осуществляется взаимодействие с Федеральной антимонопольной службой, Федеральной налоговой службой, Федеральной службой государственной статистики, Прокуратурой Санкт-Петербурга.

Результатами исполнения государственной функции являются составленные акты проверок, выдача предостережений о недопустимости нарушений обязательных требований, а также возбуждение производства по делу об административном правонарушении.

Комитет по тарифам Санкт-Петербурга осуществляет 13 видов (функций) регионального контроля (надзора) [5], в том числе за:

- регулируемые государством ценами (тарифами);
- целевым использованием финансовых средств;
- соблюдением стандартов раскрытия информации;
- применением цен на лекарственные препараты, включенные в перечень жизненно необходимых и важнейших лекарственных препаратов (ЖНВЛП) и т.д.

Для каждого вида регионального контроля (надзора) существует распоряжение Комитета по тарифам Санкт-Петербурга об утверждении Административного регламента Комитета по тарифам Санкт-Петербурга по исполнению государственной функции.

Изучив каждый Административный регламент, было выявлено, что существует «стандартный» набор административных процедур (действий), который осуществляется Комитетом вне зависимости от вида регионального контроля (надзора):

- выдача предостережения о недопустимости нарушения обязательных требований;
- осуществление плановых проверок соблюдения объектами контроля обязательных требований;
- осуществление внеплановых проверок соблюдения объектами контроля обязательных требований;
- подготовка и направление межведомственного запроса в Федеральную налоговую службу, Федеральную службу государственной статистики о предоставлении сведений и документов, необходимых Комитету для принятия решения в рамках исполнения государственной функции;
- рассмотрение акта плановой или внеплановой проверки.

Помимо вышеперечисленных основных административных процедур, существуют «дополнительные» административные процедуры. Их специфика зависит от вида государственного контроля. Примерами могут быть следующие процедуры:

- систематическое наблюдение и анализ информации о соблюдении организациями стандартов раскрытия информации при контроле за соблюдением стандартов раскрытия информации;
- прием и регистрация ежеквартальных и ежегодных отчетов организаций о выполнении инвестиционных программ при контроле за реализацией инвестиционных программ субъектов электроэнергетики;
- прием и регистрация данных в формате шаблонов единой информационно-аналитической системы (ЕИАС) о величине оптовых и розничных надбавок к фактическим отпускным ценам, установленным производителями лекарственных препаратов, на лекарственные препараты, включенные в перечень ЖНВЛП и т.д.

Анализ исследования позволяет сделать заключение, что можно построить типовой (унифицированный) бизнес-процесс, который будет отражать всю контрольно-надзорную деятельность Комитета по тарифам Санкт-Петербурга. Бизнес-процесс описан с помощью методологии функционального моделирования IDEF0 (рисунок).

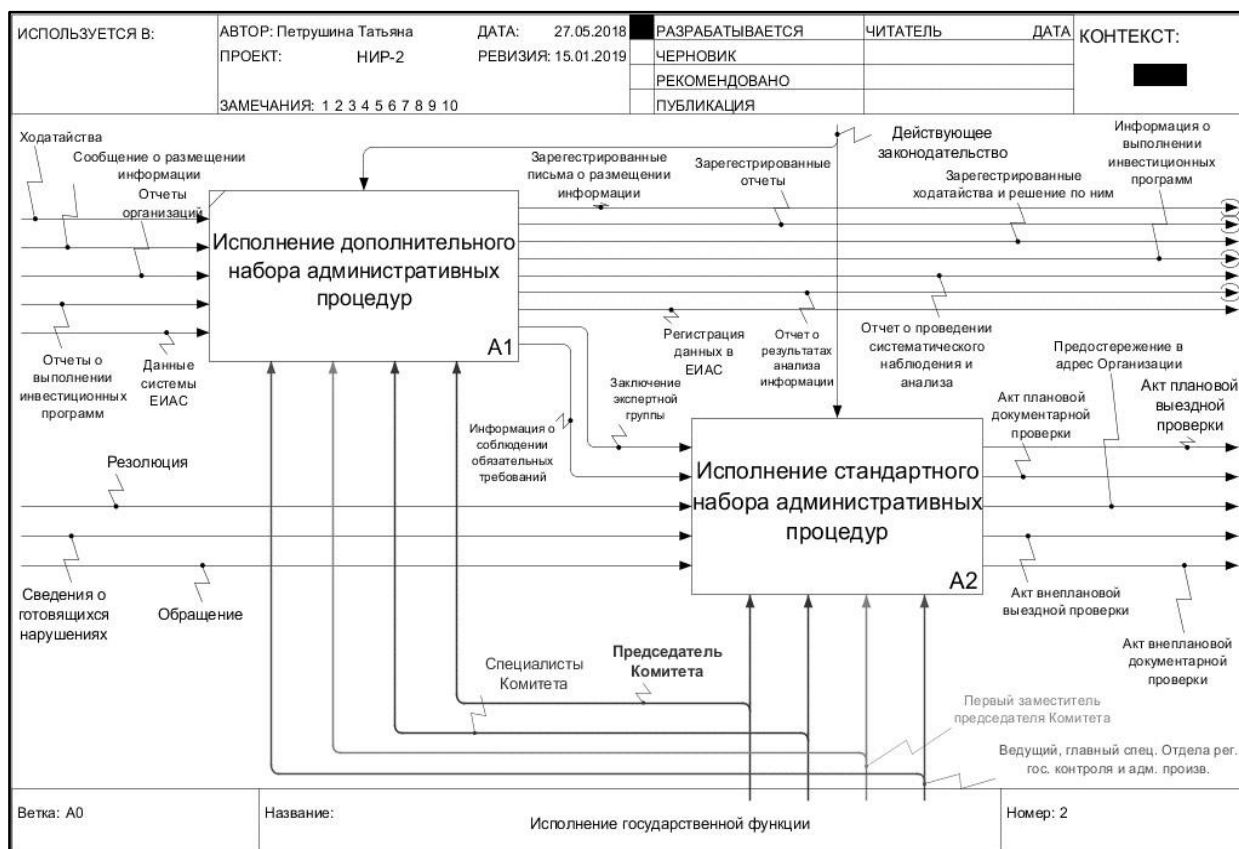


Рисунок. Исполнение административных процедур

Следующим шагом к автоматизации бизнес-процесса является подготовка требований к функциям и задачам, которые будут выполняться информационной системой.

Для основного набора административных регламентов были выявлены основные требования к функциям:

- в результате создания функции «Осуществление плановых (внеплановых) проверок» должно быть осуществлено информационное взаимодействие: с органами Прокуратуры Санкт-Петербурга с целью согласования ежегодного плана проверок; между ответственными лицами контрольно-надзорного органа для подготовки и визирования приказа о проведении проверки; с проверяемыми лицами для направления уведомлений и получения необходимых сведений;
- в результате создания функции «Взаимодействие с заинтересованными сторонами» должно быть осуществлено взаимодействие с Федеральной налоговой службой и Федеральной службой государственной статистики с целью получения сведений из реестров, которые ведутся этими службами.

В ходе выполнения данной работы были проанализированы все виды (функции) государственного контроля (надзора), описанные в административных регламентах Комитета.

Было выявлено, что есть «стандартный» набор административных процедур, которые выполняются вне зависимости от вида регионального контроля (надзора), а есть «дополнительные» административные процедуры, специфика которых зависит от вида государственного контроля.

На основании вышеперечисленных выводов и результатов был разработан унифицированный (типовой) бизнес-процесс контрольно-надзорной деятельности в области государственного тарифного регулирования в Санкт-Петербурге.

Для его дальнейшей автоматизации были выявлены требования к основным функциям, которые будут осуществляться с помощью информационной системы.

Полученные результаты будут использованы для проектирования информационной системы контрольно-надзорной деятельности в сфере государственного тарифного регулирования в Санкт-Петербурге, в частности для проектирования логической и физической моделей данных.

Литература

1. Реформирование контроля и надзора обретает четкие контуры [Электронный ресурс]. – Режим доступа: <http://ac.gov.ru/events/013707.html> (дата обращения 21.01.2019).
2. Реформа контрольно-надзорной деятельности [Электронный ресурс]. – Режим доступа: <http://ac.gov.ru/projects/otherprojects/013591.html> (дата обращения 21.01.2019).
3. Торош О.И., Плучевская Э.В. Моделирование и автоматизация бизнес-процессов как основа эффективного развития предприятия // Управление человеческими ресурсами – основа развития инновационной экономики. – 2011. – № 3. – С. 630–635.
4. Останина Д.А., Мирошниченко М.А. Автоматизация бизнес-процессов как способ повышения эффективности подразделений крупных корпораций // Сборник материалов VII Международной научно-практической конференции «Экономика знаний: стратегические проблемы и решения». – 2015. – С. 299–304.
5. Об утверждении Порядка организации и осуществления регионального государственного контроля (надзора) в области регулирования цен (тарифов) на территории Санкт-Петербурга [Электронный ресурс]. – Режим доступа: <https://www.gov.spb.ru/law/?d&nd=822402533&nh=8> (дата обращения 15.01.2019).



Рубцова Валерия Сергеевна

Год рождения: 1994

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С4255

Направление подготовки: 09.04.03 – Прикладная информатика

e-mail: vsrubtsova@gmail.com



Кононова Ольга Витальевна

Год рождения: 1964

Университет ИТМО, институт дизайна и урбанистики,
к.э.н., доцент

e-mail: kononolg@yandex.ru

УДК 004.9

ОСОБЕННОСТИ ПЕРЕХОДА АВТОМАТИЗИРОВАННОЙ ИНФОРМАЦИОННОЙ СИСТЕМЫ ЭЛЕКТРОННОГО СОЦИАЛЬНОГО РЕГИСТРА НАСЕЛЕНИЯ НА ВЗАИМОДЕЙСТВИЕ ПОСРЕДСТВОМ СИСТЕМЫ МЕЖВЕДОМСТВЕННОГО ЭЛЕКТРОННОГО ВЗАИМОДЕЙСТВИЯ ВЕРСИИ 3

Рубцова В.С.

Научный руководитель – к.э.н. Кононова О.В.

В работе приводятся результаты исследования процесса перехода на взаимодействие автоматизированной информационной системы электронного социального регистра населения с органами власти посредством системы межведомственного электронного взаимодействия версии 3, а также результаты анализа часто встречающихся (типовых) ошибок при реализации данного процесса.
Ключевые слова: межведомственное взаимодействие, СМЭВ, социальная защита населения, ЭСРН.

Исследование было проведено с целью выявления особенностей системы межведомственного электронного взаимодействия (СМЭВ) версии 3 и типовых ошибок системы, возникающих при переходе на взаимодействие посредством СМЭВ 3.

О масштабе задачи могут говорить следующие данные: на федеральном уровне в электронном виде предоставляются 320 услуг, на региональном уровне насчитывается около 13 тысяч позиций, а на муниципальном уже более 25 тысяч.

На данный момент в Российской Федерации функционирует две версии СМЭВ: вторая (СМЭВ 2) и третья (СМЭВ 3). Основное отличие третьей версии системы в наличии единого сервиса СМЭВ. Через такой сервис взаимодействуют все ведомства, формируя вид сведений – информацию, которую они хотят предоставить или получить, и через единый сервис публикуют ее для всех остальных. Таким образом, все участники взаимодействия обращаются к единому электронному сервису: потребители – для отправки запроса и получения на них ответа, а поставщики, в свою очередь, для получения запросов и отправки ответов. Преимущество по сравнению с предыдущими версиями заключается в наличии механизма очереди электронных сообщений, позволяющего участникам межведомственного взаимодействия обмениваться данными в асинхронном режиме с гарантированной доставкой запрошенных данных. С помощью единого электронного сервиса СМЭВ реализуется каждая операция получения/передачи сообщений, как с запросом, так и с ответом.

Асинхронное взаимодействие также является важным отличием двух версий системы. В предыдущей версии СМЭВ 2 обмен данными между участниками межведомственного взаимодействия происходил следующим образом: одно ведомство отправляет запрос другому и сразу же ожидает результат обработки запроса (рис. 1). Такой метод работы приводит к тому, что в дневные часы возникают пиковые нагрузки, а в ночные часы, наоборот, оборудование простаивает.

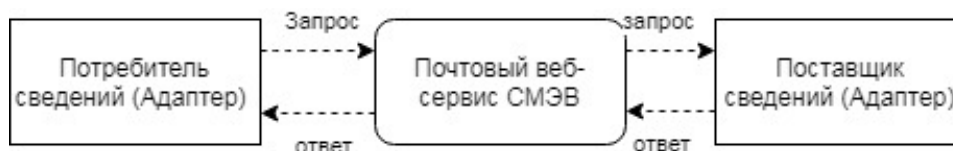


Рис. 1. Схема работы в синхронном режиме

Асинхронное взаимодействие (рис. 2) необходимо в том случае, когда обработка запроса в информационной системе (ИС) поставщика требует большее количество времени, чем период ожидания ответа СМЭВ и ИС потребителя.

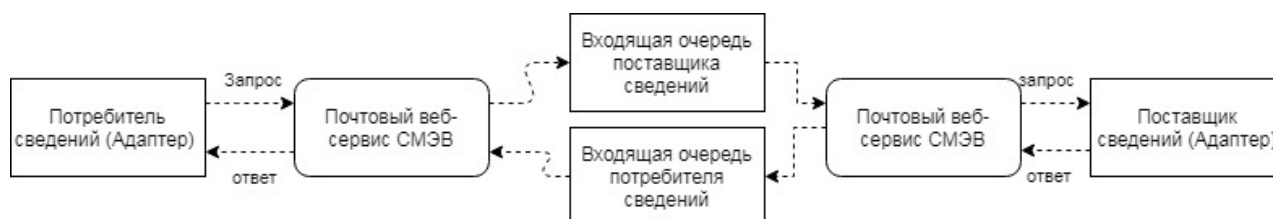


Рис. 2. Схемы работы в асинхронном режиме

Проанализировав нормативно-правовую документацию, а также научные публикации [1–5] можно выделить следующие преимущества, важные для перехода на СМЭВ 3:

1. простота подключения. Для подключения к СМЭВ 3 есть возможность использовать библиотеку типовых решений по взаимодействию со СМЭВ. Также потребитель может реализовать единый адаптер для любых видов сведений, а не для каждого в отдельности;
2. простота разработки видов сведений. Для обмена сведениями используются стандартные протоколы обмена сообщениями, а также типовой набор данных, передаваемый в теле сообщения;
3. пассивная очередь, гарантирующая доставку сообщений. Гарантируется доставка сообщений, благодаря асинхронному типу взаимодействия, а также минимизация непроизводительной нагрузки на систему поставщика сведений.

Автоматизированная ИС «Электронный социальный регистр населения Вологодской области» предназначена для:

- формирования, ведения и использования базы данных, в которой содержится информация о гражданах Российской Федерации, имеющих право на предоставление мер социальной поддержки в соответствии с законодательством;
- автоматизации деятельности органов социальной защиты населения в области предоставления гражданам мер социальной поддержки.

Для межведомственного взаимодействия на платформе разработана подсистема внешнего информационного обмена. Подсистема внешнего информационного обмена предназначена для обмена информацией с другими организациями, предоставляющими государственные услуги в электронном виде с целью подтверждения права на пользование определенными мерами социальной поддержки.

На данный момент в электронном социальном регистре населения Вологодской области (ЭСРН ВО) используется СМЭВ 2. Взаимодействие происходит следующим образом. Модуль электронных сервисов обеспечивает прием заявлений на оказание услуг в электронной форме с Единого портала государственных и муниципальных услуг, их обработку, при необходимости отправку дополнительных электронных запросов в смежные ведомства, а

также отправку промежуточных и конечных результатов обработки заявок на единый портал государственных услуг (ЕПГУ). Модуль электронных сервисов выступает в роли адаптера, размещенного на центральном узле, обрабатывающего входные данные.

После перехода на СМЭВ 3 взаимодействие должно выглядеть следующим образом. Модуль информационного взаимодействия ЕПГУ передает данные о заявлении мер социальной поддержки (МСП), заявителе и о связанных с ним личных делах посредством веб-сервиса. Центральный сервер автоматизированной системы «Электронный социальный регистр населения Вологодской области» (АС ЭСРН ВО) (Адаптер) обеспечивает прием данных заявления и связанных с ним личных дел. Для взаимодействия с внешними организациями посредством веб-сервисов, в том числе и с ЕПГУ, используется распределенная схема взаимодействия «Сервер органа социальной защиты населения – Адаптер – ЕПГУ».

Далее, в рамках взаимодействия АИС ЭСРН, установленная на районных серверах, направляет сообщение-запрос на Адаптер. Адаптер, в свою очередь, пересылает сообщение-запрос через СМЭВ в адрес ИС федеральных органов исполнительной власти и получает сообщение-ответ. Затем адаптер осуществляет маршрутизацию и пересылку сообщений-ответов в соответствующий АИС ЭСРН, установленный на районных серверах, где должна производиться дальнейшая обработка.

Согласно нормативно-правовой документации, были выявлены требования к взаимодействию с сервисами СМЭВ. Модуль, осуществляющий взаимодействие, должен иметь возможность создания и настройки сервисов по методическим рекомендациям СМЭВ 3, в рамках которого будет предоставляться информация по запросам органов власти.

Также должна быть возможность регистрации и обработки данных для заявлений, поступающих из внешних систем. Для работы системы со СМЭВ 3 должны быть реализованы:

- функции, которые обеспечат возможность создания условий и настройки доступа к единому адаптеру;
- функции, обеспечивающие создание запроса, приема и обработки запрашиваемой информации;
- виды сведений для каждой услуги, предусматривающей межведомственное взаимодействие.

Помимо прочего, требуется реализовать инструмент разработки и тестирования ф-сведений СМЭВ 3. По сути, требуется создать свой эмулятор, который бы работал аналогично эмулятору СМЭВ 3, т.е. локальный эмулятор.

Кроме вышеперечисленных требований, могут возникнуть индивидуальные требования к доработке, на основе анализа ошибок при переходе аналогичных систем.

Согласно результатам анализа перехода на СМЭВ 3 аналогичных систем, а также данным мониторинга, были выделены наиболее часто встречающиеся ошибки. В целом, ошибки можно разделить на три категории:

1. ошибки системы;
2. ошибки разработки;
3. ошибки настройки вида сведений (ВС).

Первый вид ошибок, ошибки системы, включают в себя ошибки из-за несоответствия системы требованиям СМЭВ 3. Такая ошибка может возникнуть, например, когда отсутствует необходимый ВС или отсутствует доступ к Единому электронному сервису СМЭВ.

Второй вид ошибок – ошибки разработки. Они появляются, когда разработчик, при создании ВС, допустил ошибку. Примером такой ошибки может служить ошибка вида «Не найден вид сведений», которая появляется при несовпадении наименования корневого элемента на тестируемом клиенте и в эталонном запросе. Или ошибка, допущенная при разработке XSD-схемы, выражающаяся в неправильном указании типа данных.

Также частыми ошибками являются ошибки из категории «Ошибки настройки ВС». Например, при настройке какого-либо из видов сведений, был неправильно указан формат

типа данных. Для поля СНИЛС была указана длина в 9 символов, хотя страховой номер состоит из 11 символов. Такие ошибки, как правило, являются самыми простыми и решаются очень быстро.

На основании проведенного анализа выполнен следующий этап исследования, а именно, разработка рекомендаций для перехода ЭРСН на взаимодействие посредством СМЭВ 3, рассмотрено большее количество типовых ошибок, на основании которых будут разработаны рекомендации для своевременного устранения ошибок такого рода или исключения их появления при переходе на взаимодействие системы со СМЭВ 3.

В ходе проделанного исследования были рассмотрены теоретические положения, касающиеся системы межведомственного электронного взаимодействия, изучены нормативно-правовые акты, регулирующие взаимодействие государственных органов посредством этой системы, а также был проведен анализ статистической информации об использовании СМЭВ органами социальной защиты.

В работе были выявлены проблемы СМЭВ 3, несоответствия системы ЭРСН для работы со СМЭВ третьей версии, была проанализирована нормативная документация, и выявлены требования к интеграции системы со СМЭВ 3, типизированы ошибки при интеграции аналогичных систем социальной защиты населения со СМЭВ 3.

Литература

1. Авакян Н.С. Место системы межведомственного взаимодействия (СМЭВ) в электронном правительстве РФ // Инновационная экономика: информация, аналитика, прогнозы. – 2014. – № 1-2. – С. 3–4.
2. Тюшняков В.Н., Тюшнякова И.А. Проблемы внедрения региональной системы межведомственного электронного взаимодействия // Economic sciences. – 2015. – № 12. – С. 1690–1694.
3. Ломакин А.А. Совершенствование информационного обмена в сфере социальной защиты населения Санкт-Петербурга // Управленческое консультирование. Актуальные проблемы государственного и муниципального управления. – 2014. – № 3. – С. 33–42.
4. Рубцов В.В. О межведомственном взаимодействии в реализации социальной и образовательной инклюзии для социально уязвимых групп населения // Психологическая наука и образование. – 2016. – Т. 21. – № 1. – С. 87–93.
5. Забродин Ю.М., Рубцов В.В. Аprobация, внедрение и применение профессиональных стандартов при решении задач образования и социальной сферы: этапы и перспективы // Межведомственные модели оказания социальных и образовательных услуг и практика апробации и применения профессиональных стандартов работников образования и социальной сферы. – 2016. – С. 13–30.



Смирнова Полина Владиславовна

Год рождения: 1995

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С4255

Направление подготовки: 09.04.03 – Прикладная информатика

e-mail: smirnovapv0125@ya.ru

УДК 004.9

АНАЛИЗ ПОРТАЛА «НАШ ГОРОД МОСКВА»

Смирнова П.В.

Работа выполнена в рамках темы НИР № 418229 «Исследование социальной эффективности портала «Наш город Москва» в Центре технологий электронного правительства Университета ИТМО.

В работе приведены результаты анализа социальной эффективности портала «Наш город Москва». В результате исследования портал показал себя как эффективный инструмент решения городских проблем, позволяющий решать волнующие граждан вопросы в короткие сроки. В то же время проведенный анализ позволил выделить ряд факторов, которые способствуют снижению мотивации граждан к использованию портала и уменьшают потенциальную аудиторию пользователей.

Ключевые слова: социальная эффективность, электронное участие, портал, информационные технологии, городское планирование.

Активно развивающиеся в России порталы электронного участия способствуют повышению уровня участия граждан в решении различных государственных и общественных вопросов при помощи информационно-коммуникационных технологий. Особую популярность среди жителей получили порталы по вопросам городского планирования и благоустройства, однако с научно-практической точки зрения остается нерешенным вопрос оценки социальной эффективности их функционирования. Социальная эффективность – это реализованная на практике степень ожидания интересов и потребностей людей. Возрастает роль анализа порталов электронного участия, поскольку это позволит оценить, насколько правильно выбраны направления этой деятельности и какой результат она приносит [1–5].

Портал «Наш город Москва» был разработан с целью улучшения качества жизни горожан и облика Москвы посредством участия москвичей в жизни столицы. Портал создан для контроля состояния инфраструктуры, проезжей части, транспортных узлов, освещения, уборки, а также благоустройства территорий и других вопросов, чтобы своевременно устранять возникшие сбои.

За исследуемый период система мониторинга собрала 174 504 сообщений, опубликованных на портале. Стоит отметить, что интерес граждан к portalу значительно возрос с момента его создания, о чем свидетельствует рост количества сообщений, публикуемых на портале.

Анализ категорий обращений граждан позволил определить, что наиболее актуальными вопросами для жителей Москвы являются неубранная дворовая территория, ямы/выступы на проезжей части/тротуаре и неисправное освещение в подъездах.

Ознакомиться с решениями по сообщениям граждан на портале позволяет специальный раздел «Результаты», где с помощью фильтрации за каждый месяц можно детально ознакомиться с каждым сообщением, официальным ответом и фотоотчетом по проведенным работам. При этом у автора обращения есть возможность подтвердить статус решения

проблемы, однако, отсутствует возможность оставить комментарий на предоставленный ответ и решение со стороны органов власти.

Анализ статусов обращений граждан продемонстрировал, что на момент проведения исследования на портале было решено менее половины заявленных проблем (0,34%), а количество решенных проблем, подтвержденных гражданами, также составило чуть менее половины от общего количества обращений (0,4%).

В то же время важно отметить, что с момента создания портала наблюдается положительная динамика роста количества решенных проблем. При этом также зафиксирован рост обращений, решение по которым было подтверждено гражданами. Таким образом, несмотря на невысокую долю решенных проблем на текущий момент, портал демонстрирует тенденцию как к повышению активности пользователей портала, так и к росту решенных проблем, волнующих граждан.

На портале имеется раздел рейтингования пользователей, в рамках которого пользователи могут получать очки за различную активность на портале. К примеру, очки могут быть получены за первое опубликованное сообщение, за каждое опубликованное сообщение о проблеме, за каждую признанную органом власти проблему и за первое использование кнопки «поделиться в социальных сетях».

Таким образом, рейтингование позволяет повысить мотивацию граждан к использованию портала, а также способствует дополнительному информированию общественности о деятельности данной площадки. Помимо рейтингования самих пользователей, на портале реализован рейтинг волонтеров, поликлиник, общественного транспорта, а также дворов и домов.

Рассматривая вариативность каналов продвижения портала «Наш город Москва», стоит отметить наличие мобильного приложения, поддерживающего работу с каждым типом операционной системы, а также поддержание активности в социальных сетях (Facebook, Twitter, ВКонтакте, Instagram). Кроме этого, пользователи могут сами повлиять на продвижение портала путем распространения листовок, которые можно оставить на своем доме. Для этого на портале реализован специальный раздел «Расскажи друзьям», где пользователи могут выбрать и распечатать листовки из ряда предложенных вариантов.

Несмотря на простой и понятный интерфейс портала и упрощенную процедуру регистрации, стоит отметить, что на портале не реализован раздел с краткой и емкой информацией по работе портала, схематично объясняющий порядок оформления обращения и его дальнейшего рассмотрения. На портале имеется раздел «Часто задаваемые вопросы», однако размещенной информации может быть недостаточно для пользователей, которые в первый раз столкнулись с работой портала. Кроме того, стоит отметить низкую вариативность каналов поддержки пользователей, так как обратная связь реализована только через форму обратной связи на портале. Среди факторов, влияющих на снижение активности пользователей, можно отметить то, что портал не предоставляет возможности для участия особых категорий граждан. К примеру, на сегодняшний день на портале отсутствует мультязычная версия, а также версия для слабовидящих.

Анализ социальной эффективности портала «Наш город Москва» продемонстрировал, что наиболее высокие оценки портал получил по стадиям «Формирование условий для участия и повестки дня» (75,25%) и «Мониторинг реализации принятых решений» (62,5%) (рис. 1). Максимальные значения были достигнуты по таким индикаторам, как «Нормативная регламентация, легитимность», «Удобство поиска ресурса», «Информирование, распространение информации», «Наличие мотивационных механизмов для привлечения новых пользователей на платформу», а также «Рейтингование пользователей» и «Экономия временных затрат на осуществление взаимодействия с органами власти». Это говорит о том, что портал наиболее эффективно функционирует с точки зрения информирования общественности о предлагаемых государством решениях и создания предпосылок для гражданского участия, а также позволяет наблюдать за ходом реализации решений и решать

заявленные проблемы в более короткие сроки, чем при традиционном взаимодействии граждан с городскими органами власти.

Самые низкие оценки были получены по индикаторам стадии «Интеграция вклада общественности в предложения/идеи» (8,65%). Значительный отрыв от других стадий связан, прежде всего, с тем, что портал не позволяет гражданам обсуждать и комментировать решения, которые были предоставлены со стороны органов власти, а также проанализировать вклад граждан в принятие или реализацию тех или иных решений.

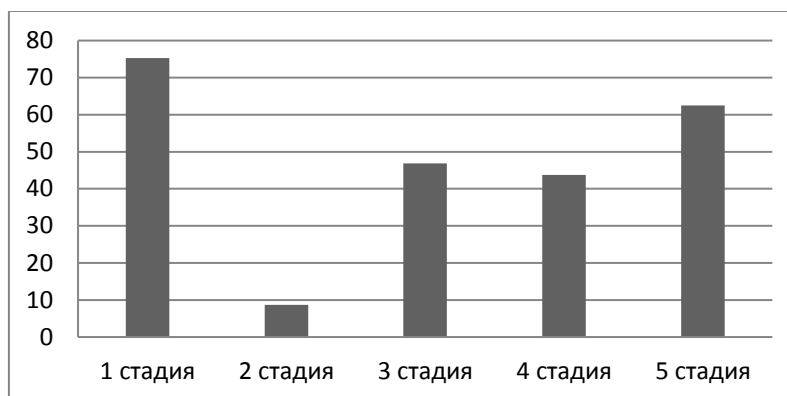


Рис. 1. Оценки портала «Наш город Москва» по стадиям принятия политических решений, %

Анализируя отдельные измерения социальной эффективности, стоит отметить, что замечен достаточно большой отрыв социального измерения в сравнении с оценками, полученными по организационному и технологическим измерениям (рис. 2).

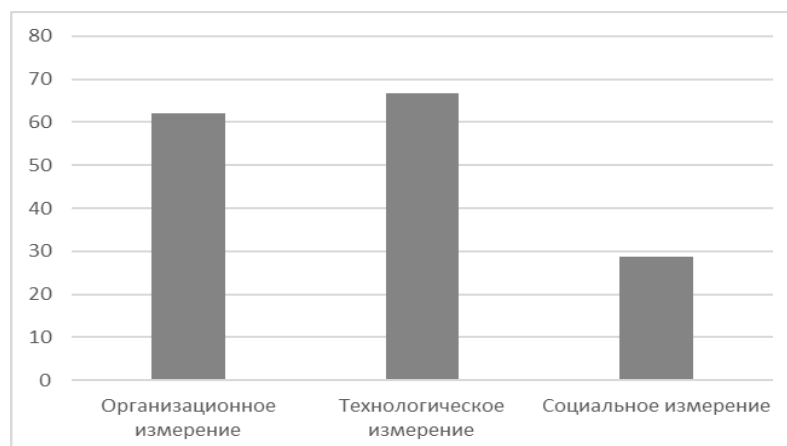


Рис. 2. Оценки портала «Наш город Москва» по измерениям социальной эффективности, %

Наиболее высоко были оценены индикаторы технологического измерения (66,75%), что свидетельствует о потенциальной готовности целевой аудитории к использованию портала.

По организационному измерению портал также получил более половины от максимально возможной оценки (62,14%), что говорит о легитимности деятельности портала и высоком уровне организации цикла принятия решений.

Наименее высоко были оценены индикаторы социального измерения (28,83%). Данный отрыв связан с низкими оценками по таким индикаторам, как «Наличие волонтерского движения, группы поддержки», «Возможность анализа вклада граждан», «Активность сбора голосов/мнений/комментариев», «Наличие дополнительных механизмов сбора мнений граждан», «Возможность принимать участие в выборе средств реализации принятых решений» и «Формирование новых возможностей для общественной активности». В то же время стоит отметить, что максимальные оценки были зафиксированы по индикаторам

«Наличие мотивационных механизмов для привлечения новых пользователей на платформу» и «Экономия временных затрат на осуществление взаимодействия с органами власти».

Таким образом, анализируя социальную эффективность портала «Наш город Москва» по социальному измерению, стоит отметить, что на портале успешно реализована мотивационная система, способная повысить интерес и активность граждан на портале, а также привлечь новых пользователей. Кроме того, портал позволяет гражданам решать волнующие проблемы в короткие сроки и экономить время на взаимодействии с органами власти. Однако при этом на портале не реализованы дополнительные возможности участия и сбора мнений граждан, что может во многом влиять на отсутствие сплоченного сообщества вокруг деятельности портала и не способствует повышению активности пользователей.

Итоговый индекс (45,19%) позволяет отнести портал «Наш город Москва» к среднему уровню социальной эффективности.

Литература

1. Medaglia R. eParticipation research: Moving characterization forward (2006–2011) // Government Information Quarterly. – 2012. – V. 29. – P. 346–360.
2. Голубева А.А., Ишматова Д.Р. Электронная демократия в России: формирование традиции политической осведомленности и участия // Вопросы государственного и муниципального управления. – 2012. – № 4. – С. 50–65.
3. Домнина А.В. Развитие механизмов электронной демократии в современной Российской Федерации // Вестник Нижегородского университета им. Н.И. Лобачевского. – 2014. – № 3-2. – С. 70–73.
4. Кузнецова Е.С., Богданова А.С. Оценка эффективности в сфере проектов государственного и муниципального управления // Вестник Мурманского государственного технического университета. – 2012. – № 1. – С. 195–198.
5. Курочкин А.В. Государственное управление и инновационная политика в условиях сетевого общества: дис. ... д-ра полит. наук. – СПб., 2014. – 295 с.



Смирнова Полина Владиславовна

Год рождения: 1995

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С4255

Направление подготовки: 09.04.03 – Прикладная информатика

e-mail: smirnovapv0125@ya.ru



Кононова Ольга Витальевна

Год рождения: 1964

Университет ИТМО, институт дизайна и урбанистики,
к.э.н., доцент

e-mail: kononolg@yandex.ru

УДК 004.9

ТЕХНОЛОГИИ ГЕЙМИФИКАЦИИ: ОБЗОР НАУЧНЫХ ПУБЛИКАЦИЙ

Смирнова П.В., Кононова О.В.

Научный руководитель – к.э.н. Кононова О.В.

В работе представлены результаты обзора подборки материалов на английском языке за 2014–2018 гг., отобранных по ключевому термину «gamification». Были выбраны 175 источников, отражающих смысл понятия. Данные, полученные в результате обзора, позволили выявить существенные характеристики понятия «геймификация» и аспекты, в которых упоминается данная технология. Анализ проводился при помощи инструментов Voyant Tools и Tropes.

Ключевые слова: контент анализ, геймификация, аналитические системы анализа текста, Tropes, Voyant Tools.

Термин «gamification» был введен в научный оборот в 2010 г. Есть много определений этого термина. Наиболее известным из них является «использование игровых элементов в неигровых условиях» [1].

Исследователи описывают множество сфер применений для геймификации, которые связаны с решением социальных, общественных, культурных и гражданских проблем [2]. Методы геймификации также используются в образовании, здравоохранении, культуре, туризме, электронном правительстве, гражданском электронном участии, сборе городских данных, и других сферах.

Мировое сообщество признает важность геймификации, как научного явления и важность проведения исследований в этой области [3]. Для того чтобы оценить динамику научных интересов к технологии геймификации предлагается обзор свода зарубежных научных публикаций. Анализ терминологического ландшафта научных публикаций позволяет:

- выделить перспективные предложения и практику по использованию геймификации в различных сферах жизни общества;
- сравнить мировые тенденции в использовании технологии геймификации;
- сосредоточить внимание исследователей на ряде научно-значимых результатов;
- выявление траекторий развития перспективных научных направлений.

Для проведения анализа текстов, найденных с использованием поисковых запросов по термин-концепту «геймификация» в базах ScienceDirect, Web of Science и Scopus был выбран

301 источник. Экспериментальная часть исследования включила использование многоязычных сред Voyant-Tools (<https://voyant-tools.org/>) и Tropes High Performance Text Analysis (<http://www.semantic-knowledge.com/tropes.htm>).

Voyant Tools – веб-приложение для анализа текста, поддерживающее множество языков и позволяющее посчитать частотность слов, построить по ним облака тегов, а также линейные графики, отражающие распределение ключевого слова по корпусу текстов.

Tropes High Performance Text Analysis – средство для текстового анализа, разработанное для семантической классификации, извлечения ключевых слов, лингвистического и качественного анализа.

Для сравнения полученных данных за разные годы был произведен анализ. На рисунке показано количество публикаций, которые были отобраны для анализа.

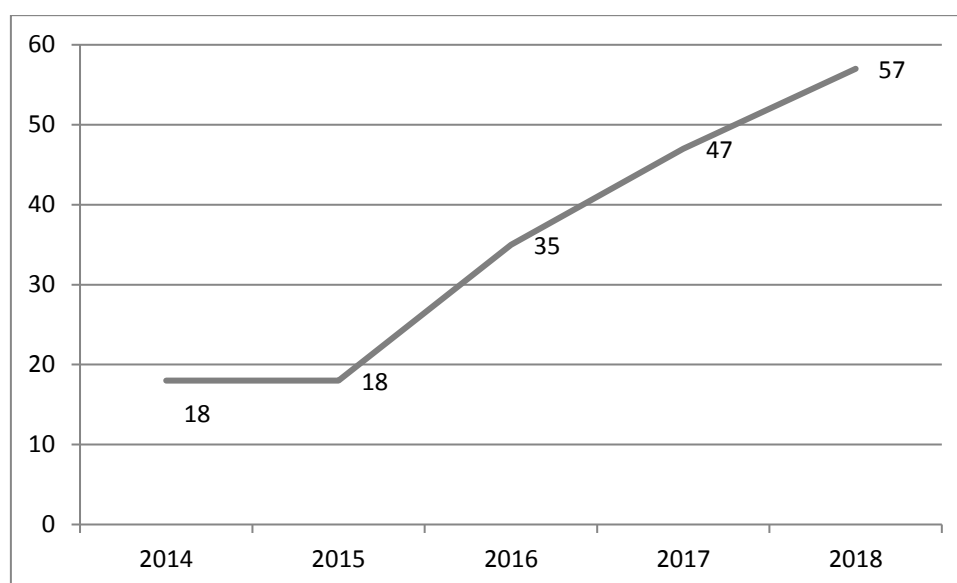


Рисунок. Количество статей с 2014 г.

Была выявлена частота упоминаний терминов, которые употребляются вместе с ключевым термин-концептом «геймификация». Наиболее часто встречающиеся термины помимо ключевого термин-концепта: Games, Learning, Motivation, Study, Elements, Social, Use, Research, Person, Time, Behavior, Students (таблица). Терминологический ландшафт может быть полезен для более тщательного изучения того, как термины используются в разных контекстах, установлении перспективных и востребованных направлений исследований.

Результаты обзора показали, что, начиная с 2014 года, научные публикации имели обзорный характер и описывали повышение эффективности процесса [4, 5]. В 2015–2016 гг. термин рассматривался в аспекте применения элементов геймификации в обучении. В 2017 году появился тренд использования геймификации в таких контекстах, как здравоохранение и управление персоналом. В 2018 году – появился тренд применения геймификации в сопровождении процессов управления городской средой, электронного участия и вовлеченности граждан в процессы принятия решений.

Таблица. Терминограмма по ключевому термину «Gamification»

2014		2015		2016		2017		2018	
Термин	Частота	Термин	Частота	Термин	Частота	Термин	Частота	Термин	Частота
Social	786	Learning	518	Learning	925	Research	1291	Design	1699
Use	331	Games	504	Use	786	Games	1092	Study	1522
Research	270	Use	358	Students	734	Learning	1064	Research	1516
Games	268	Elements	285	Social	683	Study	1023	Students	1367
Learning	263	Based	278	Games	663	Students	829	Games	1317

2014		2015		2016		2017		2018	
Information	228	Motivation	277	Design	599	Based	813	Learning	1250
Based	221	Students	268	Health	561	Participants	802	Motivation	1119
Time	209	Design	267	Research	541	Design	798	Based	1110
Study	197	Social	263	Model	517	Motivation	785	Use	1092

Терминограмма, содержащая 9 из наиболее часто встречающихся слов (Топ 9) для массива документов (корзины), была получена с использованием средств для текстового анализа. Терминограмма показывает количество наиболее часто встречающихся понятий в тексте.

Таким образом, к настоящему времени исследователи смещают свое внимание с мотивации и элементов, и сосредотачиваются на применении геймификации в конкретных сферах деятельности.

Подход заключался в расширении таксономии, определении отношений между терминами, определении синонимичных, иерархических и ассоциативных связей между терминами. Такой подход позволяет определить звенья предметных областей с опорными (ключевыми) понятиями и терминами.

Данное исследование являлось частью работы с Санкт-Петербургским межрегиональным ресурсным центром (МРЦ). Результаты исследования будут использованы при разработке электронного курса «Эффективная презентация» с применением элементов геймификации для обучения государственных служащих в МРЦ.

Литература

1. Deterding S., Dixon D., Khalad R., Nacke, L. From game design elements to gamefulness: Defining 'gamification' // Proceedings of the 15th International Academic MindTrek Conference: Envisioning Future Media Environments. – 2011. – P. 9–15.
2. Vanolo A. Cities and the politics of Gamification // Cities. – 2018. – V. 74. – P. 320–326.
3. Mueller C., Klein U., Hof A. An easy-to-use spatial simulation for urban planning in smaller municipalities // Computers, Environment and Urban Systems. – 2018. – V. 71. – P. 109–119.
4. Kampker A., Deutskens C., Deutschmann K., Maue A., Haunreiter A. Increasing Ramp-up Performance By Implementing the Gamification Approach // Procedia CIRP. – 2014. – V. 20. – P. 74–80.
5. Romano M., Díaz P., Aedo I. Emergency Management and Smart Cities: Civic Engagement Through Gamification // Third International Conference, ISCRAM-med. – 2016. – P. 3–14.

**Тарасенко Тарас Денисович**

Год рождения: 1996

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С4155Направление подготовки: 09.04.03 – Прикладная информатика

e-mail: kipesh96@gmail.com

**Видясов Евгений Юрьевич**Университет ИТМО, институт дизайна и урбанистики,
аспирантНаправление подготовки: 41.06.01 – Политические науки
и регионоведение

e-mail: vidyasov@lawexp.com

УДК 004.9**ИССЛЕДОВАНИЕ ПРОЦЕССОВ СБОРА ДАННЫХ О ПАЦИЕНТАХ В СИСТЕМЕ
ЗДРАВООХРАНЕНИЯ САНКТ-ПЕТЕРБУРГА И ЛЕНИНГРАДСКОЙ ОБЛАСТИ****Тарасенко Т.Д.****Научный руководитель – аспирант Видясов Е.Ю.**

В работе приведены результаты исследования регламентов информационного взаимодействия с единой государственной информационной системой здравоохранения и опыта информатизации пациента для выявления данных о пациенте, по которым следует осуществлять его уведомление.

Ключевые слова: электронное здравоохранение, обмен данными лабораторных исследований, медицинская информационная система, медицинская организация, электронная медицинская карта, система информирования пациента.

Основным инструментом информатизации здравоохранения в России является Единая государственная информационная система здравоохранения (ЕГИСЗ). Как написано в концепции разработки данной системы, утвержденной в апреле 2011 г., ЕГИСЗ – «совокупность информационно-технологических и технических средств, обеспечивающих информационную поддержку методического и организационного обеспечения деятельности участников системы здравоохранения». Особые места в ЕГИСЗ занимают подсистемы интегрированной электронной медицинской карты (ИЭМК) и обмена данными лабораторных исследований (ОДЛИ) [1–5].

Подсистема ИЭМК – место хранения интегрированных данных о пациентах и их обращениях в медицинские организации (МО). Сервис обеспечивает сбор, хранение и передачу гражданам информации о результатах предоставления им медицинской помощи в учреждениях здравоохранения. Также сервис обменивается данными с другими сотрудничающими сторонами.

Подсистема ОДЛИ предназначена для автоматизации процессов передачи направлений на лабораторные исследования в лабораторную информационную систему (ЛИС) и дальнейшего получения результатов. Информационное обеспечение процесса осуществляют: медицинская информационная система (МИС) медицинской организации (как источник информации о назначении и получатель результатов исследования), ЛИС (как получатель информации о назначении и источник результатов исследований) и сервис ОДЛИ (как информационная шина, обеспечивающая информационный обмен, и как региональное хранилище информации по лабораторным исследованиям).

В силу законодательства медицинские организации имеют статус операторов персональных данных, так как они осуществляют сбор, систематизацию, накопление, хранение, уточнение, обновление, изменение, распространение и уничтожение информации о пациентах. У МО возникают обязанности касательно работы с полученными данными. Перед получением медицинской помощи пациент дает ряд согласий в письменной форме, за исключением экстренных случаев госпитализации, когда пациенту необходима защита жизни или здоровья и получение согласия пациента невозможно:

- согласие на обработку персональных данных;
- информированное добровольное согласие на медицинское вмешательство;
- согласие на получении информации по каналам связи (СМС-рассылка или сообщение на электронную почту).

В случае если состояние пациента не позволяет ему выразить свою волю, а медицинское вмешательство неотложно, вопрос о его проведении решают консилиум или лечащий врач. Согласно Федеральному закону от 27.07.2006 № 152-ФЗ «О персональных данных», в случае если пациент изъявляет желание получать уведомления через СМС- рассылку или электронную почту, ему необходимо дать согласие на получении информации по каналам связи.

В период стремительного роста электронных коммуникаций, значительного увеличения интенсивности обмена информацией, уменьшения времени для принятия решений и их реализацией перед МО ставится задача соответствия развития технологий информационного взаимодействия непрерывным изменениям внешней и внутренней среды. В связи с этим разрабатываются системы информирования, главной целью которых является совершенствование коммуникации МО и пациента за счет формирования информационного пространства, отвечающего достижению целей лечебного учреждения в целом.

Система MEDVOX – это интеллектуальная платформа автоматической обработки входящих и исходящих звонков для МО. Робот-оператор берет на себя функции операторов контакт-центров и сотрудников регистратур лечебных учреждений. Данная система успешно применяется в нескольких медицинских учреждениях страны, например, в детской поликлинике № 1 Петропавловска-Камчатского и в Красноярской краевой клинической больнице. Компанией-разработчиком данной системы является компания S2S Next, она разрабатывает инновационные продукты в области речевых технологий, а также внедряет речевые интерфейсы в уже существующие клиентские сервисы.

Система доприемного информирования пациента (СДИП) включает две формы: непосредственную и дистанционную. Непосредственно перед приемом проводится предварительное информирование пациента о планируемых вопросах, которые будут заданы врачом на предстоящем приеме. Это позволяет пациенту: обдумать ответы, повысив их информативность и качество; сократить длительность непосредственного опроса пациента; увеличить время на консультирование; повысить приверженность к лечению. Дистанционная форма СДИП (СМС, e-mail и др.) предлагает пациенту учесть предпочтения в одежде, облегчающие процедуру осмотра, напоминает о необходимости выполнения подготовительных мероприятий перед диагностическими процедурами и др. По мнению авторов, такой подход позволит улучшить работу амбулаторно-поликлинического звена, минимизировать психологический дискомфорт у пациента и повысит его удовлетворенность работой поликлиники.

На сегодняшний день существует широкий спектр МИС, которые в той или иной степени осуществляют информирование пациентов. Так МИС «Ариадна» осуществляет уведомление пациентов о готовности результатов лабораторных исследований и отправка PDF-файлов с результатами исследования на электронную почту. В МИС компании «Smart Delta System» существует модуль «SMS информирования пациентов», который осуществляет информирование пациента о записи на прием, последующее напоминание о приеме и имеет ряд дополнительных возможностей, например, поздравление с днем рождения, напоминание

о наличии у пациента неоплаченного долга, о сгорании бонусного сертификата и т.д. В целом рынок МИС многообразен и не структурирован, в основном разработки МИС ведутся под конкретные потребности заказчика.

В ходе данного исследования были выявлены основные потребности информирования пациентов лечебными учреждениями:

- напоминание пациентам о запланированном приеме и запрос подтверждения;
- отмена или перенос приема;
- информирование пациентов;
- опросы пациентов с целью оценки качества обслуживания;
- уведомление с целью контроля регулярности принятия лекарственных средств;
- уведомление пациентов о необходимости сдачи анализов;
- уведомление определенных возрастных групп с целью проведения профилактических мероприятий;
- анкетирование;
- уведомление о готовности лабораторного исследования;
- рассылка результатов лабораторных исследований;
- уведомление уязвимых групп граждан, с целью контроля их жизненных показателей;
- уведомление о необходимости прохождения вакцинации.

В ходе изучения систем информирования пациентов были выявлены основные приоритетные данные для уведомления пациентов. Была изучена структура передаваемых данных из МИС в ЕГИСЗ. Данные, передаваемые в ЕГИСЗ на сегодняшний день: фамилия (врач, пациент, представитель), имя (врач, пациент, представитель), отчество (врач, пациент, представитель), идентификационный номер специальности врача, дата начала приема, дата окончания приема, идентификатор медицинской организации, код заболевания по МКБ-10, идентификатор должности медицинского работника, наименование специальности, наименование должности (не используется при передаче данных в сервис), дата вакцинации, код иммунобиологического препарата, наименование лечения, описание схемы лечения, дата начала лечения, дата окончания лечения, время проведения инструментального исследования, код типа инструментального исследования, дата окончания оказания услуги, дата начала оказания услуги, код услуги по региональной номенклатуре, наименование услуги, сведения о препарате, курсовая доза, суточная доза, количество дней лечения, дата выдачи рецепта /назначения на препарат, тип выдачи препарата, наименование препарата, лекарственная форма препарата, способ введения медикамента, разовая доза, код лекарственного средства, диагноз, информация о направлении, идентификатор учреждения, из которого осуществляется направление, идентификатор учреждения, куда направлен пациент, состав заявки на лабораторное исследование, сведения о биоматериале, дата и время сбора биоматериала, код лаборатории, статус выполнения заявки.

В таблице показаны необходимые данные для уведомления пациентов, которые необходимо передавать в ЕГИСЗ.

Таблица. Ключевые данные, необходимые для уведомления

Потребность информирования	Необходимые данные
Напоминание пациентам о запланированном приеме и запрос подтверждения	ФИО пациента, ФИО врача, специальность врача, дата приема, телефон/e-mail
Отмена или перенос приема	ФИО пациента, ФИО врача, специальность врача, дата приема, телефон/e-mail
Информирование пациентов	Содержание информирования формируется вручную, но при этом существует отбор пациентов по критериям; телефон/e-mail

Потребность информирования	Необходимые данные
Опросы пациентов с целью оценки качества обслуживания	Содержание опросов формируется вручную/автоматически для каждой специальности врача; телефон/e-mail
Уведомление с целью контроля регулярности принятия лекарственных средств	ФИО пациента, ФИО врача, медикаменты, дата принятия; телефон/e-mail
Уведомление пациентов о необходимости сдачи анализов	ФИО пациента, ФИО врача, дата назначенной сдачи; телефон/e-mail
Уведомление определенных возрастных групп с целью проведения профилактических мероприятий	Содержание уведомления формируется вручную, но при этом существует отбор пациентов по критериям; телефон/e-mail
Анкетирование	Содержание анкеты формируется вручную, но при этом существует отбор пациентов по критериям; телефон/e-mail
Уведомление о готовности лабораторного исследования	ФИО пациента, статус лаб. исследования, дата отправки
Рассылка результатов лабораторных исследований	ФИО пациента, статус лаб. исследования, дата отправки
Уведомление уязвимых групп граждан, с целью контроля их жизненных показателей	ФИО пациента, диагноз; содержание уведомления формируется вручную, но при этом существует отбор пациентов по критериям
Уведомление о необходимости прохождения вакцинации	ФИО пациента, телефон/e-mail, дата вакцинации, код вакцинации

В ходе сравнения передаваемых и необходимых для передачи данных в сервисы ЕГИСЗ было выявлено, что данные о телефонах и электронных почтах пациентов не передаются, однако вся информация о случаях обслуживания и о лабораторных исследованиях успешно передается. Для возможности реализации сервиса уведомлений на единой площадке ЕГИСЗ необходимо добавить в обязательные поля для передачи из МИС:

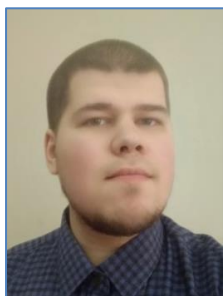
- телефон для оповещений;
- электронная почта для оповещений.

В работе представлены промежуточные результаты исследования в рамках магистерской диссертации «Разработка технического задания модуля уведомления пациентов для государственной информационной площадки здравоохранения Санкт-Петербурга и Ленинградской области». На следующем этапе исследования будут предложены варианты реализации уведомления пациентов с технической и с психологической точки зрения. Будут рассмотрены варианты реализации отправки сообщений и их содержательная часть.

Литература

- Агаларова Л.С. Совершенствование технологии работы участковых терапевтов городских поликлиник в условиях модернизации здравоохранения // Вестник ДГМА. – 2015. – № 1(14). – С. 42–46.
- Булыгина Н.В., Исаенкова Е.А. Опыт автоматизации процесса информационного сопровождения застрахованных лиц при оказании медицинской помощи // Научно-практический журнал. – 2018. – Т. 21. – № 1. – С. 84–88.
- Зарубина Т.В., Швырев С.Л., Соловьев В.Г., Раузина С.Е., Родионов В.С., Пензин О.В., Сурин М.Ю. Интегрированная электронная медицинская карта: состояние дел и перспективы // Врач и информационные технологии. – 2016. – № 2. – С. 35–44.

4. Зингерман Б.В., Шкловский-Корди Н.Е., Карп В.П., Воробьев А.И. Интегрированная электронная медицинская карта: задачи и проблемы // Врач и информационные технологии. – 2015. – № 1. – С. 24–34.
5. Концепция создания Единой государственной информационной системы в сфере здравоохранения. Утверждена Министром здравоохранения и социального развития Российской Федерации. Приказ № 364 от 28 апреля 2011 г.



Тропников Александр Сергеевич

Год рождения: 1996

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С41551

Направление подготовки: 09.04.03 – Прикладная информатика

e-mail: petyabegen@mail.ru



Низомутдинов Борис Абдуллохонович

Год рождения: 1990

Университет ИТМО, институт дизайна и урбанистики,
ведущий аналитик

e-mail: boris-wels@yandex.ru



Углова Анна Борисовна

Год рождения: 1989

Российский государственный педагогический университет
им. А.И. Герцена, кафедра психологии профессиональной
деятельности, к.психолог.н.

e-mail: anna.uglova@list.ru



Чугунов Андрей Владимирович

Год рождения: 1956

Университет ИТМО, институт дизайна и урбанистики,
к.полит.н., доцент

e-mail: chugunov@corp.ifmo.ru

УДК 004.023

ИССЛЕДОВАНИЕ МЕТОДОВ ВЫЯВЛЕНИЯ ПСИХОЭМОЦИОНАЛЬНЫХ ОСОБЕННОСТЕЙ ИНФОРМАЦИОННОГО ОБРАЗА У ПОЛЬЗОВАТЕЛЕЙ СОЦИАЛЬНЫХ СЕТЕЙ

Тропников А.С. (Университет ИТМО), **Низомутдинов Б.А.** (Университет ИТМО),
Углова А.Б. (Российский государственный педагогический университет им. А.И. Герцена)
Научный руководитель – к.полит.н. Чугунов А.В. (Университет ИТМО)

В работе представлен краткий обзор возможностей использования автоматизированных средств анализа психоэмоциональных особенностей пользователей социальных сетей и социально-психологических характеристик их информационного образа, рассмотрены существующие инструментарины обработки изобразительных и текстовых данных.

Ключевые слова: социальные сети, психоэмоциональные особенности, социально-психологические характеристики, автоматизированные средства.

Аудитория социальных сетей расширяется с каждым годом. За период с 2008 по 2018 год число пользователей одной только социальной сети Facebook увеличилось со 100 млн до 2,3 млрд. Более 42% жителей Земли активно использует одну или несколько социальных

сетей. Подобная аудитория предоставляет в сети огромное количество общедоступных данных, которые активно используются в задачах политического и коммерческого маркетинга: для выявления общественного мнения, определения потребительских предпочтений, создания таргетинговой рекламы и др.

Для специалистов различных областей виртуальная среда выступает площадкой, предоставляющей новые возможности в исследовании социально-демографических и психологических качеств человека. Основу этих возможностей задает специфика интернет-коммуникации и текстовая продукция пользователей Интернета и социальных сетей. Ее использование позволяет создать структурную модель виртуального жизненного пространства современного человека, выявить характеристики современной виртуальной коммуникации, определить факторы, опосредующие межкультурные и общественно-политические противоречия, выявить уязвимые социальные группы и т.д. Использование автоматизированных систем поможет облегчить процесс получения подобной информации и обеспечить определение корреляции с другими факторами и признаками различных объектов исследований.

На первом этапе исследования были проанализированы основные подходы к использованию автоматизированных информационных систем в социально-психологических исследованиях для выявления психоэмоциональных особенностей пользователей на основе анализа их сетевых профилей. Также был проведен сравнительный анализ этих систем и тестирование специализированного программного обеспечения.

Анализируя данные пользователей в социальных сетях, можно выявить целостную структуру текстовых и визуальных компонентов сетевого облика человека, которые составляют информационный образ пользователя и лежат в основе создания его сетевой личности. Сетевая личность представляется авторами как устойчивая система социально значимых черт, представленных в виртуальном пространстве и описывающих индивида как члена того или иного общества. Таким образом, используя общедоступные данные из социальных сетей, можно получить объемный социокультурный образ пользователя, который может быть соотнесен с реальным социально-психологическим портретом человека.

В ряде современных научных исследований исследуется возможность разработки и использования новых методов для выявления социально-психологических характеристик пользователей на основе анализа информационного контента, которые были бы также эффективны, как и традиционные психодиагностические исследования.

Большинство научных работ специализируется на различных группах данных, получаемых из социальных сетей. Некоторые работы используют только статистические данные – количества постов, фотографий, друзей и т.п. Другие делают акцент на анализе фотоизображений: их технических характеристиках, цветовой гамме, демонстрируемых эмоциях, количестве людей. Также есть серия научных статей, посвященных анализу текстовых данных, таких как комментарии, посты, записи, статусы и т.п. Текстовые данные анализируются с целью выявления эмоциональной окраски, частоты использования частей речи, ключевых слов, предмета и объекта разговора.

Следует отметить работу Михала Косинского [1], профессора Стэнфордского университета и бывшего заместителя директора Кембриджского психометрического центра. Он собирал данные пользователей Facebook через приложение, позволяющие определить психологический портрет человека. Собрав с помощью такого приложения данные более миллиона пользователей, Косинский выдвинул теорию, согласно которой, количество «лайков» способно рассказать о пользовательских предпочтениях и характере больше, чем любая другая информация. Так, по его мнению, достаточно всего 68 «лайков», чтобы с большой долей вероятности определить сексуальную ориентацию, цвет кожи и уровень интеллекта пользователя.

В исследованиях Ж. Фарнадила и Л. Хана [2] использовались методы опорных векторов, регрессионного анализа и k -ближайших соседей для поиска взаимосвязи между рядом

психоэмоциональных особенностей, выявленных у испытуемых ранее с помощью тестирования, и различными статистическими данными, полученными в ходе анализа профилей в социальных сетях. По результатам этих исследований был выдвинут ряд гипотез, предполагающих наличие взаимосвязи между, например, эмоциональной стабильностью человека и частотой обновления статуса в социальных сетях.

В работе Х.А. Швартца [3] использовался метод подсчета и анализа количества слов и частей речи в постах пользователей социальных сетей, и сравнения полученных данных с предварительно проведенным тестированием и анкетированием. В результате были представлены гипотезы, предполагающие наличие взаимосвязей между возрастом человека и частотой употребления наречий, полом и средним объемом комментариев.

Научная статья С. Росса и соавторов [4] посвящена поиску взаимосвязей между личностными чертами характера и отношением к социальным сетям при помощи проведения ряда тестирований, специализированных на выявлении доминирующей поведенческой модели из «Большой пятерки» (экстраверсия, доброжелательность, добросовестность, нейротизм, открытость опыту) и анкетирования, выявляющего отношения человека к социальным сетям.

Следует отметить публикации по проекту РФФИ «Методика выделения (аудио, видео и текстовых) признаков, характеризующих «социальный портрет» человека и психоэмоциональные особенности его информационного образа», в которых коллектив авторов [5] рассматривает формирование описания информационного образа пользователя социальной сети Вконтакте с учетом определения его психологической характеристики. В частности, описан эксперимент по определению типа темперамента на основе проанализированных с помощью разработанного программного обеспечения данных пользовательских профилей, а также, по результатам проведенного эксперимента представлены выявленные взаимосвязи между темпераментом пользователя и различными анализируемыми данными его профиля.

В ходе анализа научной литературы был выявлен набор гипотез – возможных (предполагаемых) взаимосвязей между типами обрабатываемой информацией, полученной из изображений в социальных сетях и различными характеристиками пользователя (табл. 1).

Таблица 1. Типы извлекаемой информации

Тип информации	Предполагаемая взаимосвязь
Цвет изображения	Преобладающий цвет имеет определение влияния на человеческий характер: красный – стимулятивное, синий – защищающее, зеленый – успокаивающее
Интенсивность цвета	Насыщенный цвет облегчает восприятие изображения для человека и усиливает его естественность
Пропорция динамических и статистических линий	Преобладание динамических линий на фотографии говорит об общей активности снимка и происходящих на нем действий. Статические линии используются при демонстрации монолитных, спокойных объектов, которые предпочитают снимать интроверты
Количество лиц	Большое количество лиц на фотографии свидетельствует об общей социальности человека и его активности в социуме
Угол лица	Применение неестественных поз, использование фото в профиль и острых углов расположения лица может говорить о закрытости и неуверенности человека
Эмоции	Проявляемые эмоции человеком на аватаре могут быть соотнесены с личностными чертами исследуемого человека

Опираясь на ранее проведенные исследования, становится возможным определение принципов разработки метода автоматизированного определения социально-

психологического портрета пользователя социальных сетей. На первом этапе планировалось создание психодиагностического комплекса и проведение тестирования для анализа социально-психологических характеристик пользователей социальных сетей. На втором этапе проведена автоматизированная выгрузка массива различных данных из профиля социальных сетей испытуемых с помощью средств обработки изображения и текста и расчет корреляционных взаимосвязей. Анализ взаимосвязей, полученных психодиагностических и данных профилей, позволит подойти к созданию машинного алгоритма, способного автоматически обрабатывать профили социальных сетей и определять социально-психологические характеристики информационного образа их пользователей.

Авторами был проанализирован ряд облачных платформ, предоставляющих услуги по обработке текстовых и изобразительных данных. Каждый из этих сервисов обладает различными модулями, технологией анализа и точностью. Их сравнительный анализ представлен в табл. 2.

Таблица 2. Сравнительная таблица облачных сервисов анализа текстовых и изобразительных данных

Название платформы	Microsoft Azure	Google Cloud Platform	Amazon Web Services	ISPRAS API	IQBuzz
Разработчик	Microsoft	Google	Amazon	ИСП РАН	Айкубаз
Модуль анализа изображений	+	+	+	—	—
Распознавание объектов	+	+	+	—	—
Создание тэгов	+	+	—	—	—
Определение композиции	—	+	+	—	—
Создание описания изображения	+	—	—	—	—
Определение технических характеристик	+	+	—	—	—
Распознавание лиц	+	+	+	—	—
Распознавание эмоций	+	+	+	—	—
Распознавание головных уборов и деталей лица	+	—	—	—	—
Модуль анализа текста	+	+	—	+	+
Наличие русского словаря	+	—	—	+	+
Распознавание тональности	+	+	—	+	+
Определение темы	+	—	—	+	+
Синтаксический анализ	—	—	—	+	—
Морфологический анализ	—	—	—	+	—
Определение сущностей	—	—	—	+	+

В ходе данного исследования было предложено использовать сервис Microsoft Azure для машинного анализа изображений, а сервис ISPRAS API в качестве инструмента анализа текстовых данных.

Полученные в ходе автоматизированного анализа социально-психологические характеристики информационного образа могут использоваться в самых различных направлениях, например, таргетинговой рекламе, использующей пользовательские данные для вывода рекламы конкретным потенциально-заинтересованным покупателям.

Психоэмоциональными характеристиками информационных образов также пользуются различные HR-департаменты в ряде фирм. Сотрудники этих департаментов анализируют профили социальных сетей своих потенциальных работников на момент наличия потенциальных асоциальных отклонений или наклонностей.

Подобным анализом занимаются и банковские службы при рассмотрении заявок на выдачу кредита. В ходе этого анализа выясняется благонадежность клиентов и риски стабильности выплат по кредиту.

Социально-психологические характеристики информационного образа также можно использовать при определении психологически нестабильных или социально-уязвимых людей, которым потенциально требуется профессиональная психологическая помощь. Особенно остро эта проблема стоит с подростковой аудиторией.

На следующем этапе исследования предполагается использовать отобранный инструментарий для проведения ряда экспериментов по выявлению взаимосвязей между личностными характеристиками и обработанными данными из профилей социальных сетей.

В сотрудничестве с коллегами – психологами из Российского государственного педагогического университета им. А.И. Герцена планируется провести ряд пилотных исследований. В ходе исследований будут использованы традиционные методы анализа и выявления социально-психологических характеристик пользователей.

Полученные данные о психоэмоциональных особенностях испытуемых планируется использовать в качестве сравнительных данных для последующего создания таблицы корреляций между психоэмоциональными особенностями пользователя и массивом данных, получаемых при анализе его профиля в социальных сетях.

Выбрав профили испытуемых, и выгрузив всю необходимую информацию, можно будет произвести обработку отсортированных данных с помощью ранее описанных сервисов. В ходе обработки данных и анализа выходной информации будет проведено отсеивание неподходящих или неточных сервисов.

Используя обработанную информацию, можно будет соотнести закономерности и повторяющиеся взаимосвязи с уже полученными данными в ходе ранее проведенного опроса и создать таблицу корреляции. На основе данной таблицы будет разработан алгоритм, позволяющий автоматически обрабатывать массивы выгружаемых данных и определять предполагаемые психоэмоциональные особенности пользователя.

Литература

1. Kosinski M., Stilwel D., Graepel T. Private traits and attributes are predictable from digital records of human behavior // *Proc. the National Academy of Science of the United State of America*. – 2013. – № 110. – P. 5802–5805.
2. Han L., Qiu L. Sharing emotion on Facebook: network size, density, and individual motivation // *Extended Abstracts on Human Factors in Computing Systems*. – 2012. – P. 2573–2578.
3. Schwartz H.A., Eichstaedt J.C., Kern M.L., Dziurzynski L., Ramones S.M., Agrawal M., Shah A., Kosinski M., Stillwell D., Seligman M.E.P., Ungar L.H. Personality, Gender, and Age in the Language of Social Media: The Open-Vocabulary Approach [Электронный ресурс]. – Режим доступа: <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC3783449/> (дата обращения: 06.01.2019).
4. Ross C., Orr E.S., Sisic M., Arseneault J.M., Simmering M.G., Orr R.R. Personality and motivations associated with facebook use // *Computers in Human Behavior*. – 2009. – V. 25. – № 2. – P. 578–586.
5. Крылова О.С., Власов Д.А., Шишков В.В., Алымов А.С., Ишин И.А., Колесников И.Е., Петров А.И. Описание информационного образа пользователя социальной сети с учетом его психологической характеристики // *International Journal of Open Information Technologies*. – 2018. – V. 6. – № 4. – P. 24–37.

**Беген Петр Николаевич**

Год рождения: 1996

Университет ИТМО, институт дизайна и урбанистики,
студент группы № С41551Направление подготовки: 09.04.03 – Прикладная информатика

e-mail: petyabegen@mail.ru

**Чугунов Андрей Владимирович**

Год рождения: 1956

Университет ИТМО, институт дизайна и урбанистики,
к.полит.н., доцент

e-mail: chugunov@corp.ifmo.ru

УДК 004.8**ИССЛЕДОВАНИЕ ВОЗМОЖНОСТЕЙ ПРИМЕНЕНИЯ ИСКУССТВЕННОГО
ИНТЕЛЛЕКТА В СФЕРЕ ЦИФРОВОГО ГОСУДАРСТВЕННОГО УПРАВЛЕНИЯ****Беген П.Н., Чугунов А.В.****Научный руководитель – к.полит.н. Чугунов А.В.**

Работа выполнена в рамках темы НИР «Исследование возможностей применения машинного обучения для автоматизации тематической классификации сообщений на портале «Наш Санкт-Петербург».

В работе представлен краткий обзор исследования возможностей применения искусственного интеллекта и методов машинного обучения в сфере цифрового государственного управления в ряде развитых стран, рассмотрены основные тенденции и перспективы развития технологий искусственного интеллекта, которые закреплены в программах развития и стратегических планах государств.

Ключевые слова: искусственный интеллект, машинное обучение, цифровое государственное управление, цифровые технологии, информационные технологии.

В Российской Федерации (РФ) государственная политика в сфере развития искусственного интеллекта (ИИ), применительно к задачам госуправления, впервые была отражена в утвержденной указом Президента РФ от 9 мая 2017 года «Стратегии развития информационного общества в Российской Федерации на 2017–2030 гг.». Согласно Стратегии применение в органах государственной власти новых цифровых технологий, направленных на повышение качества госуправления, является одной из основных задач применения информационно-коммуникационных технологий для развития социальной сферы, системы государственного управления, взаимодействия граждан и государства.

В данном документе [1] ИИ был отнесен к «прорывным» цифровым технологиям, т.е. таким технологиям, которые выводят на рынок новое ценностное предложение по сравнению с предыдущими, так как внедрение таких инструментов требует существенных изменений бизнес-модели, непосредственно касающейся человеческой деятельности, и технической модели функционирования объекта.

В 2018 году начался процесс очередного пересмотра приоритетов в сфере внедрения информационно-коммуникационных технологий в государственное управление и другие

сферы общественной жизни. Стимулом для этих организационных изменений стало заявление президента России В.В. Путина о необходимости разработки новой госпрограммы, которая получила название «Цифровая экономика».

Одним из основных направлений программы является «Цифровое государственное управление», ориентированное на внедрение цифровых технологий и платформенных решений в сферах государственного управления и оказания государственных услуг, в том числе в интересах населения и субъектов малого и среднего предпринимательства, включая индивидуальных предпринимателей.

Тем самым можно видеть, что в рамках приоритетов новой программы созданы организационные (и финансово-поддерживаемые) предпосылки старта проектной деятельности в сфере государственного управления с использованием такой «прорывной» технологии как ИИ.

Искусственный интеллект (Artificial Intelligence) – это область информатики, занимающаяся разработкой интеллектуальных компьютерных систем, обладающих возможностями, которые традиционно связаны с процессами человеческого разума, такими как возможность обучения, способность рассуждать и выстраивать логику, понимать язык, решать проблемы разными способами и др.

ИИ включает в себя целый комплекс стремительно развивающихся технологий и процессов. Одним из первоочередных и активно развивающихся направлений ИИ является машинное обучение (Machine Learning), которое определяется как класс методов ИИ, применяемых в областях больших данных (Big Data) и интернета вещей (IoT), изучающих и разрабатывающих алгоритмы автоматизированного распознавания образов и извлечения знаний из массивного объема данных, а также основанных на обучении аппаратных систем, на основе полученных данных, генерации прогнозных значений и рекомендаций.

Методы машинного обучения обладают широким спектром применения в различных отраслях деятельности и имеют высокую практическую значимость. Так, машинное обучение активно используется в государственном секторе, логистике и транспорте, медицине, финансах и торговле, маркетинге и продажах, страховании, военном деле и др., для оптимизации бизнес-процессов, распознавания лиц в городском потоке, прогнозировании результатов и т.п.

В сфере государственного управления ИИ помогает снизить стоимость операций, увеличить скорость работы процессов, позволяет сосредоточить больше имеющихся ресурсов на приоритетные задачи, предоставить новые методы коммуникации между обществом и органами власти (например, создание и развитие электронных порталов) и т.п. Методы машинного обучения активно используются для интеллектуального анализа данных с целью повышения уровня своей эффективности работы, значительной экономии денежных средств, обеспечении информационной безопасности.

Авторы исследования из аналитического центра «Deloitte Center for Government Insights» [2] приводят примеры использования машинного обучения для ускорения процесса работы с документами, упрощения электронного документооборота, повышения точности прогнозов для увеличения эффективности и качества стратегических документов, принимаемых государственными органами власти. А также приводят решения базовых элементарных задач, как например проблемы с регистрацией и авторизацией пользователей, частично разгружая нагрузку техподдержки ИТ-отдела.

Интересен мировой опыт применения ИИ и методов машинного обучения в сфере государственного управления, экономике и обеспечения безопасности страны. Например, полиция Нидерландов в мае 2018 года стала использовать ИИ и методы машинного обучения для расследования сложных преступлений, предварительно оцифровав более 1500 отчетов нераскрытых дел для последующего анализа записей системой машинного обучения в целях поиска достоверных доказательств совершенного противодействия. Применение ИИ, в

частности распознавание лиц, используется в департаменте полиции Детройта для расследования преступлений и быстрой поимки преступника [3].

Определенных успехов в развитии ИИ добилось Правительство США. В октябре 2016 года администрация президента США Барака Обамы подготовило отчет об использовании ИИ и методов машинного обучения в процессах государственной политики и управления. Согласно отчету был произведен запуск ИИ-систем на федеральном уровне для способствования внедрения инноваций в целях лучшего обслуживания органов власти страны на основе распределенных агентств. Данные агентства должны были обеспечить совместную работу для развития и распространения стандартов и лучших практик использования ИИ в различных операциях и процессах.

Администрация президента США Дональда Трампа продолжило развитие работ в данном направлении, предоставляя дополнительные рекомендации агентствам, определяя машинное обучение и ИИ в качестве приоритетов дальнейших научных исследований. В 2018 году на основе майского доклада Белого дома был сформирован специальный комитет по ИИ для улучшения координации федеральных усилий, связанных с ИИ, и обеспечения дальнейшего лидерства США в данной области.

Исследователи из Национальной лаборатории Ок-Ридж (Oak Ridge National Laboratory) подготовили отчет, в котором представлены результаты по оптимизации методов машинного обучения, применяемых Федеральным агентством по чрезвычайным ситуациям для поиска зданий и других техногенных объектов, уничтоженных потоком лавы на Гавайях. Еще одним примером является разработанный учеными алгоритм машинного обучения, испытанный в городе Канзас, который способен отслеживать и прогнозировать образование ям на дорожном покрытии.

Федеральные органы власти, такие как Белый дом и портал службы гражданства и иммиграции в США (U.S. Citizenship and Immigration Services), поручили ИТ-компаниям разработать персональных виртуальных помощников (или чат-ботов) – компьютерных программ, самостоятельно запускающих автоматические задания и использующихся для онлайн-системы ответов на базовые вопросы пользователей, с целью поддержания с ними диалога. Алгоритмы данных программ базируются на распознавании текста, основанном на методах машинного обучения и ИИ. Вопрос использования таких чат-ботов активно рассматривается в отчете агентства «GovLoop».

В России вопросам безопасности развития систем ИИ уделяется определенное внимание, например, исследуются проблемы возможного появления своеобразных ловушек на пути развития ИИ, и предлагаются подходы и методы их преодоления, для обеспечения более эффективной безопасности страны. Так, многие российские банки, такие как «Сбербанк», «Банк Хоум Кредит» и др., начинают активно использовать ИИ для обнаружения и предотвращения мошеннических операций.

В работе [4] представлены перспективы развития ИИ в России путем формирования и совершенствования электронного правительства, которое призвано гарантировать слаженность деятельности отдельных компонентов этого механизма для решения национальных стратегических задач, заложенных конституцией, документами стратегического и территориального планирования. С учетом возрастания динамики развития рынка и ростом неопределенности политической либо экономической ситуации, а также сложности реализуемых властями функций и услуг, растет аналитический потенциал электронного правительства, а также увеличиваются возможности решать более сложные задачи.

В Великобритании был создан единый портал открытых государственных данных (<https://data.gov.uk>), а к 2010 г. в большинстве стран мира было реализовано электронное предоставление государственных и муниципальных услуг гражданам и бизнесу.

Также правительство Великобритании поддерживает применение и развитие ИИ в бизнесе и государственном управлении в соответствии со Стратегией цифровой экономики

2015–2018 гг., утвержденной в начале 2015 г. государственной организацией «Innovate UK». Государство поддерживает научно-исследовательские работы и опытно-конструкторские работы по цифровым технологиям в учебных и исследовательских организациях в соответствии с Цифровой стратегией Великобритании (так называемой «Белой книгой») – политическим документом, опубликованным в марте 2017 года Министерством культуры, средств массовой информации и спорта.

В США также имеется подобный политический документ «Национальный стратегический план исследований и разработок в области искусственного интеллекта». Данный документ является стратегией Штатов по разработке и применению ИИ в государственном управлении и экономике страны, и содержит стратегический план исследований и разработок, финансируемых Федеральным правительством в области ИИ. Тенденции отражены в упомянутом ранее стратегическом докладе «Готовясь к будущему искусственного интеллекта», подготовленном подкомитетом Конгресса США по машинному обучению и ИИ (MLAI), где представлено текущее состояние ИИ, его существующие и потенциальные приложения и вопросы, которые прогресс ИИ ставит перед обществом и государственной политикой. Также в докладе были представлены 23 рекомендации конкретных дальнейших действий со стороны федеральных агентств и других участников, заинтересованных во внедрении ИИ в различные сферы деятельности.

В заключение следует отметить, что в России в 2019 году планируется сделать первые шаги по стандартизации требований к системам ИИ. В частности, ведется разработка национального стандарта «Информационные технологии. Искусственный интеллект. Общие положения», который установит единые требования к архитектуре систем ИИ и их функциональным характеристикам. В этом стандарте также будут представлены типовые примеры применения технологий ИИ [5].

В результате работы были исследованы примеры использования ИИ в сфере цифрового государственного управления, экономики и безопасности. Определены основные направления и перспективы развития технологий ИИ и методов машинного обучения в государственной сфере, закрепленных в программах развития и стратегических планах ряда стран.

Представленное обзорное исследование является первым этапом научно-исследовательской работы, предполагающей применение методов и технологий ИИ для автоматизации работы с системами классификаторов и справочников государственных информационных систем. Обозначенные тенденции позволяют сформировать видение развития ИИ в таких смежных областях, входящих в сферу государственного управления, как развитие электронного правительства, формирование электронных порталов и предоставление государственных услуг гражданам страны.

На следующем этапе исследования предполагается проведение анализа и отбора основных методов машинного обучения, а также их дальнейшее использование в разработке автоматической системы классификации корпусов текста на различные категории.

В сотрудничестве с компанией «Нетрика», базирующейся на ИТ-экспертизе по управлению проектами в государственном секторе, планируется разработка и последующий запуск информационной системы автоматической классификации сообщений граждан на портале «Наш Санкт-Петербург», где каждый житель Санкт-Петербурга может направить сообщения о проблемах, связанных с жилищно-коммунальным хозяйством и благоустройством города, состоянием дорог и тротуаров, незаконными объектами строительства и торговли и т.д.

Разработанная информационная система, позволит осуществить автоматизированную классификацию сообщений на основе анализа текста аннотации, которую формирует пользователь портала. Приоритетной задачей данной разработки является упрощение и автоматизация процесса подачи сообщения гражданами на портале в целях повышения эффективности и удобства работы с подачей сообщения, а также минимизации риска отклонения сообщения гражданина по причине неправильно выбранной тематики проблемы.

Литература

1. Васин С.Г. Искусственный интеллект в управлении государством // Управление. – 2017. – № 3(17). – С. 5–10.
2. Eggers W.D., Schatsky D., Viechnicki P. AI-augmented government. Using cognitive technologies to redesign public sector work [Электронный ресурс]. – Режим доступа: <https://www2.deloitte.com/content/dam/Deloitte/uk/Documents/public-sector/deloitte-uk-dup-AI-in-government.pdf> (дата обращения: 18.12.2018).
3. London D. Powering AI for government [Электронный ресурс]. – Режим доступа: <https://www.cio.com/article/3279546/artificial-intelligence/powering-ai-for-government.html> (дата обращения: 20.01.2019).
4. Соколов И.А., Дрожжинов В.И., Райков А.Н., Куприяновский В.П., Намиот Д.Е., Сухомлин В.А. Искусственный интеллект как стратегический инструмент экономического развития страны и совершенствования ее государственного управления. Часть 2. Перспективы применения искусственного интеллекта в России для государственного управления // International Journal of Open Information Technologies. – 2017. – № 5(9). – С. 77–101.
5. В России началась разработка стандартов в области квантовых коммуникаций, искусственного интеллекта и умного города [Электронный ресурс]. – Режим доступа: <http://d-russia.ru/v-rossii-nachalas-razrabotka-standartov-v-oblasti-kvantovyh-kommunikatsij-iskusstvennogo-intellekta-i-umnogo-goroda.html> (дата обращения: 19.01.2019).

СОДЕРЖАНИЕ

Направление «ИНФОРМАЦИОННЫЕ ТЕХНОЛОГИИ ПРОГРАММИРОВАНИЕ».....	4
Большаков Н.А. Применение искусственных нейронных сетей при детектировании акустических событий.....	5
Виноградова Т.Б. Сравнительный анализ методов определения эмоциональной окраски коротких текстов	8
Зено Б. Обучение распутанного представления с использованием генеративных состязательных сетей	11
Казиева Н. Функциональные схемы алгоритмов штрихового кодирования для задач лицевой биометрии	14
Лизунова И.А. анализ тональности текстов с использованием нейросетевых моделей	18
Макаров Р.Н. Предсказание грамматических и просодических характеристик текста для систем TTS	22
Мамаев Н.К. Анализ функциональных возможностей диалоговой программной библиотеки DeepPavlov	25
Мирзаянова С.В. Использование нейронных сетей при отборе признаков в задачах машинного обучения	28
Мироненко А.А. Анализ алгоритмов трекинга лиц с корректировкой ошибок трекинга.....	31
Рюмин Д.А., Аксёнов А.А., Карпов А.А. Автоматическое обнаружение лиц для человеко-машинного взаимодействия	33
Столбов М.Б., Куан Ч.Т. Адаптивный алгоритм MVDR для двухэлементной микрофонной решетки	38
Фетисова Т.А. (Санкт-Петербургский государственный университет), Черных Г.А. (Санкт-Петербургский государственный университет). Разработка диалогового ассистента для получения текстовой информации из структурированных баз знаний	42
Ховричев М.А., Черных И.А., Нугманова А.А., Булушева А.В., Кабаров В.И. Применение контекстно-зависимых методов для извлечения синонимов и концептов из неструктурированных текстов на естественном языке.....	46
Черных Г.А., Латынина В.О. Статистический анализ эмоциональной динамики литературных текстов на русском языке	51
Направление «ИНФОКОММУНИКАЦИОННЫЕ ТЕХНОЛОГИИ»	55
Алексеев Н.Б. Технологии и методы поисковой оптимизации	56
Амосова А.Ф. Обзор современных платформ для разработки систем бизнес-аналитики	59
Бойко Н.К. Применение BI-системы для автоматизации формирования оперативной экономической отчетности	62
Говоров А.И., Слипень Е.В. Проблематика синтаксического анализа SQL-кода.....	65
Гордиенко А.А. Анализ современных дополнений и расширений в браузерах.....	69
Долгачев Р.Ф. Анализ процесса классификации изображения с помощью машинного обучения	72
Загряжская Н.И. Методы поиска синонимов при обработке естественного языка.....	75

Зенцова Е.С., Иванов С.Е., Лавская Л.В. Технологии и средства разработки медицинской информационно-аналитической системы	79
Карлов Б.А., Ананченко И.В. Исследование технологии слияния данных	84
Кузьмина О.Ю. Применение декомпозиции при планировании финансовых показателей.....	87
Купина А.П. Обзор проблем интеграции существующих мессенджеров в единую систему обработки обращений.....	91
Минич Н.С. Технологии и методы автоматизированного тестирования кода при разработке распределенных систем	95
Пац К.М. Сравнение методов расчета свободной энергии Гиббса для катализаторов на основе одновалентного марганца.....	98
Попов Н.С. Анализ фрактальных свойств сетевого трафика в инфокоммуникационных системах.....	102
Пряничников А.А. Возможности взаимодействия с Wi-Fi сетями в мобильной операционной системе iOS	105
Синицына А.А. Анализ методик оценки экономической безопасности предприятия	108
Смирнов А.Е. Анализ современных решений для построения сети интернета вещей.....	111
Соломонова Ю.В., Хлопотов М.В. Разработка рекомендательной системы с учетом тематических профилей пользователей социальной сети	116
Терентьев А.В. Анализ существующих технологий беспилотных автомобилей	121
Тимоненков Д.Ю. Опыт кастомизации CRM-системы на примере разработки виджета импорта/экспорта цифровой воронки.....	124
Хмеленко А.Ю. Системный анализ и оптимизация бизнес-процессов частной медицинской клиники	127
Шаркова Е.И. Исследование методов отображения виртуальных графических моделей в системах дополненной реальности	132
Щаникова К.Е., Говоров А.И. Оценка необходимой площади лесных насаждений для обеспечения заданного региона кислородом с учетом потребностей населения.....	135
Ямщиков Д.В. Блокчейн как технология транзакций: возможности развития в сфере транснациональных платежных систем	138
Направление «ДИЗАЙН И УРБАНИСТИКА»	142
Бочкарев К.О. Обнаружение незаконной рекламы на фотографиях фасадов зданий.....	143
Вишератин А.А., Мухина К.Д., Струков А.М. Разработка интеллектуальной платформы мониторинга состояния городской среды на основе разнородных данных...	146
Вычужанин П.В., Калюжная А.В. Робастная калибровка параметров численной модели ветрового волнения SWAN.....	151
Деева И.Ю. Многомасштабное моделирование цифрового образа человека в киберпространстве	156
Деревицкий И.В. Предсказательное моделирование хронического инсулиннезависимого диабета	160
Дунаенко С.С., Смирнов Е.В. Применение методов агентного моделирования для повышения безопасности движения пешеходов	163
Елховская Л.О. Персонализированное предсказательное моделирование жизненных показателей пациентов с хроническими заболеваниями в условиях ограниченности наблюдений	168
Киселева В.О., Видясова Л.А. Определение доверия в контексте использования информационных технологий: аналитический обзор	173

Петрушина Т.Д. Автоматизация контрольно-надзорной деятельности в области государственного тарифного регулирования в Санкт-Петербурге	177
Рубцова В.С. Особенности перехода автоматизированной информационной системы электронного социального регистра населения на взаимодействие посредством системы межведомственного электронного взаимодействия версии 3	181
Смирнова П.В. Анализ портала «Наш город Москва»	185
Смирнова П.В., Кононова О.В. Технологии геймификации: обзор научных публикациях	189
Тарасенко Т.Д. Исследование процессов сбора данных о пациентах в системе здравоохранения Санкт-Петербурга и Ленинградской области.....	192
Тропников А.С. (Университет ИТМО), Низомутдинов Б.А. (Университет ИТМО), Углова А.Б. (Российский государственный педагогический университет им. А.И. Герцена). Исследование методов выявления психоэмоциональных особенностей информационного образа у пользователей социальных сетей	197
Беген П.Н., Чугунов А.В. Исследование возможностей применения искусственного интеллекта в сфере цифрового государственного управления	202

АЛЬМАНАХ НАУЧНЫХ РАБОТ МОЛОДЫХ УЧЕНЫХ УНИВЕРСИТЕТА ИТМО Том 3

В авторской редакции

Редакционно-издательский отдел Университета ИТМО

Дизайн обложки

Н.А. Потехина

Зав. РИО

Н.Ф. Гусарова

Редактор

Л.Н. Точилина

Подписано к печати 10.10.2019

Заказ № 4238

Тираж 100 экз.